

经济统计分析方法及预测

——附实用计算机程序

罗积玉 邢 瑛 编著

清华大学出版社

内 容 简 介

本书从理论上和应用上系统阐述了多变量统计分析方法在经济统计分析和预测中的应用。内容丰富,全书包括十二章,有多元线性回归、多元逐步回归、岭回归、多因变量对多自变量回归、主成分分析、因子分析、系统聚类分析、动态聚类分析、模糊聚类分析、多组判别分析、逐步判别分析、对应分析、典型相关分析、时间序列分析、马尔科夫模型分析等。

本书编写特点注重于实际应用,列举了大量实例,并附有完整的微机实用程序。该书可供从事于计量经济研究、经济管理、计划、统计、人口分析、人口研究等工作人员参考用,也可作高等院校财经、管理、统计等专业高年级学生和研究生教学参考用书,还可以作为从事多元统计分析应用和计算机应用的工作人员、研究人员参考用书。

经济统计分析方法及预测

——附实用计算机程序

罗积玉 邢 瑛 编著

☆

清华大学出版社出版

(北京 清华园)

北京京辉印刷厂印刷

新华书店北京发行所发行

☆

开本, 787×1092 1/16 印张, $24\frac{3}{4}$ 字数, 6.4千字

1987年8月第1版 1987年8月第1次印刷

印数, 00001—10000

统一书号, 15235·301 定价, 5.05元

目 录

第一章 绪论	1
第二章 数学准备	5
第一节 基本概念	5
一、自变量与因变量	5
二、总体和样本	6
三、样本均值与方差	6
四、参数与统计量	6
五、参数估计	10
第二节 矩阵代数	10
一、矩阵的概念	10
二、分块矩阵	12
三、正交矩阵与正交变换	14
四、线性方程组求解	15
五、特征值与特征向量	17
第三节 多元变量及其表征数	21
一、多元变量的总体	21
二、多元变量的表征数	21
三、多元正态分布	22
第三章 回归分析	24
第一节 多元线性回归分析	24
一、多元线性回归的数学模型	25
二、参数 β 的最小二乘估计	26
三、回归方程及回归系数的显著性检验	29
四、复相关系数与偏相关系数	32
五、回归模型的变量子集合的选择	34
六、柯布-道格拉斯生产函数的回归估计	35
七、应用实例	37
八、多元线性回归分析计算程序	40
第二节 逐步回归分析	46
一、“最优”回归方程的选择	46
二、引入变量与剔除变量的依据	47
三、逐步回归的矩阵变换算法	48
四、逐步回归的计算步骤	51
五、应用实例	55

六、逐步回归分析计算程序	37
第三节 趋势面分析	64
一、趋势面分析的一般数学方法	64
二、趋势面分析的适度检验	71
三、趋势面分析计算程序	73
第四节 岭回归分析	80
一、岭回归的方法原理	81
二、一个简单选择K值的方法	81
三、计算步骤	82
四、岭回归分析计算程序	86
第五节 多个自变量对多个因变量的回归分析	94
一、多对多回归分析的数学模型	94
二、最小二乘估计	95
三、多对多回归分析的计算步骤	97
四、多对多回归分析计算程序	102
第四章 主成分分析	109
第一节 主分量的几何解释	109
第二节 主分量的导出	110
第三节 主成分分析的计算过程	112
第四节 主成分分析的应用	116
第五节 主成分分析计算程序	120
第五章 因子分析	125
第一节 方法原理	125
第二节 因子载荷的统计意义	130
第三节 方差最大正交旋转	131
第四节 Promax 斜旋转	134
第五节 因子得分	138
第六节 应用实例	140
第七节 因子分析计算程序	144
第六章 聚类分析	156
第一节 聚类分析的方法及变量类型	156
第二节 系统聚类分析	157
一、数据变换处理	158
二、相似系数和距离	162
三、聚类方法	168
四、系统聚类法的统一公式	175
五、剩余信息的剔除	176
六、谱系分类的确定	176
七、应用实例	177

八、系统聚类分析计算程序	181
第三节 动态聚类分析	192
一、选择凝聚点	193
二、初始分类	193
三、分类函数	194
四、主要计算步骤	195
五、动态聚类分析计算程序	197
第四节 模糊聚类分析	201
一、模糊集	201
二、模糊分类关系	293
三、模糊聚类	205
四、应用实例	208
五、模糊聚类分析计算程序	213
第七章 判别分析	219
第一节 两组判别分析	219
一、方法原理	219
二、判别函数的检验	222
三、计算步骤	222
四、两组判别分析计算程序	225
第二节 多组判别分析	232
一、方法和原理	232
二、计算步骤	235
三、应用实例	239
四、多组判别分析计算程序	241
第三节 逐步判别分析	248
一、逐步判别分析的理论基础	248
二、引入和剔除变量计算	250
三、建立判别函数	252
四、逐步判别计算步骤	253
五、应用实例	257
六、逐步判别分析计算程序	260
第八章 对应分析	270
第一节 方法原理	270
一、对应于 R 型分析情况	270
二、对于 Q 型分析情况	274
第二节 计算步骤	276
第三节 应用实例	280
第四节 计算程序	284
第九章 典型相关分析	290

第一节	方法原理	290
一、	典型变量及典型相关系数的概念	290
二、	典型变量及典型相关系数的求法	291
三、	典型变量的性质	293
四、	各原始变量与典型变量的相关系数	294
第二节	计算步骤	294
第三节	应用实例	300
第四节	典型相关分析计算程序	303
第十章	时间序列分析	314
第一节	时间序列的基本概念	314
第二节	时间序列分析预测	315
一、	时间序列的趋势分析	315
二、	时间序列的周期性预测	318
三、	平稳时间序列预测	321
第三节	时间序列分析计算程序	328
第十一章	马尔科夫概型分析	335
第一节	马尔科夫过程	335
第二节	离散型的平稳马尔科夫链	336
一、	转移概率和转移概率矩阵	336
二、	切普曼-柯尔莫各洛夫方程	337
三、	绝对概率	338
四、	遍历性	341
第三节	马尔科夫链的转移概率矩阵的估计	343
第四节	转移概率矩阵 U 的谱分解	344
第五节	置换分析	346
一、	左置换	346
二、	右置换	347
三、	互置换	347
第六节	马尔科夫性质的检验	347
第七节	应用实例	348
第八节	马尔科夫概型分析计算程序	352
第十二章	综合应用实例分析	367
附录		379
参考文献		387

第一章 绪 论

国民经济是一个多质多变量的立体网络体系，国民经济的内涵及其构成和发展不是盲目的，杂乱无章的，而是具有整体性、相关性、有序性等规律。因此可以对社会主义经济过程进行定量分析，从量上研究国民经济的各种结构、发展速度、比例关系等具体量的联系及其变动规律性。

一、经济现象是质和量的统一

社会经济现象是质和量的统一。经济现象的质和量的特征共同存在于社会经济运动的过程中，各种经济现象总是通过质变和量变这两种形式来显示它们的运动特征。

大体上说来，经济学中存在着三类问题：

1. 质的问题。这类问题不直接表现为数量，对它们的认识主要依靠直观推理和逻辑分析。政治经济学就是属于以研究质的问题为主的，如生产资料所有制；生产中人与人的关系；产品的分配形式等就属于这种问题。但有时对这类问题辅之以定量分析也是必要和可能的，如就某一经济系统来说，它的目的、职能和生产性质等问题，很难以数量关系表现出来，对这样一些显示出质的规律性的问题，借助于某些基本的数学原理进行分析，比如将最优策略方法用于该经济系统，它可能对该经济的个别子系统是最优的策略，而对整个系统则不一定是最优的决策，从而从理论上论证了局部和全局的关系。

2. 量的问题。这是属于计量经济学研究的范畴。计量经济学是一门从数量上研究物质资料生产、交换、分配、消费等经济关系和经济活动规律及其应用的科学。计量经济学是把经济理论和社会经济实际两者结合起来，研究各种经济量如社会需求、供给、生产、投资、消费等方面的各种数值关系和数值规律，建立、估计和检验各类经济模型。它的目的有三个：一是结构分析，即应用计量经济模型对经济变量之间的关系作出定量的量度；二是预测未来，即应用这些模型对某些经济变量的未来数值作出比较准确的预测；三是政策评价，即不同政策方案的选择，这样使我们以计划经济为主的国家，能根据定量分析结果制定出科学的计划、合理的政策，使之恰当地安排未来的经济活动。

3. 质-量问题。在这类问题中，量的规律性和质的规律性交织在一起，而且其数量关系常不能直接表现出来，往往需要用间接方法来解决。如果教育经济学中没有定量化研究，人们对教育的经济效益就不能得到具体的认识。例如对北京八个厂矿1800多名性别、年龄、工龄、健康状况相同的人进行调查，经过统计分析处理，得到结论是，在技术复杂部门，只具有小学文化程度的工人完成定额假定为100，则初中、高中和大学文化程度的工人完成定额分别为106、113、118，经过数理统计计算得出初步估计：在我国国民收入增长额中，大约有25%是由教育投资所带来的收益。

既然任何经济现象都是质和量的统一，那么我们在研究它们时，就必须在研究其质的规律性的同时，研究其量的规律性。质是量的基础，而一定的量是质存在的必要条件，定性分析是定量分析的基础，而定量分析则使定性分析更加准确可靠，使人们对质的规律性的认识

更加深入全面，从而能深刻揭示经济现象的本质和它们相互之间的关系。

对经济学中的定量化研究必须要有正确的认识，即不能过分注重于经济现象的质而忽视了量，因而对经济学中运用数学方法处理持否定态度；而又不能过分注重经济现象的量而忽视了质，因而夸大了数学在经济学中的作用，把研究经济学，陶醉于一些数学概念和公式之中。这两种倾向都会在理论上和实践上带来有害的后果。

在经济学中应用数学方法对经济数据进行统计分析处理，使经济学研究定量化、科学化、模型化，这是经济学自身发展的必然要求。但是定量化研究作为经济学研究的一个重要侧面，应该在马克思主义理论指导下，正确地揭示社会经济现象及其相互之间存在的数量关系以及质和量的关系，从而为国民经济的科学论证提供定性和定量的理论依据。

二、经济研究的特点及处理方法

按照马克思主义的观点，一定社会的经济现象，不仅和它有关的其它现象一起构成一个普遍联系着的统一体，而且在它的内部也存在着许多彼此密切联系，互相影响（互相制约或推动）的因素。在一定条件下，一些因素推动着（或制约）另外一些与之联系的因素发生变化。这种状况表明在经济现象的内部和外部联系中存在着一定的因果关系，因而人们才能够利用这种因果关系来进行工农业生产和各种经济活动，来制定各种有关的经济政策以引导经济活动朝着人们预期的方向和规模发展。因此因果性是经济活动所固有的基本特点之一。

在经济研究中所取得的统计数据都是采取统计抽样方法获得，被抽取的变量符合随机变量的要求。例如对某城市的职工进行工资待遇抽样调查，被抽取的每一个职工的工资有高有低，他们的工资并不是互相影响的，即某甲的工资高低并不受某乙或某丙的工资高低的影响，工资额这样的变量就是随机变量。经济领域内的随机现象大量存在，因而随机性也是经济现象所具有的基本特征。

综上所述，对于经济现象可以从两个途径，即随机性和因果性途径来探求和发现它的规律性。

为了探求经济现象的内部和外部联系中所固有的随机性和因果性特征，通常采用数理统计分析方法，进行统计分析，并用图表形式或以数学方程表现的函数关系来实现。计量经济学主要是运用数学方程式组成的经济数学模型来模拟经济领域中的数量联系和关系。相对于错综复杂的经济关系，经济数学模型只是现实经济关系的抽象，但它却存在着近似模拟这些现实关系的可能性。无论是用数学方程还是使用矩阵代数等数学方法作为研究手段，都要求根据现实统计资料或抽样调查资料，运用数理统计分析方法来估算参数、简化结构，估计和调整随机误差，探求和发现经济活动的规律性和因果关系。

由于社会主义经济建设需要解决大量的实际问题，在电子计算机，特别是微型计算机日益发展和广泛使用的今天，在经济研究、统计领域、人口研究等经济管理部门都需要应用统计分析方法来解决经济现象所固有的随机性和因果性两方面的问题。而多变量统计分析是处理多因素、多指标特征问题的最有成效的实用统计分析方法。

在实际问题中，许多研究工作者在研究工作中往往需要确定如何选择和使用某一给定的统计分析方法去求解问题，采用多变量统计分析进行数据处理、建立宏观或微观经济数学模型，需要研究如下几方面的问题：

1. 简化数据结构、选择变量子集合。

当给定一组抽样数据或统计数据之后，所选取的变量或指标是否一定需要这么多：是否可以找到一个基本结构；或者在众多变量所构成的集合中，可以找到一个子集合，而且是最佳变量子集合，该子集合所包含的变量能够反映总体的结构。这种简化结构的处理对我们研究经济现象之间的规律、构造模型、简化抽样调查方案等方面都有着极其重要的意义。这种处理方法的目的是“从树木看森林”，抓住主要矛盾，舍弃次要因素，从而获得最佳效果。

采用主成分分析、因子分析、对应分析等方法都可以达到简化结构的目的。

2. 进行数值分类处理、进行分类研究、构造分类模式。

在经济管理的许多研究中，往往需要把性质相近的经济现象归为一类，这样在大量错综复杂的经济特征分别归类以后，便于找到它们内部的规律性，同一类内的经济特征有相同的规律性，不同类经济特征有不同的规律性。过去许多研究都是单因素定性分类处理，比如对人口生育模式函数研究，人们通常按照地域分类，即按平原、丘陵、山区、高原分为四大类，这种分类方法反映不出政治、经济、文化、风俗、民族等社会因素的差异，因而需要采用多变量的数值分类处理。通常用聚类分析、判别分析等方法解决这样的问题。

3. 构造模型和模型外推。

计量经济学的目的在于探索客观经济过程的数量规律，把所要考察的对象描述成能够用数学方程表达的形式，即建立计量经济模型。计量经济模型实质上是数理统计中回归模型的应用和发展。在多元分析应用中有两大类模型，一类是预测模型，另一类是描述性模型。对于预测模型通常采用回归分析方法解决，如果预测量是一个因变量，可应用多元线性回归、多元逐步回归、拟线性回归分析等方法处理；如果预测量是多个因变量，可以应用多因变量对多自变量的三角回归分析处理。对于描述性模型，通常采用聚类分析等方法处理。

4 研究时间序列变化趋势。

在经济活动中，一般都是用时间序列数字来反映经济发展状况，找出影响时间序列的各种因素，测定其对时间序列影响的程度。在西方，时间序列分析主要是调查整个经济活动中商业循环的性质以及影响它的主要原因。由于西方国家有不可克服的基本矛盾，他们的经济发展是波动的，最好的情况称为“顶峰”，最差情况叫“低谷”，从“顶峰”到“低谷”是萧条；从“低谷”到“顶峰”是复甦，这种循环对他们经济影响很大，所以西方国家很重视经济的时间序列分析。

时间序列分析是我们研究经济活动发展趋势的最重要统计方法之一。用时间序列分析方法和马尔科夫概型分析可以处理随时间变化规律的经济活动。

如何选择一个合适的统计分析方法来解决实际问题，图（1-1）给出了他们之间的对应关系。通常一个问题不是孤立地寻求一种统计分析方法来解决，而是综合应用多种统计分析方法来解决。例如为了构造计量经济模型，构造的方法和步骤是：首先根据经济理论和定性分析结果设计理论模型；根据对实际经济活动的观察，抽取相应的统计数据，并对该数据进行初步精炼；然后利用统计分析方法（如相关分析、主成分分析等）研究被抽取变量之间的相关关系，选择最佳变量子集合；在此基础上构造计量经济模型；然后对模型进行优化处理和在实际经济活动中开展应用。图（1-2）给出了构造计量经济模型的方法和过程。

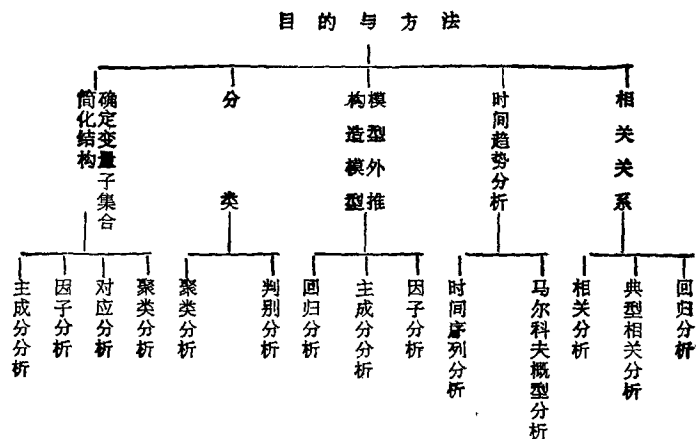


图 1-1 统计分析方法应用说明

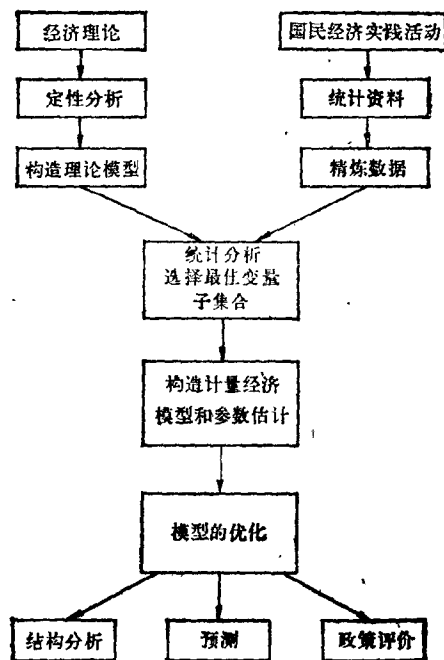


图 1-2 计量经济模型构造方法

第二章 数学准备

第一节 基本概念

多变量统计分析是研究经济领域、人口分析等许多领域的多维随机变量的统计规律。作为多元统计分析的基础，并能更好地理解数理统计学所研究的问题，本节将介绍数理统计中经常用到的一些基本概念和术语。

一、自变量与因变量

1. 自变量：所谓自变量（独立变量）的含义是这些变量可以设置为特定值或在特定值时可以被观测到。如果这些变量为特定值，则它们就是完全确定的数学变量。假若自变量的取值可以观测到，但不能控制，则它们是随机变量（它们可能是、也可能不是统计独立的）。例如粮食产量与气温、降水量、播种面积有关。其中气温、降水量、播种面积为自变量，而播种面积可以设置为特定值，它是一个完全确定的量，气温和降水量的取值可以被观测到，但不能控制，它们是随机变量。

2. 因变量：因变量可看作为随机变量，不过它依赖于一个或多个自变量，或者是它们的函数。上例中粮食产量是因变量，它依赖于气温、降水量和播种面积。

自变量和因变量如何确定，对具体问题要具体分析。根据经济理论，结合实际进行分析，找出事物的主要矛盾，将那些统计资料比较容易获得，变化规律可循的量作为我们需要构造的模型的自变量；相应的与所选定的自变量有相依函数关系的量作为因变量。

二、总体和样本

1. 总体：通常把所要调查研究的事物或现象的全体叫做总体，或称为母体。而组成总体的各个元素称为个体。一个总体中所含个体的数量称为总体的大小或总体容量。例如要研究北京市居民户的收支情况，那么北京市所有居民户的收支情况就是一个总体，而每一个居民户的收支情况就是个体。

总体可分为有限总体与无限总体。例如北京市的居民户是可以计数的，叫有限总体；相反，投掷若干颗骰子，期望得出某些点子，我们可以无限次地投掷，这样的总体称为无限总体。

2. 样本：在总体中抽取出来的一部分个体的集合称为样本。例如，在北京市居民户中随机抽取1500户来进行职工家庭收支调查，这些户就组成一个样本。

为了对总体进行研究，必须对总体进行抽样（从总体中抽取样本的过程）观测，一般要进行多次观测。只有对随机现象进行足够多次的观测，被研究的随机现象的规律性才能呈现出来。每进行一次观测，就得到一次观测值 $(x_1, x_2, \dots, x_p)$ ，这样 $(x_1, x_2, \dots, x_p)$ 就构成 p 维随机向量。

既然要从样本来研究总体，当然希望样本能愈近似地代表总体愈好。但样本不应是选择而来，只能是用随机抽样的方法从总体中抽取样本，只有这样，通过多次观测，才能比较合理地解释总体。数理统计是以随机抽样为基础的，如果抽样不是随机的，则数理统计分析方法是无用的。

三、样本均值与方差

在数理统计中起重要作用的是总体的数字特征。我们研究的目的就是寻找总体的某些特征。在实际问题中，我们无法得到总体的分布函数，也不能对总体的数字特征进行直接计算，而是以样本的一些数字特征来推断总体的特征。

均值是表示数据集中性特征的数值，它只能应用于性质相同的总体中。方差是衡量数据波动的特征，方差越大，波动越大，方差越小，波动越小。

设 $x_1, x_2, \dots, x_N$ 是一组独立的随机样本，则样本均值为

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i \quad (2.1.1)$$

样本方差为

$$s^2 = \frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2 \quad (2.1.2)$$

样本标准差为

$$s = \sqrt{\frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2} \quad (2.1.3)$$

四、参数与统计量

1. 参数：最常见的一元正态分布的密度函数为

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2} \quad (\sigma > 0) \quad (2.1.4)$$

式中的 μ, σ 是可变常数，称之为参数。它们分别代表总体的平均值和方差。对于一个特定的正态总体， μ, σ 为固定值，但对于不同的正态总体， μ, σ 又是可变的。当 μ, σ 变化时，正态分布曲线也发生变化， μ 和 σ 值唯一地确定了正态总体的分布。

2. 统计量：作为总体特征值的参数，事先是未知的，它们只能分别由样本的相应特征值来代替，这种由样本导出的数是由样本确定的，因此它是样本的函数，称之为统计量。而总体的参数是由相应的统计量估计的。

五、参数估计

通过对实际样本的统计计算来估计总体的参数，叫作参数估计。

参数估计分为点估计和区间估计。所谓点估计就是估计研究总体各参数的可能取值。所谓区间估计，就是估计研究总体各种参数取值的可能范围。

1 点估计

估计量 T 是样本 $x_1, x_2, \dots, x_n$ 的函数, 它不可以包含有未知参数, 就是说, T 是一个作估计用的统计量, 是随机向量。当随机抽取一组样本观测值 $(x_1, x_2, \dots, x_n)$ 后, 用 $T(x_1, x_2, \dots, x_n)$ 作为未知参数 θ 的估计值。对不同的样本观测值, 所得的估计值是不相同的。

构造估计量的方法很多, 如矩估计法, 最大似然估计法, 最小二乘估计法等, 在经济学上常用最大似然估计和最小二乘估计。

(1) 最大似然估计

最大似然估计法就是当某一事件出现的概率很大, 则可以推断在某一次观测试验中, 这个事件一定会发生。

假设已知总体是正态分布, 随机抽取一组样本: $x_1, x_2, \dots, x_n$, 由于它们都是从正态分布的总体中抽取的, 所以这些随机变量都遵从相同的正态分布 $N(\mu, \sigma^2)$ 。

对于样本 $x_1, x_2, \dots, x_n$, 其似然函数为

$$\begin{aligned} L &= \prod_{i=1}^n \left[\frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x_i - \mu)^2}{2\sigma^2}} \right] \\ &= \left(\frac{1}{\sigma\sqrt{2\pi}} \right)^n e^{-\sum_{i=1}^n \frac{(x_i - \mu)^2}{2\sigma^2}} \end{aligned} \quad (2.1.5)$$

最大似然估计就是使其似然函数 L 达到最大。对(2.1.5)式取对数, 并令

$$L(\mu, \sigma^2) = -\frac{1}{2\sigma^2} \sum_{i=1}^n (x_i - \mu)^2 - n \ln \sigma - n \ln \sqrt{2\pi}$$

这样可用数学分析中求极大值的方法, 求 μ 和 σ^2 的极大似然估计。即

$$\begin{cases} \frac{\partial L(\mu, \sigma^2)}{\partial \mu} = \frac{1}{\sigma^2} \sum_{i=1}^n (x_i - \mu) = 0 \\ \frac{\partial L(\mu, \sigma^2)}{\partial \sigma^2} = \frac{1}{2\sigma^4} \sum_{i=1}^n (x_i - \mu)^2 - \frac{n}{2\sigma^2} = 0 \end{cases}$$

解之得:

$$\begin{cases} \hat{\mu} = \frac{1}{n} \sum_{i=1}^n x_i = \bar{x} \\ \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = s^2 \end{cases} \quad (2.1.6)$$

由此说明, 对于随机抽取的正态分布总体的样本, 用其均值 $\bar{x}$ 和方差 s^2 去估计总体的均值和方差, 使之概率最大。

用使概率最大求估计量的方法称为最大似然估计法, 用最大似然估计法求出的估计量 $\hat{\mu}$ 、 $\hat{\sigma}^2$ 称为最大似然估计量。

(2) 最小二乘估计

最小二乘估计法就是使拟合误差的平方和达到最小的方法。

设 (y_i, x_i) , $i = 1, 2, \dots, n$ 是满足简单线性模型: $y_i = \alpha + \beta x_i + e_i$ 的 n 个点, α 和 β 的最小二乘估计量是使误差平方和 $\sum_{i=1}^n e_i^2$ 极小化的 α 和 β 的值。令

$$L(\alpha, \beta) = \sum_{i=1}^n e_i^2 = \sum_{i=1}^n (y_i - \alpha - \beta x_i)^2$$

用数学分析中求极值的方法, 解得:

$$\begin{aligned} \hat{\alpha} &= \bar{y} - \hat{\beta} \bar{x} \\ \hat{\beta} &= \frac{\sum_{i=1}^n (y_i - \bar{y})(x_i - \bar{x})}{\sum_{i=1}^n (x_i - \bar{x})^2} \end{aligned} \quad (2.1.7)$$

2. 估计量好坏的标准

对于同一个未知参数, 可以构造出不同的估计量、如何鉴别、比较估计量的好坏, 这就要进一步研究点估计的性质, 以决定估计量的选取和提供求得优良估计的方法。

衡量一个估计的好坏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, 而必须从总体出发, 即在多次观测中, 我们希望好的估计量都集中在真参数值的附近波动。

(1) 无偏性

对任意 n 和 θ , 参数 θ 的估计值 $\hat{\theta}$ 的数学期望为 θ , 即 $E(\hat{\theta}) = \theta$, 满足这种要求的统计量 $\hat{\theta}$ 称为 θ 的无偏估计量。

(2) 一致性

一致性是指样本大小增加时, 估计量更接近总体参数。例如 10000 人总体的平均工资为 μ , 抽取 1000 人估计的平均工资 $\bar{x}$ 接近于 μ , 如果样本容量从 1000 人增加到 5000 人, 则平均工资估计量 $\bar{x}$ 更接近于 μ 。

(3) 有效性

设 $\hat{\theta}$ 及 $\hat{\theta}'$ 都是 θ 的无偏估计量, 如果方差

$$D\hat{\theta} \leq D\hat{\theta}'$$

那么就说 $\hat{\theta}$ 较之 $\hat{\theta}'$ 有效。

即一个有效估计量是具有最小方差 σ^2 的估计量, 对固定的 n , $D\hat{\theta}$ 的值达到最小, 就称 $\hat{\theta}$ 为 θ 的有效估计值。

在数理统计中, 极大似然估计量是有效的、充分的估计量。在正态分布下, μ, σ^2 的极大似然估计量是:

$$\begin{aligned} \hat{\mu} = \bar{x} &= \frac{1}{n} \sum_{i=1}^n x_i \\ \hat{\sigma}^2 = s^2 &= \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 \end{aligned}$$

因为 $E(\hat{\mu}) = E(\bar{x}) = E\left(\frac{1}{n} \sum_{i=1}^n x_i\right) = \frac{1}{n} \sum_{i=1}^n E(x_i)$, 又因为样本 $x_1, x_2, \dots, x_n$ 是独立的

且服从 $N(\mu, \sigma^2)$ 分布的随机变量, 所以有

$$E(\hat{\mu}) = \frac{1}{n} \sum_{i=1}^n \mu = \mu$$

即说明样本均值 $\bar{x}$ 是总体均值 μ 的无偏估计量。

$$\begin{aligned}
\text{同理 } E(s^2) &= E\left(\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2\right) \\
&= \frac{1}{n} E\left\{\sum_{i=1}^n [(x_i - \mu) - (\bar{x} - \mu)]^2\right\} \\
&= \frac{1}{n} E\left\{\sum_{i=1}^n [(x_i - \mu)^2 - 2(x_i - \mu)(\bar{x} - \mu) + (\bar{x} - \mu)^2]\right\} \\
&= \frac{1}{n} E\left\{\sum_{i=1}^n (x_i - \mu)^2 - 2(\bar{x} - \mu)(n\bar{x} - n\mu) + n(\bar{x} - \mu)^2\right\} \\
&= \frac{1}{n} E\left[\sum_{i=1}^n (x_i - \mu)^2 - 2n(\bar{x} - \mu)^2 + n(\bar{x} - \mu)^2\right] \\
&= \frac{1}{n} E\left[\sum_{i=1}^n (x_i - \mu)^2 - n(\bar{x} - \mu)^2\right] \\
&= \frac{1}{n} \sum_{i=1}^n E(x_i - \mu)^2 - E(\bar{x} - \mu)^2
\end{aligned}$$

$$\begin{aligned}
\text{因为 } E(\bar{x} - \mu)^2 &= D(\bar{x}) = D\left(\frac{1}{n} \sum_{i=1}^n x_i\right) = \frac{1}{n^2} D\left(\sum_{i=1}^n x_i\right) \\
&= \frac{1}{n^2} n \cdot \sigma^2 = \frac{\sigma^2}{n}
\end{aligned}$$

$$\text{所以 } E(s^2) = \sigma^2 - \frac{\sigma^2}{n} = \frac{n-1}{n} \sigma^2$$

由此可见 s^2 不是 σ^2 的无偏估计量, 但

$$\hat{\sigma}^2 = \frac{n}{n-1} s^2$$

为 σ^2 的无偏估计量。

3 参数的区间估计

由于样本的随机性和研究对象的变异性, 无论是根据样本进行总体参数估计或是进行某种假设检验, 都不可避免地要发生误差, 为此要研究这个误差在某一范围内的概率问题, 这就是区间估计问题。

如果已知某一参数服从于某种分布, 则可利用样本分布函数估计出这个参数落在某一区间的概率, 这就是区间估计, 这个概率称为置信水平, 这个区间称为置信区间。

例如正态分布均值的置信区间。从正态总体中随机抽取一个容量为 n 的样本, $\bar{x}$ 为样本均值, s^2 为样本方差, $\bar{x}$ 和 s^2 是相互独立的, 则 μ 的置信区间可以运用 $\bar{x}$ 和 s^2 的联合分布来构造。

设统计量 t 为

$$t = \frac{\bar{x} - \mu}{\sigma/\sqrt{n}}$$

如果我们希望找置信水平为 0.975 的置信区间, 可以找到数 a , 使得

$$P\{-a < t < a\} = 0.975$$

这 a 可查正态分布表而得到。

因为 $-a < t < a$

所以 $\bar{x} - a \frac{\sigma}{n} < \mu < \bar{x} + a \frac{\sigma}{n}$

故有

$$P\left\{\bar{x} - a \frac{\sigma}{n} < \mu < \bar{x} + a \frac{\sigma}{n}\right\} = 0.975$$

由此得在置信水平为0.975时, μ 的置信区间为

$$\left(\bar{x} - a \frac{\sigma}{n}, \bar{x} + a \frac{\sigma}{n}\right)$$

查表得 $a = 1.96$ 。

在标准正态分布中, 置信区间的长度为 $2 \cdot a\sigma/n$, 如果要缩小置信区间范围, 提高区间估计的精确度, 唯一的办法是增大样本容量, 只有 n 定得足够大, 置信区间的范围可以取得任意小。

第二节 矩阵代数

矩阵代数是在多元统计分析及计量经济学研究中所必须应用的基本数学方法。本书广泛地使用向量和矩阵的表示法, 为此, 简略介绍如下。

一、矩阵的概念

1. 矩阵和向量的定义

一个矩阵就是许多元素作一长方形的排列。例如

$$X = \begin{bmatrix} 1 & -3 & 5.2 \\ 6 & 2 & 7 \end{bmatrix} \quad (2.2.1)$$

上式是2行3列的矩阵。矩阵通常用大写字母表示。

由 $m \times n$ 个数 a_{ij} ($i = 1, 2, \dots, m; j = 1, 2, \dots, n$)排列成的一个有 m 行、 n 列的矩阵

$$A = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \cdots & \cdots & \cdots & \cdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{bmatrix} \quad (2.2.2)$$

上式称为 $m \times n$ 阶矩阵。行数与列数相同的矩阵称为方阵, 如

$$B = \begin{bmatrix} b_{11} & b_{12} & \cdots & b_{1n} \\ b_{21} & b_{22} & \cdots & b_{2n} \\ \cdots & \cdots & \cdots & \cdots \\ b_{n1} & b_{n2} & \cdots & b_{nn} \end{bmatrix} \quad (2.2.3)$$

方阵中下标重复的元素 $b_{11}, b_{22}, \dots, b_{nn}$ 称为对角线元素。

如果一方阵只有对角线元素不为零, 其余元素均为零, 则称它为对角线矩阵, 如

$$\text{diag}(a_1, a_2, \dots, a_n) = \begin{bmatrix} 3 & 0 & 0 \\ 0 & 2.5 & 0 \\ 0 & 0 & -7 \end{bmatrix} \quad (2.2.4)$$

若一个 n 行、 n 列的对角线矩阵，其对角线元素都等于1。即

$$I = \begin{pmatrix} 1 & 0 & 0 & \cdots & 0 \\ 0 & 1 & 0 & \cdots & 0 \\ \cdots & \cdots & \cdots & \cdots & \cdots \\ 0 & \cdots & \cdots & \cdots & 1 \end{pmatrix} \quad (2.2.5)$$

被称为 n 阶单位矩阵。

如果一个 n 阶方阵中，当 $i > j$ 时所有元素均为零，或当 $i < j$ 时所有元素均为零，那么它被称作三角矩阵，如

$$C = \begin{pmatrix} c_{11} & c_{12} & \cdots & c_{1n} \\ 0 & c_{22} & \cdots & c_{2n} \\ \cdots & \cdots & \cdots & \cdots \\ 0 & \cdots & \cdots & c_{nn} \end{pmatrix}, \quad D = \begin{pmatrix} d_{11} & 0 & \cdots & 0 \\ d_{21} & d_{22} & \cdots & 0 \\ \cdots & \cdots & \cdots & \cdots \\ d_{n1} & d_{n2} & \cdots & d_{nn} \end{pmatrix}$$

一个矩阵的转置矩阵是将行和列彼此交换而得。如矩阵(2.2.2)的转置矩阵为

$$A' = \begin{pmatrix} a_{11} & a_{21} & \cdots & a_{m1} \\ a_{12} & a_{22} & \cdots & a_{m2} \\ \cdots & \cdots & \cdots & \cdots \\ a_{1n} & a_{2n} & \cdots & a_{mn} \end{pmatrix} \quad (2.2.6)$$

如果原来的矩阵和转置矩阵所有对应的元素都完全相同，则称它为对称矩阵。如

$$S = \begin{pmatrix} 5 & -2 & 2.1 \\ -2 & 3.7 & 6 \\ 2.1 & 6 & 9 \end{pmatrix} = S'$$

一个矩阵的所有元素若全为零，则它被称作零矩阵，记为 θ 。

只有一列的矩阵，如

$$E = \begin{pmatrix} e_1 \\ e_2 \\ \vdots \\ e_n \end{pmatrix} \quad (2.2.7)$$

称为列矩阵，也叫作列向量。同样行向量就是只有一行的矩阵，如

$$Z = (z_1, z_2, \cdots, z_n)$$

设 A 是一个 p 阶方阵，如果其两边的圆括号改为直线，把 $p \times p$ 个数据括在一起，就称为方阵 A 的行列式，记为 $|A|$ 。矩阵的行列式是一个数。

2. 逆矩阵

逆矩阵除了有加法、减法、乘法等代数运算外，还有与矩阵乘法相对应的逆运算，且在本书各章中应用得最广泛，为此引进逆矩阵的概念。

设 A 为 n 阶方阵，如果有一个 n 阶方阵 B ，使得

$$AB = BA = I$$

则称 B 是 A 的逆矩阵，记作 A^{-1} 。逆矩阵具有以下性质：

(1) 两方阵 A 、 B 乘积的逆矩阵，等于各因子方阵的逆阵交换次序后的乘积，即

$$(AB)^{-1} = B^{-1}A^{-1}$$