

学校的理想装备

电子图书·学校专集

校园网上的最佳资源

现代地理学中的数学方法



前 言

本人在大学本科学学习数学专业，而到了研究生阶段则改学地理学，毕业后一直在兰州大学地理系从事地理数量方法的教学与科研工作。本书就是笔者结合自己近年来的教学与科研工作，在总结国内外学者有关研究成果的基础上写成的。希望它能对从事地理学、生态学、人口学、经济学、城市科学和农业科学方面的科研人员及高等院校师生有一点参考价值。

本书初稿完成之后，于 1994 年 7 月在呼和浩特市召开的全国高校计量地理学与 GIS 教学研讨会上，向与会代表作了介绍，得到了大家的肯定和鼓励，并根据代表们提出的意见，作了修改、补充。在本书的写作过程中，得到中国科学院院士、兰州大学地理系主任李吉均教授及本人的导师艾南山教授的指导和鼓励。本书的出版，得到了兰州大学教务处及高等教育出版社地理编辑室的大力支持，书中插图由高等教育出版社郝林清绘。谨此，一并致谢！

由于笔者水平所限，本书一定存在不少错误与缺点，敬请读者批评指正。

徐建华

1994 年 10 月于兰州大学

内容提要

本书是关于地理数量方法方面的一部新作。全书十二章：1.绪论；2.统计分析方法；3.线性规划方法；4.多目标规划方法；5.随机型决策方法；6.AHP决策分析方法；7.网络分析方法；8.控制论及其应用；9.模糊数学方法；10.灰色系统方法；11.系统动力学方法；12.投入产出分析方法。本书对于从事地理学、生态学、经济学、人口学、环境学、农业科学、城市科学方面的科研人员以及高等院校师生均有一定的参考价值，亦可作为综合性大学和高等师范院校地理系高年级本科生及研究生的教材或教学参考书。

第一章 绪 论

数学方法，不仅是人们进行数学运算和求解的工具，而且能以严密的逻辑和简洁的形式描述复杂的问题，表达极为丰富的实质性思想。数学方法是现代地理学研究中必不可少的重要方法之一。

第一节 现代地理学中的数学方法

——它的形成和发展

地理学是一门古老的学科，早在我国战国前后和古希腊、古罗马时代就开始萌芽，至今已有 2000 多年的发展历史。纵观地理学的发展史，可划分为三个基本阶段：①古代地理学，是农牧业社会的产物，以地理知识的记载为主体；②近代地理学，是工商业社会的产物，是一种对各种地理现象进行条理化归纳，并对其间的关系作解释性描述的多分支知识体系；③现代地理学，是新的科学技术社会，即信息社会的产物，它把地理环境及其与人类活动的相互关系看作统一的整体，采用定性与定量相结合的方法，规范研究与实证研究并举，以解释各种地理现象的内在机制并预测其未来演变的科学。

地理学，自其产生之日起，就与数学有着不解之缘。在古代，地理学与数学之源泉科学——几何学，几乎都是研究地表的。正象《辞海》关于几何学的解释那样：“古代埃及为兴建尼罗河水利工程，曾经进行过测地工作，它逐渐发展为几何学”。因此，在来自希腊文的西方文字中，几何学有“测地”之意，如其英文为 Geometry，与地理学(Geography)、地貌学(Geomorphology)、地植物学(Geobotany)、地生态学(Geoecology)等术语有着一个共同的前缀 Geo。在古代地理学时期，人们为了测算河流长度、山体高度，计算土地面积，不得不运用几何学原理和方法。古希腊学者艾拉托塞尼(Eratosthenes)测算地球周边，就是运用了几何学原理和方法。在近代地理学时期，经济学中的区位论被移植到地理学中，开了地理学运用分析数学之先河。本世纪 20—30 年代，地理学研究中的统计方法开始萌芽，并开始进行地理要素的统计概括和相关关系探讨。这些事实充分说明，数学方法对于地理学家来说，并不陌生。但是，在古代地理学中，运用数学方法仅仅是为了描写地理事件，地理事实和记载地理知识；在近代地理学中，运用数学方法，又只是局限于对地理现象的解释性描述。而在现代地理学中运用数学方法，则是为了更进一步深入地进行定量化研究，以揭示地理现象发生、发展的内在机制及运动规律，从而为地理系统的预测及优化调控提供科学依据。现代地理学中的数学方法的出现，反映了地理学朝着定量化方向发展的新趋势。这种新趋势就是在地理学研究中，以定量的精确判断来补充定性的文字描述的不足；以抽象的、反映本质的数学模型去刻画具体的、庞杂的各种地理现象；以对过程的模拟和预测来代替对现状的分析和说明；以合理的趋势推导和反馈机制分析代替简单的因果关系分析；以最新的定量化技术革新地理学的传统研究方法。

现代地理学中的数学方法的产生与形成，可以追溯到本世纪 50 年代末期。今天，我们所说的现代地理学中的数学方法，就是在 50 年代末开始的计量运动的基础上进一步发展的产物。

一、现代地理学发展史上的计量运动

近代地理学的发展，曾形成了三种主要学派，即①由赫特纳 (A.Hettner) 首创，哈特向 (R.Hartshorne) 继承和发展了的区域学派；②由洪堡 (Alexander Von Humboldt) 和李特尔 (Karl Ritter) 创建，李希霍芬 (F.Richthofen) 继承和发展，拉采尔 (Friedrich Ratzel) 等代表“决定论”，白兰士 (Paul Vidal de la Blache) 和白吕纳 (J.Brunhes) 等代表“或然论”的人地关系学派；③由施吕特尔 (O.Schlüter) 提出，帕萨格 (S.Passarge)、苏尔 (C.O.Sauer) 等阐发的景观学派。到本世纪 40 年代，由于老的人地关系学派日趋落后，而景观学派的理论体系又尚未成熟，因而区域学派就成了当时地理学的主流学派。该学派的主要观点是，地理学的研究对象是区域，研究目标是描述和解释地球表面区域的差异性；在地理学中不存在法则，地理学只能以区域为单元进行类型研究；专论地理学是地理学研究的起点，区域地理学是地理学研究的终点；区域地理的样板，包括区域内的地质、地形、气候、水文、动植物与人类各要素及其相互关系。在赫、哈二氏的倡导下，经马东 (E.de Martonne)、惠特利西 (D.S.Whittlesey)、詹姆斯 (P.E.James) 等地理学家的努力，在西方着实出现了一个区域地理发展的黄金时代。区域地理范式也由此而变成了传统地理学的科学范式。

但是，自本世纪 50 年代以来，区域学派的观点开始受到质疑。一些学者认为，对于区域的描述冗长、乏味、没有生气；对于许多区域的划分，特别是划分大区域，都是很幼稚的、不成熟的、不科学的，区域研究当属于小范围的研究。向区域范式提出最尖锐、最直接批评的是德籍旅美地理学家舍费尔 (F.K.Schaefer)，1953 年他发表了一篇题为“地理学中的例外论”的文章，抨击了哈特向的地域独特主义观点，即“例外主义”观点。他认为，把区域地理作为专论地理成果的综合是妄自尊大，不切合实际的；在区域地理著作中没有引人注目的深刻见解；地理学应该是解释现象，而不应该是罗列现象。解释现象必须有法则，应该把地理现象看成是法则的实例。地理学的目的应该与其它科学有相似之处：都是追求、探索法则的。

舍费尔等人对区域学派的批评与否定，拉开了现代地理学发展史上的计量运动的帷幕。在舍费尔的学术思想的影响下，从本世纪 50 年代末期开始，首先在美国掀起了建立地理学法则的热潮。然而，究竟怎样建立地理学法则？不同学者从不同的角度作了探索，但一般都是将数学、物理学、社会学、经济学的理论和方法引入地理学，探求地理事物的空间格局，其共同之处在于都是开展地理学定量化研究，建立定量模式。这种定量化研究之热潮，就是所谓的计量运动。

计量运动，主要是由美国地理学家发起的，早期主要集中在几所大学。由于各校所持观点不同，研究方向不同，从而形成了各种不同的学派。其中，主要有如下三大学派：

(1) 衣阿华的经济派。该学派的主要代表人物是舍费尔和麦卡尔蒂 (H.McCarty)。此学派受经济学影响较深，着重探讨经济区位现象间相互内在联系及其组合类型。舍氏深受杜能 (J.H.von Thünen)、廖什 (A.Lösch)、克里斯塔勒 (W.Christaller) 及胡佛 (E.Hoover) 等区位论学者和区域经济学家的影响，他花费了大量的精力去翻译和宣传廖什的《区位经济学》，极力倡导建立地理学法则。麦卡尔蒂于 1954 年出版了《对经济地理理论的探讨》一书，认为生产布局理论解释有两种：其一，为因果解释，但是影响生产布

局的变量如此之多，无法处理，所以这种解释是行不通的；其二，为结合联系的解释，从结合的观点出发，只要发现两种现象常常同时出现，就无需探讨其内在因果关系，而只需探讨现象之间分布的结合律。这一学派尤其重视相关分析与回归分析等统计分析方法在人文地理学中的应用。

(2)威斯康星的统计派。早在 1943 年，该校地理系研究生威弗尔 (J.Weaver)就发表了“论美国大麦生产与气候的关系”一文，他运用相关分析、多元回归分析等方法去鉴定气候参数对大麦产量的影响，并用计算方法进行作物布局规划。后来，罗宾逊 (A.H.Robinson)领导一个研究小组，继续发展统计分析方法。1961 年，该校的社会学家东坎 (O.D.Duncan)和仇佐里 (R.P.Cuzzori)完成了巨著《统计地理学》。该学派以发展和应用统计分析方法为其主要特征。

(3)普林斯顿的社会物理学派。该学派的领袖人物是天文学家司徒瓦特 (J.Q.Stewart)。1950 年，司徒瓦特尝试着把物理学原理应用于社会现象的研究之中，创立了颇具特色的社会物理学派。通过比较研究，司氏发现，在许多社会问题研究中，可以借鉴物理学中已经建立起来的规律、定量模式和研究方法。他成功地借鉴物理学中的万有引力定律研究了人口分布的规律，发表了题为“与人口分布和均衡有关的经验数学法则”的论文。司氏认为，社会量纲与自然量纲是极相似的，具有一致性。他还在普林斯顿大学创建了社会物理学实验室。受此学派影响，引力模型、位势模型、空间相互作用模式得到了许多地理学家，特别是理论地理学家的青睐。

无论从美国还是从全世界来看，现代地理学发展史上的计量运动的兴起，首先要归功于加里森 (William L.Garrison)及其领导的华盛顿小组。加氏是第一个把地理学的理论和方法建立在定量基础上的倡导者和实践者，是第一本《计量地理学》教材的作者。他第一个率先在华盛顿大学举办了地理计量方法研讨班，从推广中心地方论、交通网络理论和统计方法等开始，培养了贝里 (B.J.L.Berry)、邦吉 (W.Bunge)、戴西 (M.F.Dacey)、盖提斯 (A.Getis)、马尔布 (D.F.Marble)、毛里尔 (R.L.Morril)、奈斯丘恩 (J.D.Nystuen)、托布勒 (W.R.Tobler)等现代地理学名家。

促进计量运动的还有美国区域科学协会和瑞典地理学定量化研究的影响。美国区域科学协会是由经济、地理、社会、城市与区域规划、建筑及工程等各个学科的学者组成的，其发起人为艾萨德 (Walter Isard)。该协会组织了大量的学术活动，编辑出版了《区域科学年鉴》，因此，该协会成为美国计量运动的源地之一。瑞典学者哈格斯特朗 (Torsten Hägerstrand)是著名的地理计量学者，早在本世纪 30 年代，哈氏领导的隆德学派就开始了对空间扩散模式的探讨。50 年代，他曾受加里森之邀请到华盛顿大学为地理计量方法研讨班授课。他还组织了美国和瑞典地理学家与克里斯塔勒会面，交流学术思想。哈氏的努力对于促进计量运动的发展和向全世界扩散起到了重要作用。

到了本世纪 60 年代，计量运动不胫而走，在短短几年时间里几乎传遍了整个世界。世界各国地理学家纷纷响应，涌现出一大批著名的学者和学派。如英国，由于受计量运动的影响，出现了以乔莱 (R.J.Chorley)、哈格特 (P.Haggett)和哈威 (D.Harvey)等为代表的剑桥学派，该学派以理论造诣高深而著称。随着计量运动的发展，应运而生了各种组织与学术刊物。1964 年，国际地理学联合会 (IGU)设立了地理学计量方法委员会 (Commission on

Quantitative Methods in Geography)；1967年，英国地理学会设立了地理教学采用模型和计量技术委员会（Standing Committee on the Role of Models and Quantitative Techniques in Geographical Teaching）；1968年，日本成立了计量地理学研究委员会，1973年又改称理论、计量地理学委员会。1963年，英国出版了《地理学计量资料杂志》，1969年，美国出版了《地理分析——国际理论地理学》杂志。我国，由于历史的原因，未能赶上计量运动的“黄金时代”，地理学的定量化进程是从本世纪70年代末、80年代初才开始的，但是其发展速度和势头却是十分喜人的。

二、现代地理学中的数学方法的发展阶段

现代地理学中的数学方法，作为一门新的方法论学科，其历史并不算长，但是发展速度是十分惊人的。自50年代末期开始的计量运动以来，现代地理学中的数学方法已经历了三个发展阶段。

第一阶段，大致从50年代末到60年代末期，是现代地理学中的数学方法发展的初期阶段。其主要特点是把统计学方法引入地理学研究领域，构造一系列统计量来定量地描述地理要素的分布特征，比较普遍地应用各种概率分布函数、平均值、方差、标准差、变异系数等统计特征参数以及简单的两要素间的一元线性回归分析方法。从今天的观点来看，这些方法是比较浅易的。但是，它却给长期以来只是定性描述的地理学带来了可喜的变化。许多过去无法准确确定的概念，如分布中心、区域形状、地理要素分布的集中和离散程度等都有了定量指标；许多地理要素之间的相关关系，可以定量地表示了。这一时期，出现了许多专门探讨和介绍数学方法（主要是数理统计方法）的地理专著，如东坎和仇佐里合著的《统计地理学》（1961）、加里森和马布里合著的《计量地理学》（1967）、金（L.J.King）所著的《地理学统计分析》（1969）等。

第二阶段，包括60年代末期到70年代末期的十年时间，属中期阶段。该阶段的特征是多元统计方法和电子计算机技术在地理学研究中的广泛应用。地理学研究对象的多因素、复杂结构和动态特征都使简单的统计方法无能为力，为此就必须寻求解决复杂的地理问题的有效方法。正是在这一时期，电子计算机的生产已经工业化，使用计算机的方法也从一般人很难掌握的机器语言程序发展到高级算法语言程序。随着计算机科学的这种变化，多元统计方法如雨后春笋般地发展起来，成为数理统计学中特别有生命力的分支之一。过去用手算很难完成的复杂计算问题，运用计算机很快就可以得出结果。以电子计算机技术为手段，许多地理学家熟练地掌握了多元统计方法，具备了分析复杂的地理问题的能力。在自然地理学、经济地理学和人文地理学中，以电子计算机为工具，运用多元统计方法使许多复杂问题得到了相当满意的解决。

第三阶段，从70年代末期开始，是现代地理学中的数学方法走向更加成熟和更加完善的阶段。不但包括了概率论与数理统计方法，还包括了运筹学中的规划方法、决策方法、网络分析方法，以及数学物理方法、模糊数学方法、分维几何学方法、非线性分析方法等，而且也包括了计量经济学中的投入产出分析方法等。更值得一提的是，在这一阶段，地理学中的数学方法的发展与现代系统科学紧密地结合起来了。系统理论、系统分析方法、系统优化方法、系统调控方法等被引进了地理学研究领域。系统科学原理和方法的引入，促进了地理学向着具有更加严密的理论结构和现代化方法的方向发

展，从而使以发展地理学方法论为己任的现代地理学中的数学方法更加明显地具有系统科学的性质与理论性的色彩。同时，电子计算机应用技术的发展，特别是 GIS 技术的成熟，为数学方法在现代地理学中的应用提供了更加先进的技术手段，从而使其应用的范围更加广阔。

第二节 现代地理学中的数学方法

——评价与应用

一、现代地理学中的数学方法的评价

1963年,鲍顿(I.Burton)用“计量革命”一词,对自本世纪50年代末期开始的以数学方法在地理学研究中的应用为内涵的计量运动作了形容,认为,此后将不再是革命了,因为它已经成为现代地理学研究的主流方向之一。不过,这种认识,并未完全统一。因为现代地理学中的数学方法的引入,一方面推动了传统地理学研究方法的变革,另一方面却产生了重数量分析,轻区域、生态研究等问题。由此产生了一场波及整个地理学界的大辩论。以至到了本世纪70年代后期,还有人提出要重新评价计量运动,重新认识地理学中的数学方法。有人认为,数学方法只能用来研究地理要素之间的数量关系及地理事物的分布形态,而不能揭示复杂的地理现象形成的机制。又有人认为,地理学的定量化,其实质就是地理学的科学化、现代化。

随着计量运动的发展,对于现代地理学中的数学方法,产生了三种观点。第一是逆计量运动之潮流,反对地理学定量化研究,认为地理现象,尤其是人文、社会经济地理现象十分复杂,不能用简单的数学方法来解释。持这种观点的地理学者,对数学方法采取拒绝和否定态度。如英国地理学家史密斯(David Smith)和奥格登(Philip Ogden),他们曾这样评价计量革命:“这种所谓的革命,实际上是很保守的,因为它把‘空间’作为地理学研究的基础和实质的化身,同时却忽视了一些社会、经济结构的变化,因而成了故弄玄虚,并把现象当作本质”。有人还把计量运动说成是“数学癖的十年”。甚至还有人大声疾呼:地理学有可能陷入“数学决定论”的危险。第二种观点与前一种针锋相对,推崇地理学定量化,认为数学方法不仅是一种分析技术,而且能够导出普遍性规律,能够解决地理学传统研究方法所不能解决的理论问题。持这种观点的有德国地理学家克里斯塔勒(W.Christaller)、美国地理学家邦吉(W.Bunge)、英国地理学家乔莱(R.Chorley)、哈格特(P.Haggett)等。代表性著作有邦吉的《理论地理学》(1962)、哈格特的《人文地理学的区位分析》(1965)、乔莱和哈格特的《地理模型》(1967)等。此外,在芬兰、日本、加拿大、新西兰、印度和前苏联等国家也出现了一批推崇数学方法的地理学者及其代表性论著。第三种观点是介于“定量化”和“反定量化”之间的“非定量化”的观点。这种观点认为,数学方法只是地理学研究方法之一,它只能用来研究地理要素之间的数量关系及地理事物的空间格局,但是不能用它来描述和解释地理规律,不能导出地理学理论。不过,这种观点并不是固定不变的,它具有较大的摇摆性。当地理学定量化研究取得较大进展时,它便宣扬数学方法,强调数学方法在地理学研究中的重要性;当地理学定量化研究遇到困难,出现问题时,它便否定数学方法,贬低数学方法。持这种观点的地理学者,中国有,外国也有。

笔者认为,正确认识与公正合理地评价数学方法在地理学研究中的地位与作用,不仅对于地理数学方法本身,而且对于整个地理学的健康发展都有着十分重要的意义。

对于现代地理学中的数学方法的评价与认识,笔者提出以下几点看法。

(1)世界上的任何事物都可以用数值来度量。在地理学研究中,一切地

理要素，例如区域的规模、城市的位置、道路的长短、气温的高低、雨量的多少、山高水深、人口增减、物产丰欠等等，均可以用数量来表示。对各种地理要素的分布及其间的相互关系，均可以用数学方法进行定量分析与研究。运用数学方法研究地理现象，可以作出确定性解释和精确预测与判断。在现代地理学研究中运用数学方法，有着传统方法无法比拟的优点。

(2)在现代地理学中，传统方法与数学方法之间并没有不可逾越的鸿沟，传统方法是数学方法的基础，数学方法是传统方法发展的必然结果和重要补充。传统方法与数学方法的区别在于：传统方法研究地理问题的程序为：考察、收集资料→根据已有的概念体系条理化→归纳、概括→建立理论与法则；而数学方法研究地理问题的程序为：观察实践→先期模式→提出假设→对资料进行筛选→建立模式→反复检验→建立理论和法则。传统方法所采用的推理方式以综合归纳为主。而数学方法所采用的推理方式以理论演绎为主，传统方法与数学方法的有机结合，是地理学研究现代化的必不可少的条件。这两种方法在现代地理学中的作用不可相互替代

(3)数学方法，不仅是人们进行数学运算和求解的工具，而且能以严密的逻辑和简洁的形式描述复杂的问题，表达极为丰富的实质性思想。对于现代地理学，数学方法不仅是应用地理学研究中的预测、决策、规划及优化设计的工具，而且也是理论地理学研究中进行逻辑推理和理论演绎的手段。

(4)客观上讲，在地理学研究中，任何方法都有其局限性，数学方法当然也不例外。一方面，对于某些地理问题，目前人们还不知道该用什么样的数学方法去处理，这是外部局限性；另一方面，单纯地用数学方法去分析、研究地理问题，究竟可以达到什么样的深度，这是内部局限性。只有正确地认识这些局限性，并不断地寻求克服它们的途径与措施，才能使地理学中的数学方法得到不断的发展和完善。

(5)现代地理学中的数学方法的形成和发展，离不开电子计算机，它与电子计算机应用技术密切相关。一方面，电子计算机的产生和发展，是现代地理学中的数学方法形成与发展的一个十分重要的条件。另一方面，现代地理学中的数学方法的发展，又为电子计算机技术在地理学中的应用提供了更加广阔的领域。被誉为地理学的第三代语言的现代地理学技术——地理信息系统(GIS)，就是数学方法与电子计算机应用技术在现代地理学研究领域内相互结合、相互渗透的产物。这一技术，可以说是在现代地理学中综合运用数学方法和电子计算机技术的一个成功的典范。

二、数学方法在现代地理学研究中的应用

(一)数学方法应用的一些方面

现代地理学，是一门研究地理环境及其与人类活动之间相互关系的综合性、交叉性学科。它以分布、形态、类型、关系、结构、联系、过程、机制等概念构筑其理论体系，注重的是地理事物的空间格局与地理现象的发生、发展及变化规律，追求的目标是人地系统的优化——即人口、资源、环境与社会经济协调发展。所采用的研究方法，是定性方法与定量方法相结合、综合归纳与理论演绎方法并用、规范与实证研究方法并举。

数学方法，不仅是现代地理学研究中的理论演绎与逻辑推理的工具，而且也是定量分析、模拟运算、预测、决策、规划及优化设计的手段。在现代地理学研究的各个分支领域中，它能被按照不同的要求与方式应用。其应用的方面，主要包括：

(1)分布型分析。这类研究，主要是对地理要素的分布特征及规律进行定量分析。譬如，运用平均值、方差、标准差、变异系数、峰度、偏度等统计量描述地理要素的分布特征；运用概率函数研究地理要素的分布规律；等等。

(2)相互关系分析。这类研究，主要是对地理要素、地理事物之间的相互关系进行定量分析。譬如，运用统计相关分析方法定量地揭示地理要素之间的相关程度；运用灰色关联分析方法揭示地理事物之间相互联系的密切程度；运用回归分析方法给出地理要素之间相关关系的定量表达式；运用投入产出分析方法定量分析区域经济系统中各个产业之间的相互联系；等等。

(3)类型研究。主要是对地理事物的类型和各种地理区域进行定量划分。譬如，运用模式识别方法、判别分析方法、聚类分析方法等定量地研究土地类型、地带及自然区和经济区的划分问题等。

(4)网络分析。主要是对水系、交通网络、行政区域、经济区域等的空间结构进行定量分析。在地理网络分析中，几何学方法和图论方法是常用的主要方法。譬如，交通网络中结点之间的接近度、可达性、最短路径，及最大流与最小运费流，以及行政或经济区域中的城镇体系及其等级-规模等问题的研究，均属于网络分析的范畴。

(5)趋势面分析。趋势面分析，就是运用适当的数学方法计算出一个空间曲面，并以这个空间曲面去拟合地理要素分布的空间形态，展示其空间分布规律。这种空间曲面就称之为趋势面。趋势面分析所采用的数学方法通常是回归分析方法。其分析的步骤是：首先运用回归分析方法拟合出所要分析的地理要素的趋势面方程，然后以趋势面方程计算出每一个地理测点上的该地理要素的趋势值，并以一定的间隔画出趋势等值线图。这种趋势等值线图就展示了所要分析的地理要素的空间分布规律。

(6)空间相互作用分析。主要是定量地分析各种“地理流”在不同区域之间“流动”的方向与强度。譬如，运用线性规划方法研究某个大区域中各个小区之间的货流问题；运用投入产出分析方法研究各个区域之间产品的流动及分配与消费问题；运用一些已经建立的理论模式研究不同区域之间的人口流动问题、商品购销问题；等等。

(7)系统仿真研究。就是针对复杂的地理问题——即对象系统，在对各种系统要素之间的相互关系与反馈机制分析的基础上，构造系统结构，建立描述系统的数学模型，并以适当的计算方法与算法语言将数学模型转化为计算机可以识别与运行的工作模型，通过模型的运行，对真实系统进行模拟与仿真，从而达到揭示系统的运行机制与规律的目的。实践证明，系统动力学方法是地理系统仿真研究中可供借鉴的一种有效方法，它为许多复杂地理问题的研究提供了一种可供尝试的途径。

(8)过程模拟与预测研究。任何地理事物、地理现象，都随着时间在不断地运动和变化着，即经历着特定的地理过程。这类研究，旨在通过对地理过程的模拟与拟合，定量地揭示地理事物、地理现象随时间变化的规律，从而对其未来发展趋势作出预测。在地球表层系统中，主要的地理过程包括气候过程、水文过程、生物过程、地貌过程、生态-环境过程、经济过程、社会过程、文化过程等。对于这些过程的模拟与预测研究，经常采用的数学方法有回归分析法、马尔可夫方法、灰色建模方法、系统动力学方法等。

(9)空间扩散研究。这类研究，旨在定量地揭示各种地理现象，包括自

然现象、经济现象、社会现象、文化现象、技术现象在地理空间上的扩散规律。譬如，坡面泥石流运动、各种污染物在水体和大气中的扩散、各种经济现象的集聚与扩散、文化与技术的传播等问题，都属于空间扩散研究的范畴。这类研究，经常采用的方法有微分建模方法、数学物理方法、蒙卡罗模拟方法等。

(10)空间行为研究。主要是对人类活动的空间行为决策进行定量的研究。譬如，资源利用与环境保护问题、经济活动的空间组织问题、产业布局的区位问题、城乡区域规划问题等都属于空间行为研究的范畴。这类研究，经常采用的数学方法有数学规划方法，如线性规划、多目标规划、多维灰色规划方法等，以及决策分析方法，如 AHP 决策分析方法、风险型决策方法、非确定型决策方法、模糊决策方法、灰色局势决策方法等。

(11)地理系统优化调控研究。主要是运用系统控制论的有关原理与方法，研究人-地相互作用的地理系统的优化调控问题，寻求人口、资源、环境与社会经济协调发展的方法、途径与措施。这类研究，经常采用的是现代控制论方法、大系统理论及灰色去余控制理论等。

(二)应用数学方法必须注意的一些问题

在现代地理学研究中，为了成功地运用数学方法，达到定量分析的目的，必须注重如下几个方面的问题：

(1)关于地理数据的筛选与质量检验问题。数学方法，在现代地理学研究中的作用是重要的，它是建立模型和进行定量分析的基本工具与先决条件。但是，地理数据却是定量地研究地理问题的基础，它在建模分析中的作用有两个方面：一是确定模型中的参数与初值；二是检验模型的正确性、合理性和有效性。没有地理数据，模型中的参数与初值将无法确定，模型的正确性、合理性和有效性将无法检验。地理数据的质量，直接影响着由模型所得出的研究结果的正确性。在地理问题的研究中，运用数学方法所建立的定量分析模型，可以被形象地看作加工原料、制造产品的“机器”或“设备”，这里的“原料”就是输入模型的原始地理数据，而“产品”便是由模型得出的研究结果。显然，“产品”的质量不仅取决于“机器”的性能，而且还依赖于“原料”的品质。如果输入的地理数据质量不高，则输入的结论就不会可靠。由此可见，在地理问题研究中，地理数据的丰富性、完备性和准确性，也是能否成功地运用数学方法的关键。所以，在运用数学方法研究地理问题时，就必须首先注重对地理数据的筛选和质量检验工作。

(2)模型的建造问题。描述地理问题的数学模型，是对地理问题进行定量地研究的依据。所以，描述地理问题的数学模型的建造，是对地理问题进行定量研究的关键环节之一。英国著名地理学家威尔逊 (A.Wilson)，曾就如何建造地理数学模型发表了自己的见解，这些见解可以归纳为如下几条：①建造一个模型，首先必须明确建模的目标，即建模者必须回答所建模型将被用来做什么？企图解决什么问题？②地理问题——即所研究的对象系统，其构成要素是什么？这些要素之间的相互联系、相互作用及其动态变化应该由什么形式的变量被模型所反映？其中哪些变量是以量化变量的形式出现的？③在各类变量中，必须明确哪些变量是可控变量？即通过对哪些变量的调控可以使系统的行为发生改变？④在模型中，如何处理时间概念？即认为被研究的对象系统是无记忆系统还是记忆系统？是建立静态模型还是建立动态模型？⑤所建模型将采用什么观点、解决哪些理论问题？与此问题有关的建立

模型的基本假设，以及所依据的理论，将要解决的问题等都将直接或间接地体现在模型之中；⑥能用于建模的有关数据、资料是什么？其可靠性如何？应采用什么样的建模技术？有现成的技术方法可供借鉴还是需要建造新模型？采用什么方法确定模型的参数？⑦所建模型的精度，以及该模型的合理性和有效性如何？采用什么方法和手段检验所建模型？威氏的这些见解道出了建造地理数学模型所必须注意的各个环节，同时也为我们提供了一个一般性的建模程序。

(3)与 GIS 相结合的问题。地理信息系统 (GIS)，是本世纪 70 年代后期发展起来的，对地理数据进行采集、输入、存储、更新、检索、管理及综合分析输出的计算机应用技术。它是以计算机为工具，综合利用定位观测数据、统计调查数据、地图数据、遥感数据等，通过一系列空间操作与分析，对地理学进行综合研究的现代化手段。数学方法，只有与 GIS 技术相结合，才能不断地提高其应用层次与水平，不断地拓宽其应用领域，充分发挥它在现代地理学研究中的作用。一方面，从数学方法的角度来看，对于一些复杂地理问题的研究，采用任何单一的数学方法和单个数学模型都是很难奏效的。解决这类问题，需要综合运用多种数学方法，建立一系列具有分析、模拟、仿真、预测、规划、决策、调控等各种功能的众多模型组成的模型系统才能完成。然而，这种模型系统的运行，不但需要大量地理数据构成的数据库的支持，还需要强有力的计算方法与计算机程序的支持，而且由模型系统运行所得到的研究结论也需要以简明扼要的形式——地图、统计图形或表格方式被输出。显然，对于模型系统的这些支持，必须由 GIS 技术才能完成。另一方面，从 GIS 的角度来看，它不仅需要运用数学方法为其建造空间分析模型，如数字地形模型 (DTM)、空间统计分析模型、叠加 (Overlay) 分析模型、缓冲 (Buffer) 区分析模型等，以及其它应用分析模型，如综合评价模型、预测模型、规划模型、决策分析模型等；而且就连 GIS 中的一些基本技术，如空间数据的编码、数据格式的转换算法、遥感数据的几何校正、数据模型与数据库的建造等都需要借助有关的数学方法来实现。近几年来所出现的一种针对一些特定的领域的面向应用对象的、高层次的智能化的地理信息系统——地理决策支持系统，就是数学方法、人工智能技术与 GIS 技术在应用地理学研究领域中相结合的成功典范。

第二章 统计分析方法

地理系统，是由多种要素相复合而构成的复杂巨系统。在这个系统中，一方面，各种要素之间存在着相互联系、相互影响和相互制约的关系；另一方面，各种要素的复合作用又使各种地理事物和地理现象表现出强烈的地域差异性。为了定量地揭示各种地理要素之间的相互关系，以及各种地理事物和地理现象所表现出来的地域分异规律，就必须采用以概率论和数量统计知识为基础的统计分析方法对地理系统进行深入的研究。本章，我们将介绍和探讨统计相关分析、回归分析、系统聚类分析、主成分分析、马尔可夫预测等统计分析方法在地理系统分析中的应用问题。

第一节 地理要素间的相关分析

地理要素之间的相关分析的任务，是揭示地理要素之间相互关系的密切程度。而地理要素之间相互关系的密切程度的测定，主要是通过对相关系数的计算与检验来完成的。一、两要素间相关程度的测定

(一)相关系数的计算与检验

1.相关系数的计算

对于两个要素 x 与 y ，如果它们的样本值分别为 x_i 和 y_i ($i=1, 2, \dots, n$)，则它们之间的相关系数被定义为：

$$r_{xy} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2} \cdot \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2}} \quad (1)$$

在 (1) 式中， $\bar{x}$ 和 $\bar{y}$ 分别表示两个要素样本值的平均值，即 $\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$ ，

$\bar{y} = \frac{1}{n} \sum_{i=1}^n y_i$ ； r_{xy} 为要素 x 与 y 之间的相关系数，它就是表示该两要素之间相关程度的统计指标，其值在 $[-1, 1]$ 区间之内。 $r_{xy} > 0$ ，表示正相关，即两要素同向发展； $r_{xy} < 0$ ，表示负相关，即两要素异向发展。 r_{xy} 的绝对值越接近于 1，表示两要素的关系越密切；越接近于 0，表示两要素的关系越不密切。

如果记：

$$L_{xy} = \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) = \sum_{i=1}^n x_i y_i - \frac{1}{n} \left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n y_i \right)$$

$$L_{xx} = \sum_{i=1}^n (x_i - \bar{x})^2 = \sum_{i=1}^n x_i^2 - \frac{1}{n} \left(\sum_{i=1}^n x_i \right)^2$$

$$L_{yy} = \sum_{i=1}^n (y_i - \bar{y})^2 = \sum_{i=1}^n y_i^2 - \frac{1}{n} \left(\sum_{i=1}^n y_i \right)^2$$

则公式 (1) 式可以进一步简化为

$$r_{xy} = \frac{L_{xy}}{\sqrt{L_{xx}L_{yy}}} \quad (2)$$

例如，某地区 1981—1990 年期间的粮食总产量 (x) 和农业总产值 (y) 数据如表 2-1 所示。试计算该地区粮食总产量与农业总产值之间的相关系数

表 2-1 某地区粮食总产量与农业产值数据

据表 2-1 计算可得：

$$L_{xy} = \sum_{i=1}^n x_i y_i - \frac{1}{10} \left(\sum_{i=1}^{10} x_i \right) \left(\sum_{i=1}^{10} y_i \right) = 30.7000$$

$$L_{xx} = \sum_{i=1}^{10} x_i^2 - \frac{1}{10} \left(\sum_{i=1}^{10} x_i \right)^2 = 360.000$$

$$L_{yy} = \sum_{i=1}^{10} y_i^2 - \frac{1}{10} \left(\sum_{i=1}^{10} y_i \right)^2 = 3.0840$$

故：

$$r_{xy} = \frac{L_{xy}}{\sqrt{L_{xx}L_{yy}}} = \frac{30.7000}{\sqrt{360.0000 \times 3.0840}} = 0.9214$$

即该地区粮食总产量与农业总产量之间的相关系数为 0.9214。

如果问题涉及到 $x_1, x_2, \dots, x_n$ 等 n 个要素，则对于其中任何两个要素 x_i 和 x_j ，我们都可以按照公式 (1) 或 (2) 式计算它们之间的相关系数 r_{ij} ，这样就得到多要素的相关系数矩阵：

$$R = \begin{bmatrix} r_{11} & r_{12} & \cdots & r_{1n} \\ r_{21} & r_{22} & \cdots & r_{2n} \\ M & M & M \\ r_{n1} & r_{n2} & \cdots & r_{nn} \end{bmatrix} \quad (3)$$

显然，由公式 (1) 或 (2) 式容易知道：

(1) $r_{ii}=1$ ($i=1, 2, \dots, n$)，即每一个要素 x_i 与它自己本身的相关程度最大；

(2) $r_{ij}=r_{ji}$ ($i, j=1, 2, \dots, n$)，即第 i 个要素 (x_i) 对第 j 个要素 (x_j) 的相关程度，与第 j 个要素 (x_j) 对第 i 个要素 (x_i) 的相关程度相等。

2. 相关系数的检验

当要素之间的相关系数求出之后，还需要对所求得的相关系数进行检验。这是因为，这

表 2-2 检验相关系数 $\rho=0$ 的临界值 (r_a) 表

$$p \{ |r| > r_a \} = \alpha$$

α f	0.10	0.05	0.02	0.01	0.0001
1	0.98769	0.99692	0.999507	0.999877	0.999998
2	0.90000	0.95000	0.9800	0.99000	0.999000
3	0.8054	0.8783	0.93433	0.95873	0.991160
4	0.7293	0.8114	0.8822	0.91720	0.97406
5	0.6694	0.7545	0.8329	0.8745	0.95047
6	0.6215	0.7067	0.7887	0.8343	0.92493
7	0.5822	0.6664	0.7493	0.7977	0.8982
8	0.5494	0.6319	0.7155	0.7646	0.8721
9	0.5214	0.6021	0.6851	0.7348	0.8471
10	0.4973	0.5760	0.6581	0.7079	0.8233
11	0.4762	0.5529	0.6339	0.6835	0.8010
12	0.4575	0.5324	0.6120	0.6614	0.7800
13	0.4409	0.5139	0.5923	0.6411	0.7603
14	0.4259	0.4973	0.5742	0.6226	0.7420
15	0.4124	0.4821	0.5577	0.6055	0.7246
16	0.4000	0.4683	0.5425	0.5897	0.7084
17	0.3887	0.4555	0.5285	0.5751	0.6932
18	0.3783	0.4438	0.5155	0.5614	0.6787
19	0.3687	0.4329	0.5034	0.5487	0.6652
20	0.3598	0.4227	0.4921	0.5368	0.6524
25	0.3233	0.3809	0.4451	0.4869	0.5794
30	0.2960	0.3494	0.4093	0.4487	0.5541
40	0.2573	0.3044	0.3578	0.3932	0.4896
45	0.2428	0.2875	0.3384	0.3721	0.4648
50	0.2306	0.2732	0.3218	0.3541	0.4433
60	0.2108	0.2500	0.2948	0.3248	0.4078
70	0.1954	0.2319	0.2737	0.3017	0.3799
80	0.1829	0.2172	0.2565	0.2830	0.3568
90	0.1726	0.2050	0.2422	0.2673	0.3375
	0.1628	0.1946	0.2301	0.2540	0.3211

里的相关系数是根据要素之间的样本值计算出来的，它随着样本数的多少或取样方式的不同而不同，因此它只是要素之间的样本相关系数，只有通过检验，才能知道它的可信度。

一般情况下，相关系数的检验，是在给定的置信水平下，通过查相关系数检验的临界值表来完成的。表 2-2 给出了相关系数真值 $\rho=0$ （即两要素不相关）时样本相关系数的临界值 r_a 。

在表 2-2 中，左边的 f 值称为自由度，其数值为 $f=n-2$ ，这里 n 为样本数；上方的 a 代表不同的置信水平；表内的数值代表不同的置信水平下相关

系数 $\rho=0$ 的临界值, 即 r_a ; 公式 $p = \{ |r| > r_a \} = a$ 的意思是当所计算的相关系数 r 的绝对值大于在 a 水平下的临界值 r_a 时, 两要素不相关 (即 $\rho=0$) 的可能性只有 a 。在前例中, $f=10-2=8$, 在不同的置信水平下的临界值 r_a 可以从表中查得: $r_{0.1}=0.5494$, $r_{0.05}=0.6319$, $r_{0.02}=0.7155$, $r_{0.01}=0.7646$, $r_{0.001}=0.8721$ 。由于 $r_{xy}=0.9214 > r_{0.001}=0.8721$, 这说明该地区粮食总产量 (x) 与农业总产值 (y) 不相关的概率只有 $a=0.001$, 即 0.1%, 换句话说, 该地区粮食总产量 (x) 与农业总产值 (y) 同向相关的概率达 0.999, 即 99.9%。

一般而言, 当 $|r| < r_{0.1}$ 时, 则认为两要素不相关, 这时的样本相关系数就不能反映两要素之间的关系。

(二) 等级相关系数的计算与检验

1. 等级相关系数的计算

等级相关系数, 又称顺序相关系数, 与前述相关系数一样, 它也是描述两要素之间相关程度的一种统计指标, 不过在计算方法上, 与前述相关系数的计算有所不同。等级相关系数是将两要素的样本值按数值的大小顺序排列位次, 以各要素样本值的位次代替实际数据而求得的一种统计量。实际上, 它是位次分析方法的数量化。

设两个要素 x 和 y 有 n 对样本值, 令 R_1 代表要素 x 的序号 (或位次), R_2 代表要素 y 的序号 (或位次), $d_i^2 = (R_{1i} - R_{2i})^2$ 代表要素 x 和 y 的同一组样本位次差的平方, 那么要素 x 与 y 之间的等级相关系数 (r'_{xy}) 被定义

$$r'_{xy} = 1 - \frac{6 \sum_{i=1}^n d_i^2}{n(n^2 - 1)} \quad (4)$$

例如, 我国 1985 年各省 (市, 区) 的总人口 (x) 和社会总产值 (y) 及其位次列于表 2-3。试计算总人口 (x) 与社会总产值 (y) 之间的等级相关系数。

由表 2-3 计算可知: $n = 29$, $n(n^2 - 1) = 24360$, $\sum_{i=1}^{29} d_i^2 = 1134$, 故

$$r'_{xy} = 1 - \frac{6 \sum_{i=1}^{29} d_i^2}{n(n^2 - 1)} = 1 - \frac{6 \times 1134}{24360} = 0.726$$

即: 总人口 (x) 与社会总产值 (y) 的等级相关系数为 0.726

2. 等级相关系数的检验

与相关系数一样, 等级相关系数是否显著, 也需要检验。表 2-4 给出了等级相关系数检验的临界值。

表 2-4 的内容与表 2-2 的内容相似, n 代表样本个数, a 代表不同的置信水平, 也称显著水

表 2-3 1985 年全国各省 (市, 区) 总人口与社会总产值

省(市、区)	总人口(x)及其位次		社会总产量(y)及其位次		位次差的平方 $d_i^2 = (R_{1i} - R_{2i})^2$
	人口数(万人)	位次 R_1	总产值(亿元)	位次 R_2	
北京	960	25	482.10	14	121
天津	808	26	431.92	15	121
河北	5548	7	715.30	10	9
山西	2627	19	404.40	17	4
内蒙	2007	22	252.53	23	1
辽宁	3686	12	1076.53	4	64
吉林	2298	20	422.24	16	16
黑龙江	3311	15	629.79	12	9
上海	1217	24	1055.83	5	361
江苏	6213	5	1536.66	1	16
浙江	4032	10	869.26	7	9
安徽	5156	8	577.83	13	25
福建	2713	18	345.30	19	1
江西	3460	13	363.03	18	25
山东	7965	3	1240.17	2	1
河南	7713	2	815.51	8	36
湖北	4913	9	806.42	9	0
湖南	5622	6	634.62	11	25
广东	6253	4	1113.92	3	1
广西	3873	11	315.69	21	100
四川	10183	1	1046.11	6	25
贵州	2968	17	216.84	25	64
云南	3406	14	289.26	22	64
西藏	199	29	22.24	29	0
陕西	3002	16	343.05	20	16
甘肃	2041	21	232.68	24	9
青海	407	28	54.95	27	1
宁夏	415	27	51.45	28	1
新疆	1361	23	186.50	26	9
Σ	104532	/	17193.71	/	1134

表 2-4 等级相关系数检验的临界值

n	显著水平 a		n	显著水平 a	
	0.05	0.01		0.05	0.01
4	1.000		16	0.425	0.601
5	0.900	1.000	18	0.399	0.564
6	0.829	0.943	20	0.377	0.534
7	0.714	0.893	22	0.359	0.508
8	0.643	0.833	24	0.343	0.485
9	0.600	0.783	26	0.329	0.465
10	0.564	0.746	28	0.317	0.448
12	0.506	0.712	30	0.306	0.432
14	0.456	0.645			

平，表中的数值为临界值 r_a 。在上例中， $n=29$ ，表中没有给出相应的样本数下的临界值 r_a ，但我们发现，在同一显著水平下，随着样本数的增大，临界值 r_a 减少。在 $n=28$ 时，查表可知： $r_{0.05}=0.317$ ， $r_{0.01}=0.448$ ，由于 $r'_{xy}=0.726 > r_{0.01}=0.448$ ，故 r'_{xy} 在 $a=0.01$ 的置信水平上是显著的。

二、多要素间相关程度的测定

(一)偏相关系数的计算与检验

地理系统是一种多要素的复杂巨系统，其中一个要素的变化必然影响到其它各要素的变化。在多要素所构成的地理系统中，当我们研究某一个要素对另一个要素的影响或相关程度时，把其它要素的影响视为常数（保持不变），即暂不考虑其它要素的影响，而单独研究那两个要素之间的相互关系的密切程度时，则称为偏相关。用以度量偏相关程度的统计量，称为偏相关系数。

1. 偏相关系数的计算

偏相关系数，可利用单相关系数来计算。假设有三个要素 x_1, x_2, x_3 ，其两两间单相关系数矩阵为

$$R = \begin{bmatrix} r_{11} & r_{12} & r_{13} \\ r_{21} & r_{22} & r_{23} \\ r_{31} & r_{32} & r_{33} \end{bmatrix} = \begin{bmatrix} 1 & r_{12} & r_{13} \\ r_{21} & 1 & r_{23} \\ r_{31} & r_{32} & 1 \end{bmatrix}$$

因为相关系数矩阵是对称的，故在实际计算时，只要计算出 r_{12} ， r_{13} 和 r_{23} 即可。在偏相关分析中，常称这些单相关系数为零级相关系数。对于上述三个要素 x_1, x_2, x_3 ，它们之间的偏相关系数共有三个，即 $r_{12 \cdot 3}$ ， $r_{13 \cdot 2}$ ， $r_{23 \cdot 1}$

（下标点后面的数字，代表在计算偏相关系数时，保持不变量，如 $r_{12 \cdot 3}$ 即表示 x_3 保持不变），其计算公式分别如下：

$$r_{12 \cdot 3} = \frac{r_{12} - r_{13}r_{23}}{\sqrt{(1 - r_{13}^2)(1 - r_{23}^2)}} \quad (5)$$

$$r_{13 \cdot 2} = \frac{r_{13} - r_{12}r_{23}}{\sqrt{(1 - r_{12}^2)(1 - r_{23}^2)}} \quad (6)$$

$$r_{23 \cdot 1} = \frac{r_{23} - r_{12}r_{13}}{\sqrt{(1 - r_{12}^2)(1 - r_{13}^2)}} \quad (7)$$

式 (5) — (7) 表示三个偏相关系数，称为一级偏相关系数。

若有四个要素 X_1, X_2, X_3, X_4 ，则有六个偏相关系数，即 $r_{12 \cdot 34}, r_{13 \cdot 24}, r_{14 \cdot 23}, r_{23 \cdot 14}, r_{24 \cdot 12}, r_{34 \cdot 12}$ ，它们称为二级偏相关系数，其计算公式分别如下：

$$r_{12 \cdot 34} = \frac{r_{12 \cdot 3} - r_{14 \cdot 3}r_{24 \cdot 3}}{\sqrt{(1 - r_{14 \cdot 3}^2)(1 - r_{24 \cdot 3}^2)}} \quad (8)$$

$$r_{13 \cdot 24} = \frac{r_{13 \cdot 2} - r_{14 \cdot 2}r_{34 \cdot 2}}{\sqrt{(1 - r_{14 \cdot 2}^2)(1 - r_{34 \cdot 2}^2)}} \quad (9)$$

$$r_{14 \cdot 23} = \frac{r_{14 \cdot 2} - r_{13 \cdot 2}r_{43 \cdot 2}}{\sqrt{(1 - r_{13 \cdot 2}^2)(1 - r_{43 \cdot 2}^2)}} \quad (10)$$

$$r_{23 \cdot 14} = \frac{r_{23 \cdot 1} - r_{24 \cdot 1}r_{34 \cdot 1}}{\sqrt{(1 - r_{24 \cdot 1}^2)(1 - r_{34 \cdot 1}^2)}} \quad (11)$$

$$r_{24 \cdot 13} = \frac{r_{24 \cdot 1} - r_{23 \cdot 1}r_{43 \cdot 1}}{\sqrt{(1 - r_{23 \cdot 1}^2)(1 - r_{43 \cdot 1}^2)}} \quad (12)$$

$$r_{34 \cdot 12} = \frac{r_{34 \cdot 1} - r_{32 \cdot 1}r_{42 \cdot 1}}{\sqrt{(1 - r_{32 \cdot 1}^2)(1 - r_{42 \cdot 1}^2)}} \quad (13)$$

在式 (8) 中， $r_{12 \cdot 34}$ 表示在 X_3 和 X_4 保持不变的条件下， X_1 和 X_2 的偏相关系数，其余式 (9) — (13) 依此类推。

应所考虑的要素多于四个时，则可以依次考虑，计算三级甚至更多级偏相关系数。

假若，对于某四个地理要素 X_1, X_2, X_3, X_4 的 23 个样本数据，经过计算得到了如下的单相关系数矩阵：

$$R = \begin{bmatrix} r_{11} & r_{12} & r_{13} & r_{14} \\ r_{21} & r_{22} & r_{23} & r_{24} \\ r_{31} & r_{32} & r_{33} & r_{34} \\ r_{41} & r_{42} & r_{43} & r_{44} \end{bmatrix} = \begin{bmatrix} 1 & 0.416 & 0.346 & 0.579 \\ 0.416 & 1 & -0.592 & 0.950 \\ -0.346 & -0.592 & 1 & -0.469 \\ 0.579 & 0.950 & -0.469 & 1 \end{bmatrix}$$

为了说明偏相关系数的计算方法，现以 (14) 式中的单相关系数为例，来计算一级和二级偏相关系数。为了计算二级偏相关系数，需要先计算一级偏相关系数，由 (5) 式可求得

$$r_{12 \cdot 3} = \frac{r_{12} - r_{13}r_{23}}{\sqrt{(1-r_{13}^2)(1-r_{23}^2)}} = \frac{0.416 - 0.346 \times (-0.592)}{\sqrt{(1-0.346^2)(1-0.592^2)}} = 0.821$$

同理，依次可以计算出其它各一级偏相关系数，见表 2-5。

表 2-5 一级偏相关系数

$r_{12 \cdot 3}$	$r_{13 \cdot 2}$	$r_{14 \cdot 2}$	$r_{14 \cdot 3}$	$r_{23 \cdot 1}$	$r_{24 \cdot 1}$	$r_{24 \cdot 3}$	$r_{34 \cdot 1}$	r_{34}
0.821	0.808	0.647	0.895	-0.863	0.956	0.945	-0.875	0.37

在一级偏相关系数求出以后，便可代入公式计算二级偏相关系数，如由(8)式计算可得

$$r_{12 \cdot 34} = \frac{r_{12 \cdot 3} - r_{14 \cdot 3}r_{24 \cdot 3}}{\sqrt{(1-r_{14 \cdot 3}^2)(1-r_{24 \cdot 3}^2)}} = \frac{0.821 - 0.895 \times 0.945}{\sqrt{(1-0.895^2)(1-0.945^2)}} = 0.170$$

同理，依次可计算出其它各二级偏相关系数，见表 2-6。

表 2-6 二级偏相关系数

$r_{12 \cdot 34}$	$r_{13 \cdot 24}$	$r_{14 \cdot 23}$	$r_{23 \cdot 14}$	$r_{24 \cdot 13}$	$r_{34 \cdot 12}$
-0.170	0.802	0.635	-0.187	0.821	-0.337

容易看出，偏相关系数具有下述性质：

(1)偏相关系数分布的范围在-1 到 1 之间，譬如，固定 X_3 ，则 X_1 与 X_2 间的偏相关系数满足 $-1 \leq r_{12 \cdot 3} \leq 1$ 。当 $r_{12 \cdot 3}$ 为正值时，表示在 X_3 固定时， X_1 与 X_2 之间为正相关；当 $r_{12 \cdot 3}$ 为负值时，表示在 X_3 固定时， X_1 与 X_2 之间为负相关。

(2)偏相关系数的绝对值越大，表示其偏相关程度越大。例如， $|r_{12 \cdot 3}| = 1$ ，则表示当 X_3 固定时， X_1 与 X_2 之间完全相关；当 $|r_{12 \cdot 3}| = 0$ 时，表示当 X_3 固定时， X_1 与 X_2 之间完全无关。

(3)偏相关系数的绝对值必小于或最多等于由同一系列资料所求得的复相关系数（详见后述），即 $R_{1 \cdot 23} \geq |r_{12 \cdot 3}|$ 。

2. 偏相关系数的显著性检验

偏相关系数的显著性检验，一般采用 t-检验法。其统计量计算公式为

$$t = \frac{r_{12 \cdot 34 \cdots m}}{\sqrt{1 - r_{12 \cdot 34 \cdots m}^2}} \sqrt{n - m - 1} \quad (15)$$

在 (15) 式中， $r_{12 \cdot 34 \cdots m}$ 为偏相关系数，n 为样本数，m 为自变量个数。

譬如，对于前述计算得到的偏相关系数 $r_{24 \cdot 13} = 0.821$ ，由于 $n=23$ ， $m=3$ ，

故

$$t = \frac{0.821}{\sqrt{1-0.821^2}} \sqrt{23-3-1} = 6.268$$

查 t 分布表, 可得出不同显著水平上的临界值 t_a , 若 $t > t_a$, 则表示偏相关显著; 反之, $t < t_a$, 则偏相关不显著。在自由度为 $23-3-1=19$ 时, 查表得 $t_{0.001}=3.883$, 所以 $t > t_a$, 这表明在置信度水平 $\alpha=0.001$ 上, 偏相关系数 $r_{24 \cdot 13}$ 是显著的。

(二) 复相关系数的计算与检验

严格来说, 以上的分析都是揭示两个要素 (变量) 间的相关关系, 或者是在其它要素 (变量) 固定的情况下来研究两要素间的相关关系的。但实际上, 一个要素的变化往往受多种要素的综合作用和影响, 而单相关或偏相关分析的方法都不能反映各要素的综合影响。要解决这一问题, 就必须采用研究几个要素同时与某一个要素之间的相关关系的复相关分析法。几个要素与某一个要素之间的复相关程度, 可用复相关系数来测定。

1. 复相关系数的计算

复相关系数, 可以利用单相关系数和偏相关系数求得。

设 Y 为因变量, $X_1, X_2, \dots, X_k$ 为自变量, 则将 Y 与 $X_1, X_2, \dots, X_k$ 之间的复相关系数记为 $R_{y \cdot 12 \dots k}$ 。其计算公式如下

当有两个自变量时,

$$R_{y \cdot 12} = \sqrt{1 - (1 - r_{y1}^2)(1 - r_{y2 \cdot 1}^2)} \quad (16)$$

当有三个自变量时,

$$R_{y \cdot 123} = \sqrt{1 - (1 - r_{y1}^2)(1 - r_{y2 \cdot 1}^2)(1 - r_{y3 \cdot 12}^2)} \quad (17)$$

一般地, 当有 k 个自变量时,

$$R_{y \cdot 12 \dots k} = \sqrt{1 - (1 - r_{y1}^2)(1 - r_{y2 \cdot 1}^2) \dots [1 - r_{yk \cdot 12 \dots (k-1)}^2]} \quad (18)$$

以 (14) 式所描述的四个地理要素之间的相互关系为例, 若以 X_4 为因变量, X_1, X_2, X_3 为自变量, 则可以按下式计算 X_4 与 X_1, X_2, X_3 之间的复相关系数

$$\begin{aligned} R_{4 \cdot 123} &= \sqrt{1 - (1 - r_{41}^2)(1 - r_{42 \cdot 1}^2)(1 - r_{43 \cdot 12}^2)} \\ &= \sqrt{1 - (0.579^2)(1 - 0.956^2)[1 - (-0.337)^2]} = 0.974 \end{aligned}$$

关于复相关系数的性质, 可以概括为如下几点:

(1) 复相关系数介于 0 到 1 之间, 即

$$0 \leq R_{y \cdot 12 \dots k} \leq 1$$

(1) 复相关系数越大, 则表明要素 (变量) 之间的相关程度越密切。复相关系数为 1, 表示完全相关; 复相关系数为 0, 表示完全无关。

(3) 复相关系数必大于或至少等于单相关系数的绝对值。

2. 复相关系数的显著性检验

对复相关系数的显著性检验, 一般采用 F -检验法。其统计量计算公式为

$$F = \frac{R_{y \cdot 12 \dots k}^2}{1 - R_{y \cdot 12 \dots k}^2} \times \frac{n - k - 1}{k} \quad (19)$$

在 (19)式中, n 为样本数, k 为自变量个数。对于前述计算得出的复相关系数 $R_{4.123}=0.974$, 由于 $n=23$, $k=3$, 故

$$F = \frac{0.974}{1-0.974^2} \times \frac{23-3-1}{3} = 120.1907$$

查 F -检验的临界值表 (见本书附录 II), 可以得出不同显著水平上的临界值 F_α , 若 $F > F_{0.01}$, 则表示复相关在置信度水平 $\alpha=0.01$ 上显著, 称为极显著; 若 $F_{0.05} < F \leq F_{0.01}$, 则表示复相关在置信度水平 $\alpha=0.05$ 上显著; 若 $F_{0.10} \leq F \leq F_{0.05}$, 则表示复相关在置信度水平 $\alpha=0.10$ 上显著; 若 $F > F_{0.10}$, 则表示复相关不显著, 即因变量 Y 与 K 个自变量之间的关系不密切。在上例中, $F=120.1907 > F_{0.01}=5.0103$, 故复相关达到了极显著水平。

第二节 地理要素间的回归分析

地理要素间的相关分析揭示了诸地理要素之间相互关系的密切程度。然而诸要素之间相互关系的进一步具体化,譬如某一地理要素与其它地理要素之间的相互关系若能用一定的函数形式予以近似的表达,那么其实用意义将会更大。在复杂地理系统中,某些要素的变化很难预测或控制,相反,另外一些要素则容易被预测或控制。在这种复杂地理系统中,若能在某些难测难控的要素与其它易测易控的要素之间建立一种近似的函数表达式,则就可以比较容易地通过那些易测易控要素的变化情况去了解那些难测难控的要素的变化情况。数理统计学为我们提供了回归分析方法,是研究要素之间具体的数量关系的一种强有力的手段,借助于这种方法,可以建立地理要素之间的相关关系模型——回归分析模型。

现代地理科学研究的对象是多层次多要素的复杂系统,其要素之间的相互关系,既有线性的,也有非线性的。因此,地理要素之间的回归分析模型,既有线性回归模型,也有非线性回归模型。但是在回归分析研究中,许多非线性模型都可以通过变量变换将其转化为线性模型来处理。下面我们首先来介绍地理要素之间的线性回归模型。

一、一元线性回归模型

一元线性回归模型描述的是两个要素(变量)之间的线性相关关系。假设有两个地理要素(变量) x 和 y , x 为自变量, y 为因变量。则,一元线性回归模型的基本结构形式为

$$y_a = a + bx_a + \varepsilon_a \quad (1)$$

在(1)式中, a 和 b 为待定参数; $a=1, 2, \dots, n$ 为 n 组观测数据 $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ 的下标; ε_a 为随机变量。如果记 a 和 b 分别为参数 a 与 b 的拟合值,便得一元线性回归模型

(2)式代表 x 与 y 之间相关关系的拟合直线,常称为回归直线; $\hat{y}$ 是 y 的估计值,亦称回归值。

(一)参数 a 、 b 的最小二乘估计

实际观测值 y_i 与回归值 $\hat{y}_i$ 之差 $e_i = y_i - \hat{y}_i$,刻画了 y_i 与 $\hat{y}_i$ 的偏离程度,即表示实际观测值与回归估计值之间的误差大小。参数 a 与 b 的最小二乘拟合原则要求 y_i 与 $\hat{y}_i$ 的误差 e_i 的平方和达到最小,即

$$\begin{aligned} Q &= \sum_{i=1}^n e_i^2 = \sum_{i=1}^n (y_i - \hat{y}_i)^2 \\ &= \sum_{i=1}^n (y_i - a - bx_i)^2 \rightarrow \min \end{aligned} \quad (3)$$

根据取极值的必要条件,有

$$\begin{cases} \frac{\partial Q}{\partial a} = -2 \sum_{i=1}^n (y_i - a - bx_i) = 0 \\ \frac{\partial Q}{\partial b} = -2 \sum_{i=1}^n (y_i - a - bx_i)x_i = 0 \end{cases}$$

即

$$\begin{cases} \sum_{i=1}^n (y_i - a - bx_i) = 0 \\ \sum_{i=1}^n (y_i - a - bx_i)x_i = 0 \end{cases}$$

上述方程组可以进一步写成

$$\begin{cases} na + \left(\sum_{i=1}^n x_i \right) b = \sum_{i=1}^n y_i \\ \left(\sum_{i=1}^n x_i \right) a + \left(\sum_{i=1}^n x_i^2 \right) b = \sum_{i=1}^n x_i y_i \end{cases} \quad (4)$$

方程组 (4) 式通常被称为正规方程组，它又可以被写成矩阵形式

$$\begin{bmatrix} n & \sum_{i=1}^n x_i \\ \sum_{i=1}^n x_i & \sum_{i=1}^n x_i^2 \end{bmatrix} \begin{bmatrix} a \\ b \end{bmatrix} = \begin{bmatrix} \sum_{i=1}^n y_i \\ \sum_{i=1}^n x_i y_i \end{bmatrix} \quad (4)$$

解上述正规方程组 (4) 式或 (4') 式，就可以得到关于参数 a 与 b 的拟合值：

$$\hat{a} = \bar{y} - \hat{b} \bar{x} \quad (5)$$

$$\begin{aligned} b &= \frac{L_{xy}}{L_{xx}} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2} \\ &= \frac{\sum_{i=1}^n x_i y_i - \frac{1}{n} \left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n y_i \right)}{\sum_{i=1}^n x_i^2 - \frac{1}{n} \left(\sum_{i=1}^n x_i \right)^2} \end{aligned} \quad (6)$$

在 (5) 式和 (6) 式中， $\bar{x}$ 和 $\bar{y}$ 分别为 x_i 和 y_i ($i = 1, 2, \dots, n$) 的平均值，即

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i, \quad \bar{y} = \frac{1}{n} \sum_{i=1}^n y_i。$$

建立一元线性回归模型的过程，就是用变量 x_i 和 y_i 的实际观测数据确定参数 a 和 b 的最小二乘估计值 $\hat{a}$ 和 $\hat{b}$ 的过程。现在，我们用表 2-1 中的数据，建立某地区农业总产值 (y) 与粮食总产量 (x) 之间的一元线性回归模型。

回归系数 a 和 b 的拟合值分别为

$$\begin{aligned}\hat{b} &= \frac{L_{xy}}{L_{xx}} = \frac{\sum_{i=1}^{10} x_i y_i - \frac{1}{10} \left(\sum_{i=1}^{10} x_i \right) \left(\sum_{i=1}^{10} y_i \right)}{\sum_{i=1}^{10} x_i^2 - \frac{1}{10} \left(\sum_{i=1}^{10} x_i \right)^2} \\ &= \frac{837.1 - \frac{1}{10} \times 240 \times 33.6}{6120 - \frac{1}{10} \times (240)^2} = 0.085278 \\ \hat{a} &= \bar{y} - \hat{b} \bar{x} \\ &= \frac{1}{10} \times 33.6 - 0.085278 \times \frac{1}{10} \times 240 = 1.313328\end{aligned}$$

故该地区农业总产值 (y) 与粮食总产量 (x) 之间的回归方程为

$$\hat{y} = 1.313328 + 0.085278x \quad (7)$$

(二)一元线性回归模型显著性检验

回归模型建立之后, 需要对模型的可信度进行检验, 以鉴定模型的质量。线性回归方程的显著性检验是借助于 F 检验来完成的。

在回归分析中, y 的 n 次观测值 $y_1, y_2, \dots, y_n$ 之间的差异, 可用观测值 y_i 与其平均值 $\bar{y}$ 的离差平方和来表示, 它被称为总的离差平方和, 记为

$$S_{\text{总}} = L_{yy} = \sum_{i=1}^n (y_i - \bar{y})^2 \quad (8)$$

可以证明

$$\begin{aligned}S_{\text{总}} &= L_{yy} = \sum_{i=1}^n (y_i - \bar{y})^2 \\ &= \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \sum_{i=1}^n (\hat{y}_i - \bar{y})^2 \\ &= Q + U\end{aligned} \quad (9)$$

(9) 式中, $Q = \sum_{i=1}^n (y_i - \hat{y}_i)^2$ 称为误差平方和, 或剩余平方和, 而

$$\begin{aligned}U &= \sum_{i=1}^n (\hat{y}_i - \bar{y})^2 = \sum_{i=1}^n (a + bx_i - a - b\bar{x})^2 \\ &= b^2 \sum_{i=1}^n (x_i - \bar{x})^2 = b^2 L_{xx} = bL_{xy}\end{aligned}$$

称为回归平方和。

由 (9) 式可以看出, 当 U 对 L_{yy} 的贡献越大时, Q 的影响就越小, 回归模型的效果就越好。这样, 就可以由统计量 $F = U / \frac{Q}{n-2}$

$$\quad (10)$$

衡量回归模型的效果, 显然 F 越大, 就意味着模型的效果越佳。事实上, 统计量 $F \sim F(1, n-2)$ 。在显著水平 α 下, 若 $F > F_{\alpha}(1, n-2)$, 则认为回归方程效果在此水平下显著。一般地, 当 $F < F_{0.10}(1, n-2)$ 时, 则认为方程效果不明显。对于回归方程 (7) 式, 我们有

$$S_{\text{总}} = L_{yy} = \sum_{i=1}^{10} (y_i - \bar{y})^2 = 3.0840$$

$$U = bL_{xy} = 0.085278 \times 30.7000 = 2.6180346$$

$$Q = S_{\text{总}} - U = 3.0840 - 2.6180346 = 0.4659654$$

所以

$$F = U / \frac{Q}{n-2} = \frac{2.6180346}{0.4659654 \times \frac{1}{8}} = 44.948137$$

在置信水平 $\alpha=0.01$ 下查 F 分布表得 $F_{0.01}(1,8)=11.6$ 。由于 $F=44.948137 > F_{0.01}(1,8)=11.6$ ，所以回归方程 (7) 式在置信水平 $\alpha=0.01$ 下是显著的。

二、多元线性回归模型

在多要素的地理系统中，除了在某两个要素之间存在着相互作用和影响而发生某种相关外，在若干个（多于两个）要素之间也存在着相关影响、相互关联的情况。因此，多元地理回归模型更带有普遍性的意义。

（一）多元线性回归模型的建立

假设某一因变量 y 受 k 个自变量 $x_1, x_2, \dots, x_k$ 的影响，其 n 组观测值为 $(y_a, x_{a1}, x_{a2}, \dots, x_{ak})$ ， $a=1, 2, \dots, n$ 。那么，多元线性回归模型的结构形式为：

$$y_a = \beta_0 + \beta_1 x_{a1} + \beta_2 x_{a2} + \dots + \beta_k x_{ak} + \varepsilon_a \quad (11)$$

在 (11) 式中， $\beta_0, \beta_1, \dots, \beta_k$ 为待定参数， ε_a 为随机变量。如果 $b_0, b_1, \dots, b_k$ 分别为 $\beta_0, \beta_1, \beta_2, \dots, \beta_k$ 的拟合值，则得回归方程

$$\hat{y} = b_0 + b_1 x_1 + b_2 x_2 + \dots + b_k x_k \quad (12)$$

在 (12) 式中， b_0 为常数， $b_1, b_2, \dots, b_k$ 被称为偏回归系数。偏回归系数 b_i ($i=1, 2, \dots, k$) 的意义是，当其它自变量 x_j ($j \neq i$) 都固定时，自变量 x_i 每变化一个单元而使因变量 y 平均改变的数值。

根据最小二乘法原理， β_i ($i=0, 1, 2, \dots, k$) 的估计值 b_i ($i=0, 1, 2, \dots, k$) 要使

$$\begin{aligned} Q &= \sum_{a=1}^n (y_a - \hat{y}_a)^2 \\ &= \sum_{a=1}^n [y_a - (b_0 + b_1 x_{a1} + b_2 x_{a2} + \dots + b_k x_{ak})]^2 \\ &\rightarrow \min \end{aligned}$$

由求极值的必要条件得

$$\begin{cases} \frac{\partial Q}{\partial b_0} = -2 \sum_{a=1}^n (y_a - \hat{y}_a) = 0 \\ \frac{\partial Q}{\partial b_j} = -2 \sum_{a=1}^n (y_a - \hat{y}_a) x_{aj} = 0 (j=1, 2, \dots, k) \end{cases}$$

方程组 (14) 式经展开整理后得

$$\left\{ \begin{aligned} &nb_0 + \left(\sum_{a=1}^n x_a\right)b_1 + \left(\sum_{a=1}^n x_{a_2}\right)b_2 + \cdots + \left(\sum_{a=1}^n x_{a_k}\right)b_k = \sum_{a=1}^n y_a \\ &\left(\sum_{a=1}^n x_{a_1}\right)b_0 + \left(\sum_{a=1}^n x_{a_1}^2\right)b_1 + \left(\sum_{a=1}^n x_{a_1}x_{a_2}\right)b_2 + \cdots + \left(\sum_{a=1}^n x_{a_1}x_{a_2}\right)b_k = \left(\sum_{a=1}^n x_a y_a\right) \\ &\left(\sum_{a=1}^n x_{a_2}\right)b_0 + \left(\sum_{a=1}^n x_{a_1}x_{a_2}\right)b_1 + \sum_{a=1}^n (x_{a_2}^2)b_2 + \cdots + \left(\sum_{a=1}^n x_{a_1}x_{a_k}\right)b_k = \sum_{a=1}^n x_{a_2}y_a \\ &\dots\dots\dots \\ &\left(\sum_{a=1}^n x_{a_k}\right)b_0 + \left(\sum_{a=1}^n x_a x_{a_k}\right)b_1 + \left(\sum_{a=1}^n x_{a_2}x_{a_k}\right)b_2 + \cdots + \left(\sum_{a=1}^n x_{a_k}^2\right)b_k = \sum_{a=1}^n x_{a_k}y_a \end{aligned} \right. \quad (15)$$

方程组 (15) 式称为正规方程组。如果引入以下矩阵：

$$\begin{aligned}
X &= \begin{pmatrix} 1 & x_{11} & x_{12} & \cdots & x_{1k} \\ 1 & x_{21} & x_{22} & \cdots & x_{2k} \\ 1 & x_{31} & x_{32} & \cdots & x_{3k} \\ \dots & \dots & \dots & \dots & \dots \\ 1 & x_{n1} & x_{n2} & \cdots & x_{nk} \end{pmatrix} \\
A = X^T X &= \begin{pmatrix} 1 & 1 & 1 & \cdots & 1 \\ x_{11} & x_{21} & x_{31} & \cdots & x_{n1} \\ x_{12} & x_{22} & x_{32} & \cdots & x_{n2} \\ \dots & \dots & \dots & \dots & \dots \\ x_{1k} & x_{2k} & x_{3k} & \cdots & x_{nk} \end{pmatrix} \cdot \begin{pmatrix} 1 & x_{11} & x_{12} & \cdots & x_{1k} \\ 1 & x_{21} & x_{22} & \cdots & x_{2k} \\ 1 & x_{31} & x_{32} & \cdots & x_{3k} \\ \dots & \dots & \dots & \dots & \dots \\ 1 & x_{n1} & x_{n2} & \cdots & x_{nk} \end{pmatrix} \\
&= \begin{pmatrix} n & \sum_{a=1}^n x_{a1} & \sum_{a=1}^n x_{a2} & \cdots & \sum_{a=1}^n x_{ak} \\ \sum_{a=1}^n x_{a2} & \sum_{a=1}^n x_{a1}^2 & \sum_{a=1}^n x_{a1} x_{a2} & \cdots & \sum_{a=1}^n x_{a1} x_{ak} \\ \sum_{a=1}^n x_{a2} & \sum_{a=1}^n x_{a1} x_{a2} & \sum_{a=1}^n x_{a2}^2 & \cdots & \sum_{a=1}^n x_{a2} x_{ak} \\ \dots & \dots & \dots & \dots & \dots \\ \sum_{a=1}^n x_{ak} & \sum_{a=1}^n x_{a1} x_{ak} & \sum_{a=1}^n x_{a2} x_{ak} & \cdots & \sum_{a=1}^n x_{ak}^2 \end{pmatrix} \\
Y &= \begin{pmatrix} y_1 \\ y_2 \\ \dots \\ y_n \end{pmatrix} \qquad b = \begin{pmatrix} b_0 \\ b_1 \\ b_2 \\ \dots \\ b_n \end{pmatrix}
\end{aligned}$$

表 2-7 某地区城市公共交通营运额、人口数及
工农业总产值的年平均数据

城市序号	公共交通营运额(y) (千人公里)	人口数(x_1) (千人)	工农业总产值(x_2) (千万元)
1	6825.99	1298.00	437.26
2	512.00	119.80	1283.48
3	1902.00	344.28	1128.33
4	146.00	235.56	600.58
5	2824.00	163.79	783.15
6	37.00	76.72	65.26
7	52.00	17.81	441.26
8	56.00	30.66	242.33
9	187.00	15.92	23.98
10	1065.00	345.08	371.98
11	107.00	6.70	324.40
12	173.00	28.00	262.11
13	771.00	75.00	1508.16
14	192.00	12.47	1072.27

据表 2-7 中的数据，我们有

$$X = \begin{pmatrix} 1 & x_{11} & x_{12} \\ 1 & x_{21} & x_{22} \\ 1 & x_{31} & x_{32} \\ M & M & M \\ 1 & x_{14,1} & x_{14,2} \end{pmatrix} = \begin{pmatrix} 1 & 1298.00 & 437.26 \\ 1 & 119.80 & 1283.48 \\ 1 & 344.28 & 1128.33 \\ M & M & M \\ 1 & 12.47 & 1508.16 \end{pmatrix}$$

$$Y = \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ M \\ y_{14} \end{pmatrix} = \begin{pmatrix} 6825.00 \\ 512.00 \\ 1902.00 \\ M \\ 192.00 \end{pmatrix}$$

经过计算可得

$$b = \begin{pmatrix} b_0 \\ b_1 \\ b_2 \end{pmatrix} = (X^T X)^{-1} X^T Y = \begin{pmatrix} -172.2415 \\ 5.1075 \\ 0.3636 \end{pmatrix}$$

故 y 与 x_1 及 x_2 之间的线性回归方程为

$$\hat{y} = -172.2415 + 5.1075x_1 + 0.3636x_2 \quad (17)$$

(二)多元线性回归模型的显著性检验

与一元线性回归模型一样，当多元线性回归模型建立以后，也需要进行显著性检验。

与前面的一元线性回归分析一样，因变量 y 的观测值 $y_1, y_2, \dots, y_n$ 之间

的波动或差异，是由两个因素引起的，一是由于自变量 $x_1, x_2, \dots, x_k$ 的取值不同，另一是受其它随机因素的影响而引起的。为了从 y 的总变差中把它们区分开来，就需要对回归模型进行方差分析，也就是将 y 的总的离差平方和 $S_{\text{总}}$ (或 L_{yy}) 分解成两个部分，即回归平方和 U 和剩余平方和 Q ：

$$S_{\text{总}} = L_{yy} = U + Q$$

在多元线性回归分析中，回归平方和表示的是所有 k 个自变量对 y 的变差的总影响，它可以按公式

$$U = \sum_{i=1}^n (\hat{y}_i - \bar{y})^2 = \sum_{i=1}^k b_i L_{iy}$$

计算，而剩余平方和为

$$Q = \sum_{a=1}^n (y_a - \hat{y}_a)^2 = L_{yy} - U$$

以上几个公式与一元线性回归分析中的有关公式完全相似。它们所代表的意义也相似，即回归平方和越大，则剩余平方和 Q 就越小，回归模型的效果就越好。不过，在多元线性回归分析中，各平方和的自由度略有不同，回归平方和 U 的自由度等于自变量的个数 K ，而剩余平方和的自由度等于 $n-K-1$ ，所以 F 统计量为

$$F = \frac{U / K}{Q / (n - K - 1)} \quad (18)$$

当统计量 F 计算出来之后，就可以查 F 分布表对模型进行显著性检验。在上例中，计算可得

$$\begin{aligned} S_{\text{总}} = L_{yy} &= \sum_{a=1}^{14} (y_a - \bar{y})^2 = 44521048.53 \\ U &= b_1 L_{1y} + b_2 L_{2y} = 39030046.11 \\ Q &= S_{\text{总}} - U = 5491002.42 \end{aligned}$$

故

$$F = \frac{U / K}{Q / (n - K - 1)} = \frac{U / 2}{Q / 11} = \frac{39030046.11 / 2}{5491002.42 / 11} = 39.094$$

在置信水平 $\alpha=0.01$ 下查 F 分布表知： $F_{0.01}(2, 11)=7.21$ 。由于 $F=39.094 > F_{0.01}(2, 11)=7.21$ ，所以在置信水平 $\alpha=0.01$ 下，回归方程 (17) 式是显著的。

三、非线性回归模型的建立方法

在复杂地理系统中，除了线性关系以外，要素之间的非线性关系也是大量存在的。因此，对非线性回归分析，也有必要作一些介绍。

(一) 非线性关系的线性化

前面已经讨论了线性回归模型的建立方法。在复杂地理系统研究中，对于要素之间的非线性关系，若能找到某种途径将其转化为线性关系，我们就可以借助于线性回归模型的建立方法，建立要素之间的非线性回归模型。事实上，这是可以办得到的，只要根据要素之间的关系设定新的变量，通过变量替换就可以将原来的非线性关系转化为新变量下的线性关系。譬如：

(1) 对于指数曲线 $y = de^{bx}$ ，令 $y' = \ln y$ ， $x' = x$ ，就可以将其转化为直线形式： $y' = a + bx'$ ，其中， $a = \ln d$ ；

(2)对于对数曲线 $y=a+b\ln x$, 令 $y'=y$, $x'=\ln x$, 就可以将其转化为直线形式: $y'=a+bx'$;

(3)对于幂函数曲线 $y=dx^b$, 令 $y'=\ln y$, $x'=x$, 就可以将其转化为直线形式: $y'=a+bx'$, 其中, $a=\ln d$;

(4) 对于双曲线 $\frac{1}{y}=a+\frac{b}{x}$, 令 $y'=\frac{1}{y}$, $x'=\frac{1}{x}$, 就可以将其转化为直线形式: $y'=a+bx'$;

(5) 对于S型曲线 $y=\frac{1}{a+be^{-x}}$, 令 $y'=\frac{1}{y}$, $x'=e^{-x}$, 就可以将其转化为直线形式: $y'=a+bx'$;

(6)对于幂函数乘积:

$$y = dx_1^{\beta_1} \cdot x_2^{\beta_2} \cdots x_k^{\beta_k}$$

只要令 $y' = \ln y$, $x'_1 = \ln x_1$, $x'_2 = \ln x_2$, $\cdots$, $x'_k = \ln x_k$, 就可以将其转化为直线形式:

$$y' = \beta_0 + \beta_1 x'_1 + \beta_2 x'_2 + \cdots + \beta_k x'_k$$

上式中, $\beta_0 = \ln d$;

(7)对于对数函数和: $y = \beta_0 + \beta_1 \ln x_1 + \beta_2 \ln x_2 + \cdots + \beta_k \ln x_k$

只要令 $y' = y$, $x'_1 = \ln x_1$, $x'_2 = \ln x_2$, $\cdots$, $x'_k = \ln x_k$, 就可以将其化为线性形式:

$$y' = \beta_0 + \beta_1 x'_1 + \beta_2 x'_2 + \cdots + \beta_k x'_k$$

以上这种将非线性函数关系转化为线性关系的过程称为非线性关系的线性处理。不过, 需要强调指出的是, 这种转化过程并不能保证函数关系中变量个数不变。譬如, 对于两变量的多项式

$$y = \beta_0 + \beta_1 x + \beta_2 x^2 + \cdots + \beta_k x^k$$

若令 $x'_1 = x$, $x'_2 = x^2$, $\cdots$, $x'_k = x^k$, $y' = y$, 则它就被转化为多变量的线性模型:

$$y' = \beta_0 + \beta_1 x'_1 + \beta_2 x'_2 + \cdots + \beta_k x'_k$$

(二)非线性回归模型建立的实例

通过上述分析, 我们可以得到建立非线性回归模型的一般方法: 首先通过适当的变量替换将非线性关系线性化, 然后再用线性回归分析方法建立新变量下的线性回归模型, 通过新变量之间的线性相关关系反映原来变量之间的非线性相关关系。下面, 我们结合实例, 说明非线性地理回归模型的建立过程。

例如, 黄土高原某地区 1984—1990 年期间, 小麦亩产量 (y) 与化肥使用量 (x_1), 以及农家肥 (干纯粪) 使用量 (x_2) 的数据如表 2-8 所示。试建立 y 与 x_1 及 x_2 之间的相关关系模型。

表 2-8 某地区小麦亩产量与化肥、农家肥使用量 (千克/亩)

年份	1984	1985	1986	1987	1988	1989	1990
序号	1	2	3	4	5	6	7
小麦亩产量 (y)	116.0	123.5	123.0	166.6	118.5	197.0	153.0
化肥使用量 (x ₁)	2.21	3.96	3.77	4.28	4.00	7.32	8.66
农家肥使用量 (x ₂)	108.9	127.4	110.1	121.5	137.4	139.7	130.0

从表 2-8 可以看出, 小麦亩产量 (y) 随着化肥使用量 (x₁) 及农家肥使用量 (x₂) 的增加而增加, 但肥料投入量的增长速度越来越高于小麦亩产量的增长速度, 其间的关系可用对数变化规律来模拟, 即

$$y = \beta_0 + \beta_1 \ln x_1 + \beta_2 \ln x_2 + \varepsilon \quad (19)$$

在 (19) 式中, 若令 $y' = y$, $x_1' = \ln x_1$, $x_2' = \ln x_2$, 则它可以被化为线性形式

$$y' = \beta_0 + \beta_1 x_1' + \beta_2 x_2' + \varepsilon \quad (19')$$

变量替换后, 各新变量对应的观测数据如表 2-9 所示。

根据表 2-9 中的数据, 计算可得: $\bar{x}_1' = \frac{1}{7} \sum_{a=1}^7 x_{a1}' = 1.4980$; $\bar{x}_2' = \frac{1}{7} \sum_{a=1}^7 x_{a2}' = 4.8241$; $\bar{y}' = \frac{1}{7} \sum_{a=1}^7 y_a' = 142.500$; 以及 $L_{11}' = \sum_{a=1}^7 (x_{a1}' - \bar{x}_1')^2 = 1.23464$;

表 2-9 变量替换后各新变量的对应数据

$$L_{12}' = L_{21}' = \sum_{a=1}^7 (x_{a1}' - \bar{x}_1')(x_{a2}' - \bar{x}_2') = 0.1879;$$

$$L_{22}' = \sum_{a=1}^7 (x_{a2}' - \bar{x}_2')^2 = 0.05903;$$

$$L_{1y}' = \sum_{a=1}^7 (x_{a1}' - \bar{x}_1')(y_a' - \bar{y}') = 59.7337;$$

$$L_{2y}' = \sum_{a=1}^7 (x_{a2}' - \bar{x}_2')(y_a' - \bar{y}') = 9.2856。$$

所以, 正规方程组为

$$\begin{cases} 1.23464b_1 + 0.1879b_2 = 59.7337 \\ 0.1879b_1 + 0.05903b_2 = 9.2856 \\ b_0 = 142.5 - 1.4980b_1 - 4.8241b_2 \end{cases} \quad (20)$$

解上述正规方程组 (20) 式可得

$$\begin{cases} b_0 = 40.64341 \\ b_1 = 47.388 \\ b_2 = 6.39899 \end{cases}$$

因此, (19') 式所对应的线性回归方程为

$$y' = 40.64341 + 47.388x_1' + 6.398899x_2' \quad (21)$$

而对应于 (19) 式的非线性回归方程为:

$$y=40.64341+47.388\ln x_1+6.39899\ln x_2 \quad (22)$$

第三节 系统聚类分析方法

聚类分析，亦称群分析或点群分析，它是研究多要素事物分类问题的数量方法。其基本原理是，根据样本自身的属性，用数学方法按照某些相似性或差异性指标，定量地确定样本之间的亲疏关系，并按这种亲疏关系程度对样本进行聚类。

聚类分析方法，是地理学中研究地理事物分类问题和地理分区问题的重要的数量分析方法。常见的聚类分析方法有系统聚类法、动态聚类法和模糊聚类法等。本节，我们将结合有关实例，主要介绍和探讨系统聚类分析方法在地理学研究中的应用问题。

一、聚类要素的数据处理

在聚类分析中，聚类要素的选择是十分重要的，它直接影响分类结果的准确性和可靠性。在地理分类和分区研究中，被聚类的对象常常是多个要素构成的。不同要素的数据往往具有不同的单位和量纲，因而其数值的差异可能是很大的，这就会对分类结果产生影响。因此当分类要素的对象确定之后，在进行聚类分析之前，还要对聚类要素进行数据处理。

假设有 m 个被聚类的对象，每一个被聚类对象都有 $x_1, x_2, \dots, x_n$ 个要素构成。它们所对应的要素数据可用表 2-10 给出。在聚类分析中，常用的聚类要素的数据处理方法有如下几种。

聚类对象	要 素					
	x_1	x_2	$\cdots$, x_j	$\cdots$,	x_n	
1	x_{11}	x_{12}	$\cdots$,	x_{1j}	$\cdots$,	x_{1n}
2	x_{21}	x_{22}	$\cdots$,	x_{2j}	$\cdots$,	x_{2n}
M	M	M		M		M
i	x_{i1}	x_{i2}	$\cdots$,	x_{ij}	$\cdots$,	x_{in}
M	M	M		M		M
m	x_{m1}	x_{m2}	$\cdots$,	x_{mj}	$\cdots$,	x_{mn}

(1)总和标准化。分别求出各聚类要素所对应的数据的总和，以各要素的数据除以该要素数据的总和，即

$$x'_{ij} = x_{ij} / \sum_{i=1}^m x_{ij} \quad \begin{pmatrix} i = 1, 2, \dots, m \\ j = 1, 2, \dots, n \end{pmatrix} \quad (1)$$

这种标准化方法所得的新数据 x'_{ij} 满足

$$\sum_{i=1}^m x'_{ij} = 1 \quad (j = 1, 2, \dots, n)$$

(2)标准差的标准化，即

$$x'_{ij} = \frac{x_{ij} - \bar{x}_j}{s_j} \quad \begin{pmatrix} i = 1, 2, \dots, m \\ j = 1, 2, \dots, n \end{pmatrix} \quad (2)$$

在 (2)式中，

$$\bar{x}_j = \frac{1}{m} \sum_{i=1}^m x_{ij}, \quad S_j = \sqrt{\frac{1}{m} \sum_{i=1}^m (x_{ij} - \bar{x}_j)^2}$$

由这种标准化方法所得的新数据 x'_{ij} ，各要素的平均值为0，标准差为 1，即有

$$\bar{x}'_j = \frac{1}{m} \sum_{i=1}^m x'_{ij} = 0, \quad S_j = \sqrt{\frac{1}{m} \sum_{i=1}^m (x'_{ij} - \bar{x}'_j)^2} = 1$$

(3)极大值标准化，即

$$x'_{ij} = \frac{x_{ij}}{\max\{x_{ij}\}} \quad \left(\begin{matrix} i = 1, 2, \cdots, m \\ j = 1, 2, \cdots, n \end{matrix} \right) \quad (3)$$

经过这种标准化所得的新数据，各要素的极大值为 1，其余各数值小于 1。

(4)极差的标准化，即

$$x_{ij} = \frac{x_{ij} - \min\{x_{ij}\}}{\max_i\{x_{ij}\} - \min\{x_{ij}\}} \quad \left(\begin{matrix} i = 1, 2, \cdots, m \\ j = 1, 2, \cdots, n \end{matrix} \right) \quad (4)$$

经过这种标准化所得的新数据，各要素的极大值为 1，极小值为 0，其余的数值均在 0 与 1 之间。

表 2-11 给出了某地区九个农业区的七项经济指标,其极差标准化处理后的数据如表 2-12 所示。

表 2-11 某地区九个农业区的七项经济指标数据

区代号	指 标					
	人均耕地 x1 (亩/人)	劳均耕地 x2 (亩/个)	水田比重 x3 (%)	复种指数 x4 (%)	粮食亩产 x5 (公斤/亩)	人均粮食 x6 (公斤/人)
G ₁	4.41	16.40	5.63	113.60	300.70	1036.40
G ₂	4.72	14.57	0.39	95.10	184.90	683.70
G ₃	1.84	4.74	5.28	148.50	462.30	611.10
G ₄	2.69	7.91	0.39	111.00	297.20	632.60
G ₅	1.22	3.18	72.04	217.80	816.60	791.10
G ₆	1.23	3.16	43.78	179.60	598.20	636.50
G ₇	1.12	2.72	65.15	194.70	712.60	634.30
G ₈	4.40	9.99	5.35	94.90	245.30	771.70
G ₉	2.50	6.21	2.90	94.80	282.10	574.60

表 2-12 极差标准化处理后的数据

区代号	指 标						
	x_1	x_2	x_3	x_4	x_5	x_6	x_7
G_1	0.91	1.00	0.07	0.15	0.18	1.00	0.14
G_2	1.00	0.87	0.00	0.00	0.00	0.24	0.00
G_3	0.20	0.15	0.07	0.44	0.44	0.08	0.07
G_4	0.44	0.38	0.00	0.13	0.18	0.13	0.00
G_5	0.03	0.03	1.00	1.00	1.00	0.45	1.00
G_6	0.03	0.03	0.61	0.69	0.65	0.13	0.59
G_7	0.00	0.00	0.90	0.81	0.84	0.13	1.00
G_8	0.91	0.53	0.07	0.00	0.10	0.43	0.09
G_9	0.38	0.26	0.04	0.00	0.15	0.00	0.00

二、距离和相似系数的计算

距离是事物之间差异性的测度，而相似系数则是其相似性的测度，所以距离和相似系数是聚类分析的依据和基础。当聚类要素的数据处理工作完成以后，就要计算分类对象之间的距离或相似系数，并依据距离或相似系数的矩阵结构进行聚类。

(一)距离的计算

如果我们把每一个分类对象的 n 个聚类要素看成 n 维空间的 n 个坐标轴，则每一个分类对象的 n 个要素所构成的 n 维数据向量就是 n 维空间中的一个点。这样，各分类对象之间的差异性就可以由它们所对应的 n 维空间中点之间的距离度量。常用的距离有

(1)绝对值距离

$$d_{ij} = \sum_{k=1}^n |x_{ik} - x_{jk}| \quad (i, j = 1, 2, \dots, m) \quad (5)$$

(2)欧氏距离

$$d_{ij} = \sqrt{\sum_{k=1}^n (x_{ik} - x_{jk})^2} \quad (i, j = 1, 2, \dots, m) \quad (6)$$

(3)明科夫斯基距离

$$d_{ij} = \left[\sum_{k=1}^n |x_{ik} - x_{jk}|^p \right]^{\frac{1}{p}} \quad (i, j = 1, 2, \dots, m) \quad (7)$$

(7)式中， $p \geq 1$ 。当 $p=1$ 时，它就是绝对值距离；当 $p=2$ 时，它就是欧氏距离。

(4)切比雪夫距离。当明科夫斯基距离 $p \rightarrow \infty$ 时，有

$$d_{ij} = \max |x_{ik} - x_{jk}| \quad (i, j = 1, 2, \dots, m) \quad (8)$$

选择不同的距离，聚类结果会有所差异。在地理分区和分类研究中，往往采用几种距离进行计算、对比，选择一种较为合理的距离进行聚类。

据表 2-12 中的数据，用公式 (5) 式计算可得九个农业区之间的绝对值距离矩阵如下：

$$D = (d_{ij})_{9 \times 9} = \begin{pmatrix} 0 & & & & & & & & \\ 1.52 & 0 & & & & & & & \\ 3.10 & 2.70 & 0 & & & & & & \\ 2.19 & 1.47 & 1.23 & 0 & & & & & \\ 5.86 & 6.02 & 3.64 & 4.77 & 0 & & & & \\ 4.72 & 4.46 & 1.86 & 2.99 & 1.78 & 0 & & & \\ 5.79 & 5.53 & 2.93 & 4.06 & 0.83 & 1.07 & 0 & & \\ 1.32 & 0.88 & 2.24 & 1.29 & 5.14 & 3.96 & 5.03 & 0 & \\ 2.62 & 1.66 & 1.20 & 0.51 & 4.84 & 3.06 & 3.32 & 1.40 & 0 \end{pmatrix} \quad (9)$$

(二)相似系数的计算

常见的相似系数是夹角余弦和相关系数，其计算公式如下：

(1)夹角余弦：

$$r_{ij} = \cos \theta_{ij} = \frac{\sum_{k=1}^n (x_{ik} x_{jk})}{\sqrt{\sum_{k=1}^n x_{ik}^2} \sqrt{\sum_{k=1}^n x_{jk}^2}} \quad (i, j = 1, 2, \dots, m) \quad (10)$$

在 (10) 式中，显然有： $-1 \leq \cos \theta_{ij} \leq 1$ 。

(2)相关系数：

$$r_{ij} = \frac{\sum_{k=1}^n (x_{ik} - \bar{x}_i)(x_{jk} - \bar{x}_j)}{\sqrt{\sum_{k=1}^n (x_{ik} - \bar{x}_i)^2} \sqrt{\sum_{k=1}^n (x_{jk} - \bar{x}_j)^2}} \quad (i, j = 1, 2, \dots, m) \quad (11)$$

在 (11) 式中， $\bar{x}_i$ 和 $\bar{x}_j$ 分别为聚类对象 i 和 j 各要素标准化数据的平均值。

据表 2-12 中的数据，用夹角余弦公式 (10) 式计算，可得如下的相似系数矩阵：

$$R = (r_{ij})_{9 \times 9} = \begin{pmatrix} 1 & & & & & & & & \\ 0.88 & 1 & & & & & & & \\ 0.49 & 0.38 & 1 & & & & & & \\ 0.88 & 0.94 & 0.67 & 1 & & & & & \\ 0.30 & 0.06 & 0.76 & 0.30 & 1 & & & & \\ 0.24 & 0.05 & 0.80 & 0.30 & 0.99 & 1 & & & \\ 0.20 & 0.01 & 0.71 & 0.24 & 0.98 & 0.99 & 1 & & \\ 0.93 & 0.95 & 0.45 & 0.92 & 0.21 & 0.18 & 0.14 & 1 & \\ 0.77 & 0.93 & 0.55 & 0.95 & 0.21 & 0.23 & 0.19 & 0.90 & 1 \end{pmatrix} \quad (12)$$

三、直接聚类法

直接聚类法，是根据距离或相似系数矩阵的结构一次并类得到结果，是

一种简便的聚类方法。它先把各个分类对象单独视为一类，然后根据距离最小或相似系数最大的原则，依次选出一对分类对象，并成新类。如果其中一个分类对象已归于一类，则把另一个也归入该类；如果一对分类对象正好属于已归的两类，则把这两类并为一类。每一次归并，都划去该对象所在的列与列序相同的行。那么，经过 $m-1$ 次就可以把全部分类对象归为一类，这样就可以根据归并的先后顺序作出聚类分析的谱系图。

下面，我们据距离矩阵 (9) 式，用直接聚类法对某地区的九个农业区进行聚类分析。

第一步，在距离矩阵 D 中，除对角线元素以外， $d_{49}=d_{94}=0.51$ 为最小者，故将第 4 区与第 9 区并为一类，划去第 9 行和第 9 列；

第二步，在余下的元素中，除对角线元素以外， $d_{75}=d_{57}=0.83$ 为最小者，故第 5 区与第 7 区并为一类，划掉第 7 行和第 7 列；

第三步，在第二步之后余下的元素之中，除对角线元素以外， $d_{82}=d_{28}=0.88$ 为最小者，故将第 2 区与第 8 区并为一类，划去第 8 行和第 8 列；

第四步，在第三步之后余下的元素中，除对角线元素以外， $d_{43}=d_{34}=1.23$ 为最小者，故将第 3 区与第 4 区并为一类，划去第 4 行和第 4 列，此时，第 3、4、9 区已归并为一类。

第五步，在第四步之后余下的元素中，除对角线元素以外， $d_{21}=d_{12}=1.52$ 为最小者，故将第 1 区与第 2 区并为一类，划去第 2 行与第 2 列，此时，第 1、2、8 区已归并为一类；

第六步，在第五步之后余下的元素中，除对角线元素以外， $d_{65}=d_{56}=1.78$ 为最小者，故将第 5 区与第 6 区并为一类，划去第 6 行和第 6 列，此时，第 5、6、7 区已归并为一类；

第七步，在第六步之后余下的元素中，除对角线元素以外， $d_{31}=d_{13}=3.10$ 为最小者，故将第 1 区与第 3 区并为一类，划去第 3 行和第 3 列，此时，第 1, 2, 3, 4, 8, 9 区已归并为一类。

第八步，在第七步之后余下的元素中，除去对角线元素以外，只有 $d_{51}=d_{15}=5.86$ ，故将第 1 区与第 5 区并为一类，划去第 5 行和第 5 列，此时，第 1, 2, 3, 4, 5, 6, 7, 8, 9 区均归并为一类。

根据上述步骤，我们可以作出聚类过程的谱系图 (图 2-1)。直接聚类法虽然简便，但在归类过程中是划去行和列的，因而难免有信息损失。因此直接聚类法并不是最好的系统聚类法。

四、最短距离聚类法

最短距离法，是在原来的 $m \times m$ 距离矩阵的非对角元素中找出 $d_{pq} = \min \{d_{ij}\}$ ，把分类对象 G_p 和 G_q 归并为一新类 G_r ，然后按计算公式：

$$d_{rk} = \min \{d_{pk}, d_{qk}\} \quad (k \neq p, q) \quad (13)$$

计算原来各类与新类之间的距离，这样就得到一个新的 $(m-1)$ 阶的距离矩阵；再从新的距离矩阵中选出最小的 d_{ij} ，把 G_i 和 G_j 归并成新类；再计算各类与新类的距离，这样一直下去，直至各分类对象被归为一类为止。

以下，我们据 (9) 式中的距离矩阵，用最短距离聚类法对某地区的九个农业区进行聚类分析。

第一步，在 9×9 阶距离矩阵 D 中，非对角元素中最小者是 $d_{94}=0.51$ ，故首先将第 4 区与第 9 区并为一类，记为 G_{10} ，即 $G_{10} = \{G_4, G_9\}$ 。分别按照公式 (13) 式计算 $G_1, G_2, G_3, G_5, G_6, G_7, G_8$ 与 G_{10} 之间的距离得：

$$\begin{aligned} d_{1,10} &= \min \{d_{14}, d_{19}\} = \min \{2.19, 2.62\} = 2.19 \\ d_{2,10} &= \min \{d_{24}, d_{29}\} = \min \{1.47, 1.66\} = 1.47 \\ d_{3,10} &= \min \{d_{34}, d_{39}\} = \min \{1.23, 1.20\} = 1.20 \\ d_{5,10} &= \min \{d_{54}, d_{59}\} = \min \{4.77, 4.84\} = 4.77 \\ d_{6,10} &= \min \{d_{64}, d_{69}\} = \min \{2.99, 3.06\} = 2.99 \\ d_{7,10} &= \min \{d_{74}, d_{79}\} = \min \{4.06, 3.32\} = 3.32 \\ d_{8,10} &= \min \{d_{84}, d_{89}\} = \min \{1.29, 1.40\} = 1.29 \end{aligned}$$

这样就得到 $G_1, G_2, G_3, G_5, G_6, G_7, G_8, G_{10}$ 上的一个新的 8×8 阶距离矩阵：

G_1	G_2	G_3	G_5	G_6	G_7	G_8	G_{10}	
G_1	0							
G_2	1.52	0						
G_3	3.10	2.70	0					
G_5	5.86	6.02	3.64	0				
G_6	4.72	4.46	1.86	1.78	0			
G_7	5.79	5.53	2.93	0.83	1.07	0		
G_8	1.32	0.88	2.24	5.14	3.96	5.03	0	
G_{10}	2.19	1.47	1.20	4.77	2.99	3.32	1.29	0

第二步，在上一步骤中所得到的新的 8×8 阶距离矩阵中，非对角元素中最小者为 $d_{57}=0.83$ ，故将 G_5 与 G_7 归并为一类，记为 G_{11} ，即 $G_{11} = \{G_5, G_7\}$ 。再分别按照公式 (13) 式计算 $G_1, G_2, G_3, G_6, G_8, G_{10}$ 与 G_{11} 之间的距离，可得到一个新的 7×7 阶距离矩阵：

G_1	G_2	G_3	G_6	G_8	G_{10}	G_{11}
G_1	0					
G_2	1.52	0				
G_3	3.10	2.70	0			
G_6	4.72	4.46	1.86	0		
G_8	1.32	0.88	2.24	3.96	0	
G_{10}	2.19	1.47	1.20	2.99	1.20	0
G_{11}	5.79	5.53	2.93	1.07	5.03	3.32

第三步，在第二步所得到的新的 7×7 阶距离矩阵中，非对角线元素中最小者为 $d_{28}=0.88$ ，故将 G_2 与 G_8 归并为一类，记为 G_{12} ，即 $G_{12} = \{G_2, G_8\}$ 。再分别按公式 (13) 式计算 $G_1, G_3, G_6, G_{10}, G_{11}$ 与 G_{12} 之间的距离，可得到一个新的 6×6 阶距离矩阵：

G_1	G_3	G_6	G_{10}	G_{11}	G_{12}
G_1	0				

G_3 3.10 0
 G_6 4.72 1.86 0
 G_{10} 2.19 1.20 2.99 0
 G_{11} 5.79 2.93 1.07 3.32 0
 G_{12} 1.32 2.24 3.96 1.20 5.03 0

第四步，在第三步中所得到的新的 6×6 阶距离矩阵中，非对角线元素中最小者为 $d_{6,11}=1.07$ ，故将 G_6 和 G_{11} 归并为一类，记为 G_{13} ，即 $G_{13}=\{G_6, G_{11}\}=\{G_6, (G_5, G_7)\}$ 。再按照公式 (13) 式计算 G_1, G_3, G_{10}, G_{12} 与 G_{13} 之间的距离，可得一个新的 5×5 阶距离矩阵：

	G_1	G_3	G_{10}	G_{12}	G_{13}
G_1	0				
G_3	3.10	0			
G_{10}	2.19	1.20	0		
G_{12}	1.32	2.24	1.20	0	
G_{13}	4.72	1.86	2.99	3.96	0

第五步，在第四步中所得到的新的 5×5 阶距离矩阵中，非对角线元素中最小者为 $d_{3,10}=1.20$ ，故将 G_3 和 G_{10} 归并为一类，记为 G_{14} ，即 $G_{14}=\{G_3, G_{10}\}=\{G_3, (G_4, G_9)\}$ 。再按公式 (13) 式计算 G_1, G_{12}, G_{13} ，与 G_{14} 之间的距离，可得一个新的 4×4 阶的距离矩阵：

	G_1	G_{12}	G_{13}	G_{14}
G_1	0			
G_{12}	1.32	0		
G_{13}	4.72	3.96	0	
G_{14}	2.19	1.20	2.99	0

第六步，在第五步中所得的新的 4×4 阶距离矩阵中，非对角线元素中最小者为 $d_{12,14}=1.20$ ，故将 G_{12} 与 G_{14} 归并为一类，记为 G_{15} ，即 $G_{15}=\{G_{12}, G_{14}\}=\{(G_2, G_8), (G_3, (G_4, G_9))\}$ 。再用公式 (13) 式计算 G_1, G_{13} 与 G_{15} 之间的距离，可得一个新的 3×3 阶距离矩阵：

	G_1	G_{13}	G_{15}
G_1	0		
G_{13}	4.72	0	
G_{15}	1.32	2.99	0

第七步，在第六步中所得的新的 3×3 阶距离矩阵中，非对角线元素中最小者为 $d_{1,15}=1.32$ ，故将 G_1 与 G_{15} 归并为一类，记为 G_{16} ，即 $G_{16}=\{G_1, G_{15}\}=\{G_1, (G_2, G_8), (G_3, (G_4, G_9))\}$ 。再用公式 (13) 式计算 G_{13} 与 G_{16} 之间的距离，可得一个新的 2×2 阶距离矩阵：

	G_{13}	G_{16}
G_{13}	0	
G_{16}	2.99	0

第八步，将 G_{13} 和 G_{16} 归并为一类。此时，所有分类对象均被归并为一类。
综合上述聚类过程，可以作出最短距离聚类谱系图（如图 2-2 所示）。

五、最远距离聚类法

最远距离聚类法与最短距离聚类法的区别在于计算原来的类与新类距离时采用的公式：

$$d_{rk} = \max\{d_{pk}, d_{qk}\} \quad (k \neq p, q) \quad (14)$$

对于上述同样的例子，最远距离聚类法的聚类过程如下：

第一步，在 9×9 阶距离矩阵中，非对角线元素中最小者为 $d_{94}=0.51$ ，故将 G_4 与 G_9 归并为一类，记为 G_{10} ，即 $G_{10}=\{G_4, G_9\}$ 。按照公式 (14) 式分别计算 $G_1, G_2, G_3, G_5, G_6, G_7, G_8$ 与 G_{10} 之间的距离，可以得到一个新的 8×8 阶距离矩阵：

	G_1	G_2	G_3	G_5	G_6	G_7	G_8	G_{10}
G_1	0							
G_2	1.52	0						
G_3	3.10	2.70	0					
G_5	5.86	6.02	3.64	0				
G_6	4.72	4.46	1.86	1.78	0			
G_7	5.79	5.53	2.93	0.83	1.70	0		
G_8	1.32	0.88	2.24	5.14	3.96	5.03	0	
G_{10}	2.62	1.66	1.23	4.84	3.06	4.06	1.40	0

第二步，在第一步所得到的新的 8×8 阶距离矩阵中，非对角线元素中最小者为 $d_{57}=0.83$ ，故将 G_5 与 G_7 归并为一类，记为 G_{11} ，即 $G_{11}=\{G_5, G_7\}$ 。再按照公式 (14) 式分别计算 $G_1, G_2, G_3, G_6, G_8, G_{10}$ 与 G_{11} 之间的距离，可得一个新的 7×7 阶距离矩阵如下：

	G_1	G_2	G_3	G_6	G_8	G_{10}	G_{11}
G_1	0						
G_2	1.52	0					
G_3	3.10	2.70	0				
G_6	4.72	4.46	1.86	0			
G_8	1.32	0.88	2.24	3.96	0		
G_{10}	2.62	1.66	1.23	3.06	1.40	0	
G_{11}	5.86	6.02	3.64	1.78	5.14	4.84	0

第三步，在第二步中所得到的新的 7×7 阶距离矩阵中，非对角线元素中最小者为 $d_{28}=0.88$ ，故将 G_2 与 G_8 并为一类，记为 $G_{12}=\{G_2, G_8\}$ 。再按照公式 (14) 式分别计算 $G_1, G_3, G_6, G_{10}, G_{11}$ 与 G_{12} 之间的距离，可得一个新的 6×6 阶距离矩阵如下：

	G_1	G_3	G_6	G_{10}	G_{11}	G_{12}
G_1	0					
G_3	3.10	0				

G_6	4.72	1.86	0			
G_{10}	2.62	1.23	3.06	0		
G_{11}	5.86	3.64	1.78	4.84	0	
G_{12}	1.52	2.70	4.46	1.66	6.02	0

第四步，在第三步中得到的新的 6×6 阶距离矩阵中，非对角线元素中最小者为 $d_{3,10}=1.23$ ，故将 G_3 和 G_{10} 归并为一类，记为 G_{13} ，即 $G_{13}=\{G_3, G_{10}\}=\{G_3, (G_4, G_9)\}$ 。再按照公式 (14) 式分别计算 G_1, G_6, G_{11}, G_{12} 与 G_{13} 之间的距离，可得一个新的 5×5 阶距离矩阵如下：

	G_1	G_6	G_{11}	G_{12}	G_{13}
G_1	0				
G_6	4.72	0			
G_{11}	5.86	1.78	0		
G_{12}	1.52	4.46	6.02	0	
G_{13}	3.10	3.06	4.84	2.70	0

第五步，在第四步中所得到的新的 5×5 阶距离矩阵中，非对角线元素中最小者为 $d_{1,12}=1.52$ ，故将 G_1 和 G_{12} 归并为一类，记为 G_{14} ，即 $G_{14}=\{G_1, G_{12}\}=\{G_1, (G_2, G_8)\}$ 。再按照公式 (14) 式分别计算 G_6, G_{11}, G_{13} 和 G_{14} 之间的距离，可得一个新的 4×4 距离矩阵如下：

	G_6	G_{11}	G_{13}	G_{14}
G_6	0			
G_{11}	1.78	0		
G_{13}	3.06	4.84	0	

第六步，在第五步中所得到的新的 4×4 阶距离矩阵中，非对角线元素中最小者为 $d_{6,11}=1.78$ ，故将 G_6 与 G_{11} 并为一类，记为 G_{15} ，即 $G_{15}=\{G_6, G_{11}\}=\{G_6, (G_5, G_7)\}$ 。再按照公式 (14) 式分别计算 G_{13}, G_{14} 和 G_{15} 之间的距离，可得一个新的 3×3 阶距离矩阵如下：

	G_{13}	G_{14}	G_{15}
G_{13}	0		
G_{14}	3.10	0	
G_{15}	4.84	6.02	0

第七步，在第六步中所得到的新的 3×3 阶距离矩阵中，非对角线元素中最小者为 $d_{13,14}=3.10$ ，故将 G_{13} 和 G_{14} 归并为一类，记为 G_{16} ，即 $G_{16}=\{G_{13}, G_{14}\}=\{(G_3, (G_4, G_9)), (G_1, (G_2, G_8))\}$ 。再按照公式 (14) 式计算 G_{15} 与 G_{16} 之间的距离，可得一个新的 2×2 阶距离矩阵如下：

	G_{15}	G_{16}
G_{15}	0	
G_{16}	6.02	0

第八步，将 G_{15} 与 G_{16} 归并为一类。此时，各个分类对象均已归并为一类。

综合上述各聚类步骤，可作出最远距离聚类的谱系图（如图 2-3 所示）。

六、系统聚类法计算类之间距离的统一公式

从公式 (13)和 (14)式不难看出，最短距离聚类法具有空间压缩性，而最大距离聚类法具有空间扩张性。它们的这种性质可以形象地用图 2-4 来表示。在图 2-4 中，最短距离为 $d_{AB}=d_{a_1b_1}$ ，最远距离为 $d_{AB}=d_{a_2b_2}$ 。这两种聚类方法关于类之间的距离计算可以用一个统一的式

表 2-13 八种系统聚类方法的距离参数值

方法名称	参 数				矩阵要求	空间性质
	a_p	a_q	β	γ		
最短距离法	$\frac{1}{2}$	$\frac{1}{2}$	0	$-\frac{1}{2}$	各种D	压缩
最远距离法	$\frac{1}{2}$	$\frac{1}{2}$	0	$\frac{1}{2}$	各种D	扩张
中线法	$\frac{1}{2}$	$\frac{1}{2}$	$-\frac{1}{4} \leq \beta \leq 0$	0	欧氏距离	保持
重心法	$\frac{n_p}{n_p+n_q}$	$\frac{n_q}{n_p+n_q}$	$-\frac{n_p n_q}{(n_p+n_q)^2}$	0	欧氏距离	保持
组平均法	$\frac{n_p}{n_p+n_q}$	$\frac{n_q}{n_p+n_q}$	0	0	各种D	保持
距离平方和法	$\frac{n_k+n_p}{n_k+n_r}$	$\frac{n_k+n_q}{n_k+n_r}$	$-\frac{n_k}{n_k+n_r}$	0	欧氏距离	压缩
可变数平均法	$(1-\beta) \frac{n_p}{n_r}$	$(1-\beta) \frac{n_q}{n_r}$	<1	0	各种D	不定
可变法	$\frac{(1-\beta)}{2}$	$\frac{(1-\beta)}{2}$	<1	0	各种D	扩张

子表示：

$$d_{kr}^2 = a_p d_{pk}^2 + a_q d_{qk}^2 + \gamma |d_{pk}^2 - d_{qk}^2| \quad (15)$$

当 $\gamma = -\frac{1}{2}$ 时，(15) 式就是最短距离聚类法计算类之间的距离的公式

(13) 式；当 $\gamma = \frac{1}{2}$ 时，(15) 式就是最远距离聚类法计算类之间的距离的

公式 (14) 式。

除了最短距离聚类法和最远距离聚类法外，系统聚类的方法还有多种，公式：

$$D_{kr}^2 = a_p D_{kp}^2 + a_q D_{kq}^2 + \gamma |D_{kp}^2 - D_{kq}^2| \quad (16)$$

就是八种不同系统聚类方法计算类之间距离的统一表达式。当 α 、 β 、 γ 三个参数取不同的值时，就形成了不同的聚类方法（见表 2-13），式中 n_p 是 p 类中单元的个数， n_q 是 q 类中单元的个数， $n_r = n_p + n_q$ ； β 一般取负值。

第四节 主成分分析方法

地理环境是多要素的复杂系统，在我们进行地理系统分析时，多变量问题是经常会遇到的。变量太多，无疑会增加分析问题的难度与复杂性，而且在许多实际问题中，多个变量之间是具有一定的相关关系的。因此，我们就会很自然地想到，能否在各个变量之间相关关系研究的基础上，用较少的新变量代替原来较多的变量，而且使这些较少的新变量尽可能多地保留原来较多的变量所反映的信息？事实上，这种想法是可以实现的，本节拟介绍的主成分分析方法就是综合处理这种问题的一种强有力的方法。

一、主成分分析的基本原理

主成分分析是把原来多个变量化为少数几个综合指标的一种统计分析方法，从数学角度来看，这是一种降维处理技术。假定有 n 个地理样本，每个样本共有 p 个变量描述，这样就构成了一个 $n \times p$ 阶的地理数据矩阵：

$$X = \begin{pmatrix} X_{11} & X_{12} & \cdots & X_{1p} \\ X_{21} & X_{22} & \cdots & X_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ M & M & M & M \\ \vdots & \vdots & \ddots & \vdots \\ X_{n1} & X_{n2} & \cdots & X_{np} \end{pmatrix} \quad (1)$$

如何从这么多变量的数据中抓住地理事物的内在规律性呢？要解决这个问题，自然要在 p 维空间中以加以考察，这是比较麻烦的。为了克服这一困难，就需要进行降维处理，即用较少的几个综合指标来代替原来较多的变量指标，而且使这些较少的综合指标既能尽量多地反映原来较多指标所反映的信息，同时它们之间又是彼此独立的。那么，这些综合指标（即新变量）应如何选取呢？显然，其最简单的形式就是取原来变量指标的线性组合，适当调整组合系数，使新的变量指标之间相互独立且代表性最好。

如果记原来的变量指标为 $x_1, x_2, \dots, x_p$, 它们的综合指标——新变量指标为 $x_1, x_2, \dots, x_m$ ($m \leq p$)。则

$$\begin{cases} x_1 = l_{11}x_1 + l_{12}x_2 + \dots, + l_{1p}x_p \\ x_2 = l_{21}x_1 + l_{22}x_2 + \dots, + l_{2p}x_p \\ \\ z_m = l_{m1}x_1 + l_{m2}x_2 + \dots, + l_{mp}x_p \end{cases} \quad (2)$$

在 (2) 式中, 系数 l_{ij} 由下列原则来决定:

(1) z_i 与 z_j ($i \neq j; i, j=1, 2, \dots, m$) 相互无关;

(2) z_1 是 $x_1, x_2, \dots, x_p$ 的一切线性组合中方差最大者； z_2 是与 z_1 不相关的 $x_1, x_2, \dots, x_p$ 的所有线性组合中方差最大者；……； z_m 是与 $z_1, z_2, \dots, z_{m-1}$ 都不相关的 $x_1, x_2, \dots, x_p$ 的所有线性组合中方差最大者。

这样决定的新变量指标 $z_1, z_2, \dots, z_m$ 分别称为原变量指标 $x_1, x_2, \dots, x_p$ 的第一, 第二, $\dots$, 第 m 主成分。其中, z_1 在总方差中占的比例最大, $z_2, z_3, \dots, z_m$ 的方差依次递减。在实际问题的分析中, 常挑选前几个最大的主成分, 这样既减少了变量的数目, 又抓住了主要矛盾, 简化了变量之间的关系。

从以上分析可以看出，找主成分就是确定原来变量 x_j ($j=1, 2, \dots, p$) 在诸主成分 z_i ($i=1, 2, \dots, m$) 上的载荷 l_{ij} ($i=1, 2, \dots, m; j=1, 2, \dots, p$)，从数学上容易知道，它们分别是 $x_1, x_2, \dots, x_p$ 的相关矩阵的 m 个较大的特征值所对应的特征向量。

二、主成分分析的计算步骤

通过上述主成分分析的基本原理的介绍，我们可以把主成分分析计算步骤归纳如下：

(1) 计算相关系数矩阵

$$R = \begin{pmatrix} r_{11} & r_{12} & \cdots & r_{1p} \\ r_{21} & r_{22} & \cdots & r_{2p} \\ M & & M & \\ r_{p1} & r_{p2} & \cdots & r_{pp} \end{pmatrix} \quad (3)$$

在公式 (3) 中， r_{ij} ($i, j=1, 2, \dots, p$) 为原来变量 x_i 与 x_j 的相关系数，其计算公式为

$$r_{ij} = \frac{\sum_{k=1}^n (x_{ki} - \bar{x}_i)(x_{kj} - \bar{x}_j)}{\sqrt{\sum_{k=1}^n (x_{ki} - \bar{x}_i)^2 \sum_{k=1}^n (x_{kj} - \bar{x}_j)^2}} \quad (4)$$

因为 R 是实对称矩阵 (即 $r_{ij}=r_{ji}$)，所以只需计算其上三角元素或下三角元素即可。

(2) 计算特征值与特征向量

首先解特征方程 $|\lambda I - R| = 0$ 求出特征值 λ_i ($i=1, 2, \dots, p$)，并使其按大小顺序排列，即 $\lambda_1 \geq \lambda_2 \geq \dots, \geq \lambda_p \geq 0$ ；然后分别求出对应于特征值 λ_i 的特征向量 e_i ($i=1, 2, \dots, p$)。

(3) 计算主成分贡献率及累计贡献率

主成分 z_i 贡献率： $r_i / \sum_{k=1}^p \gamma_k$ ($i=1, 2, \dots, p$)，累计贡献率： $\sum_{k=1}^m \gamma_k / \sum_{k=1}^p \gamma_k$ 。

一般取累计贡献率达 85-95% 的特征值 $\lambda_1, \lambda_2, \dots, \lambda_m$ 所对应的第一，第二，……，第 m ($m \leq p$) 个主成分。

(4) 计算主成分载荷

$$p(z_k, x_i) = \sqrt{\gamma_k} e_{ki} \quad (i, k=1, 2, \dots, p) \quad (5)$$

由此可以进一步计算主成分得分：

$$Z = \begin{pmatrix} z_{11} & z_{12} & \cdots & z_{1m} \\ z_{21} & z_{22} & \cdots & z_{2m} \\ M & & M & \\ z_{n1} & z_{n2} & \cdots & z_{nm} \end{pmatrix} \quad (6)$$

三、主成分分析实例

对于某区域地貌-水文系统，其 57 个流域盆地的九项地理要素： x_1 为流域盆地总高度 (m)， x_2 为流域盆地山口的海拔高度 (m)， x_3 为流域盆地周长 (m)， x_4 为河道总长度 (km)， x_5 为河

表 2-14 某 57 个流域盆地地理要素数据

序号	x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8
1	760	5490	1.704	2.481	30	2.785	31.8	20
2	1891	4450	2.765	4.394	30	5.833	37.0	26
3	325	5525	1.500	2.660	36	3.042	21.1	25
4	515	4760	2.750	5.320	117	4.844	30.1	98
5	513	6690	1.142	2.080	32	5.100	25.7	26
6	1570	8640	6.130	10.210	76	4.290	24.9	61
7	2210	8415	8.760	15.000	66	4.500	26.6	56
8	515	7040	1.300	1.260	13	3.500	22.2	10
9	1192	6258	8.447	30.606	286	6.500	29.1	225
10	1540	6280	5.174	11.383	82	4.070	23.3	63
11	950	8520	2.880	6.870	62	3.650	27.2	47
12	850	9460	7.480	7.790	30	4.900	11.6	24
13	1237	5937	2.046	2.993	28	2.720	29.6	19
14	553	7480	4.120	22.800	407	4.310	21.0	305
15	281	7050	3.360	8.240	83	4.190	8.20	67
16	1242	6525	3.520	7.490	51	3.790	29.2	41
17	889	7836	3.295	8.655	65	3.740	32.4	50
18	1342	5340	3.120	7.810	69	8.340	33.0	56
19	4523	4879	10.370	78.510	507	4.490	39.3	398
20	3275	6050	5.050	11.530	50	3.570	30.4	38
21	1510	5490	4.090	12.960	116	4.888	30.0	98
22	1655	5245	2.580	4.420	30	2.833	31.9	21
23	1655	5245	2.560	5.460	45	3.42	33.7	34
24	1475	4450	1.837	2.064	18	4.75	37.0	15
25	2144	4197	4.148	9.942	71	4.227	35.0	57
26	515	6650	1.050	1.260	17	5.100	27.4	14
27	834	6450	5.909	16.099	160	6.440	31.1	134
28	834	6450	5.379	10.758	110	4.630	31.1	90
29	1010	6745	4.242	13.694	109	4.430	24.6	86
30	543	6745	1.856	2.898	18	2.420	24.6	13

31	621	7099	2.273	3.863	27	4.600	24.6	21	0.278
32	1290	6745	4.924	12.993	85	4.250	27.8	69	0.947
33	955	7080	2.083	2.387	20	2.78	27.8	16	0.1930
34	885	7150	1.553	1.554	10	2.75	27.8	7	0.1290
35	847	7188	1.591	1.610	14	3.17	31.3	10	0.0940
36	798	7188	1.098	1.023	11	3.00	31.3	8	0.0645
37	1039	5961	2.727	3.295	28	5.50	29.6	24	0.2520
38	1213	5961	3.030	6.894	49	6.43	29.6	41	0.4580
39	1074	5813	2.500	2.954	30	5.33	29.6	26	0.3200
40	370	8295	1.740	2.000	21	4.33	17.8	17	0.1560
41	430	8240	2.130	2.310	14	3.75	18.9	11	0.1820
42	690	8410	1.630	1.680	12	3.25	18.9	9	0.1080
43	773	8410	2.070	2.410	18	3.83	18.9	17	0.1980
44	100	6790	0.830	1.400	25	4.40	11.4	19	0.0429
45	80	6790	0.550	0.470	10	2.75	11.4	7	0.0130
46	96	6765	0.650	0.730	15	4.00	11.4	12	0.0215
47	2490	6535	11.970	59.450	363	2.87	28.0	293	4.9300
48	1765	6575	7.350	21.760	140	3.46	26.7	114	1.9400
49	1158	6862	2.689	4.717	34	3.23	32.8	26	0.3580
50	1070	7055	2.178	3.448	26	2.70	32.8	18	0.2730
51	1495	7055	2.917	3.939	27	2.67	32.8	18	0.2995
52	1601	6949	2.803	4.205	28	3.08	32.8	21	0.3200
53	1251	5135	7.760	23.150	160	3.86	29.5	131	1.1920
54	1587	5095	6.160	17.020	119	4.71	29.9	98	1.3900
55	1230	5120	4.740	8.460	54	3.79	23.4	43	0.8110
56	1290	4960	2.040	2.800	24	6.25	37.0	21	0.1910
57	2400	4920	2.260	3.290	27	5.16	36.2	23	0.2580

道总数， x_6 为平均分叉率， x_7 为河谷最大坡度（度）， x_8 为河源数及 x_9 为流域盆地面积（ km^2 ）的原始数据如表 2-14 所示。张超先生（1984）曾用这些地理要素的原始数据对该区域地貌-水文系统作了主成分分析。下面，我们将其作为主成分分析方法在地理学研究中的一个应用实例介绍给读者，以供参考。

表 2-15 相关系数矩阵

	x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8
x_1	1.000							
x_2	-0.370	1.000						
x_3	0.619	-0.017	1.000					
x_4	0.657	-0.157	0.841	1.000				
x_5	0.474	-0.150	0.737	0.921	1.000			
x_6	0.074	-0.274	0.167	0.094	0.165	1.000		
x_7	0.607	-0.566	0.162	0.217	0.158	0.170	1.000	
x_8	0.481	-0.158	0.753	0.928	0.999	0.181	0.164	1.000
x_9	0.689	0.016	0.910	0.937	0.788	0.071	0.158	0.799

(1)首先将表 2-14 中的原始数据作标准化处理，由公式 (4) 计算得相关系数矩阵（见表 2-15）。

(2)由相关系数矩阵计算特征值，以及各个主成分的贡献率与累计贡献率（见表 2-16）。由表 2-16 可知，第一，第二，第三主成分的累计贡献率已高达 86.5%，故只需求出第一，第二，第三主成分 z_1, z_2, z_3 即可。

表 2-16 特征值及主成分贡献率

主成分	特征值	贡献率 (%)	累计贡献率 (%)
1	5.043	56.029	56.029
2	1.746	19.399	75.428
3	0.997	11.076	86.504
4	0.610	6.781	93.285
5	0.339	3.778	97.061
6	0.172	1.907	98.967
7	0.079	0.8727	99.840
8	0.014	0.1556	99.996
9	0.0004	0.0042	100.00

(3)对于特征值 $\lambda_1=5.043, \lambda_2=1.746, \lambda_3=0.997$ 分别求出其特征向量 e_1, e_2, e_3 ，并计算各变量 $x_1, x_2, \dots, x_9$ 在各主成分上的载荷得到主成分载荷矩阵（见表 2-17）。

表 2-17 主成分载荷矩阵

原变量	主 成 分			占方差的百分数 (%)
	z_1	z_2	z_3	
x_1	0.75	-0.38	-0.36	83.05
x_2	-0.25	0.82	-0.08	73.20
x_3	0.89	0.19	0.00	82.19
x_4	0.97	0.14	-0.03	96.63
x_5	0.91	0.18	0.16	88.26
x_6	0.20	-0.360	0.86	89.97
x_7	0.35	-0.80	-0.25	83.19
x_8	0.92	0.17	0.16	89.90
x_9	0.93	0.22	-0.10	92.16

从表 2-17 可以看出，第一主成分 z_1 与 $x_1, x_3, x_4, x_5, x_8, x_9$ 有较大的正相关，这是由于这六个地理要素与流域盆地的规模有关，因此第一主成分可以被认为是流域盆地规模的代表；第二主成分 z_2 与 x_2 有较大的正相关，与 x_7 有较大的负相关，而这两个地理要素是与流域切割程度有关的，因此第二主成分可以被认为是流域侵蚀状况的代表；第三主成分 z_3 与 x_6 有较大的正相关，而地理要素 x_6 是流域比较独立的特性——河系形态的表征，因此，第三主成分可以被认为是代表河系形态的主成分。

以上分析结果表明，根据主成分载荷，该区域地貌-水文系统的九项地理要素可以被归为三类，即流域盆地的规模，流域侵蚀状况和流域河系形态。如果选取其中相关系数绝对值最大者作为代表，则流域面积，流域盆地出口的海拔高度和分叉率可作为这三类地理要素的代表，利用这三个要素代替原来九个要素进行区域地貌-水文系统分析，可以使问题大大地简化。

第五节 马尔可夫预测方法

对事件的全面预测，不仅要能够指出事件发生的各种可能结果，而且还必须给出每一种结果出现的概率，说明被预测的事件在预测期内出现每一种结果的可能性程度。这就是关于事件发生的概率预测。

马尔可夫 (Markov) 预测法，就是一种关于事件发生的概率预测方法。它是根据事件的目前状况来预测其将来各个时刻 (或时期) 变动状况的一种预测方法。马尔可夫预测法是地理预测研究中重要的预测方法之一。

一、几个基本概念

为了介绍马尔可夫预测法在区域开发研究中的应用，我们首先来介绍有关马尔可夫预测法的几个基本概念。

(一) 状态、状态转移过程与马尔可夫过程

1. 状态 在马尔可夫预测中，“状态”是一个重要的术语。所谓状态，就是指某一事件在某个时刻 (或时期) 出现的某种结果。一般而言，随着所研究的事件及其预测的目标不同，状态可以有不同的划分方式。譬如，在商品销售预测中，有“畅销”、“一般”、“滞销”等状态；在农业收成预测中，有“丰收”、“平收”、“欠收”等状态；在人口构成预测中，有“婴儿”、“儿童”、“少年”、“青年”、“中年”、“老年”等状态；在经济发展水平预测中，有“落后”、“较发达”、“发达”等状态；等等。

2. 状态转移过程 在事件的发展过程中，从一种状态转变为另一种状态，就称为状态转移。譬如，天气变化从“晴天”转变为“阴天”、从“阴天”转变为“晴天”、从“晴天”转变为“晴天”、从“阴天”转变为“阴天”等都是状态转移。

事件的发展，随着时间的变化而变化所作的状态转移，或者说状态转移与时间的关系，就称为状态转移过程，简称过程。

3. 马尔可夫过程 若每次状态的转移都只与前一时刻的状态有关、而与过去的状态无关，或者说状态转移过程是无后效性的，则这样的状态转移过程就称为马尔可夫过程。在区域开发活动中，许多事件发展过程中的状态转移都是具有无后效性的，对于这些事件的发展过程，都可以用马尔可夫过程来描述。

(二) 状态转移概率与状态转移概率矩阵

1. 状态转移概率 在事件的发展变化过程中，从某一种状态出发，下一时刻转移到其它状态的可能性，称为状态转移概率。根据条件概率的定义，由状态 E_i 转为状态 E_j 的状态转移概率 $P(E_i \rightarrow E_j)$ 就是条件概率 $P(E_j/E_i)$ ，即

$$P(E_i \rightarrow E_j) = P(E_j/E_i) = P_{ij} \quad (1)$$

2. 状态转移概率矩阵 假定某一种被预测的事件有 $E_1, E_2, \dots, E_n$ ，共 n 个可能的状态。记 P_{ij} 为从状态 E_i 转为状态 E_j 的状态转移概率，作矩阵

$$P = \begin{pmatrix} P_{11} & P_{12} & \cdots & P_{1n} \\ P_{21} & P_{22} & \cdots & P_{2n} \\ M & M & & M \\ P_{n1} & P_{n2} & \cdots & P_{nn} \end{pmatrix} \quad (2)$$

则称 P 为状态转移概率矩阵。

如果被预测的某一事件目前处于状态 E_i ，那么在下一个时刻，它可能由状态 E_i 转向 $E_1, E_2, \dots, E_i \dots E_n$ 中的任一个状态。所以 P_{ij} 满足条件：

$$\begin{cases} 0 \leq P_{ij} \leq 1 & (i, j = 1, 2, \dots, n) \\ \sum_{j=1}^n P_{ij} = 1 & (i = 1, 2, \dots, n) \end{cases} \quad (3)$$

一般地，我们将满足条件 (3) 的任何矩阵都称为随机矩阵，或概率矩阵。不难证明，如果 P 为概率矩阵，则对任何数 $m > 0$ ，矩阵 P_m 都是概率矩阵。

如果 P 为概率矩阵，而且存在整数 $m > 0$ ，使得概率矩阵 P_m 中诸元素皆非零，则称 P 为标准概率矩阵。可以证明，如果 P 为标准概率矩阵，则存在

非零向量 $\alpha = [x_1, x_2, \dots, x_n]$ ，而且 x_i 满足 $0 \leq x_i \leq 1$ 及 $\sum_{i=1}^n x_i = 1$ ，使得

$$\alpha P = \alpha \quad (4)$$

这样的向量 α 称为平衡向量，或终极向量。

3. 状态转移概率矩阵的计算 计算状态转移概率矩阵 P ，就是要求每个状态转移到其它任何一个状态的转移概率 P_{ij} ($i, j = 1, 2, \dots, n$)。为了求出每一个 P_{ij} ，我们采用频率近似概率的思想来加以计算。

考虑某地区农业收成变化的三个状态，即“丰收”、“平收”和“欠收”。记 E_1 为“丰收”状态， E_2 为“平收”状态， E_3 为“欠收”状态。表 2-18 给出了该地区 1950—1989 年期间农业收成的

表 2-18 某地区农业收成变化的状态转移情况

年份	1950	1951	1952	1953	1954	1955	1956	1957	1958
序号	1	2	3	4	5	6	7	8	9
状态	E_1	E_1	E_2	E_3	E_2	E_1	E_3	E_2	E_1
年份	1960	1961	1962	1963	1964	1965	1966	1967	1968
序号	11	12	13	14	15	16	17	18	19
状态	E_3	E_1	E_2	E_3	E_1	E_2	E_1	E_3	E_3
年份	1970	1971	1972	1973	1974	1975	1976	1977	1978
序号	21	22	23	24	25	26	27	28	29
状态	E_3	E_3	E_2	E_1	E_1	E_3	E_2	E_2	E_1
年份	1980	1981	1982	1983	1984	1985	1986	1987	1988
序号	31	32	33	34	35	36	37	38	39
状态	E_1	E_3	E_2	E_1	E_1	E_2	E_2	E_3	E_1

状态变化情况。以下，我们来计算该地区农业收成变化的状态转移概率矩阵。

从表 2-18 中可知，在 15 个从 E_1 出发 (转移出去) 的状态转移中，有 3 个是从 E_1 转移到 E_1 的 (即 $1 \rightarrow 2, 24 \rightarrow 25, 34 \rightarrow 35$)，有 7 个是从 E_1 转移到 E_2 的 (即 $2 \rightarrow 3, 9 \rightarrow 10, 12 \rightarrow 13, 15 \rightarrow 16, 29 \rightarrow 30, 35 \rightarrow 36, 39 \rightarrow 40$)，有 5 个是从 E_1 转移到 E_3 的 (即 $6 \rightarrow 7, 17 \rightarrow 18, 20 \rightarrow 21, 25 \rightarrow 26, 31 \rightarrow 32$)。

故

$$P_{11} = P(E_1 \rightarrow E_1) = P(E_1 | E_1) = \frac{3}{15} = 0.2000$$

$$P_{12} = P(E_1 \rightarrow E_2) = P(E_2 | E_1) = \frac{7}{15} = 0.4667$$

按照上述同样的办法计算可以得到

$$P_{21} = P(E_2 \rightarrow E_1) = P(E_1 | E_2) = \frac{7}{13} = 0.5385$$

$$P_{22} = P(E_2 \rightarrow E_2) = P(E_2 | E_2) = \frac{2}{13} = 0.1538$$

$$P_{23} = P(E_2 \rightarrow E_3) = P(E_3 | E_2) = \frac{4}{13} = 0.3077$$

$$P_{31} = P(E_3 \rightarrow E_1) = P(E_1 | E_3) = \frac{4}{11} = 0.3636$$

$$P_{32} = P(E_3 \rightarrow E_2) = P(E_2 | E_3) = \frac{5}{11} = 0.4545$$

$$P_{33} = P(E_3 \rightarrow E_3) = P(E_3 | E_3) = \frac{2}{11} = 0.1818$$

所以，该地区农业收成变化的状态转移概率矩阵为

$$P = \begin{bmatrix} 0.200 & 0 & 0.4667 & 0.333 \\ 0.538 & 5 & 0.1538 & 0.3077 \\ 0.363 & 6 & 0.4545 & 0.1818 \end{bmatrix} \quad (5)$$

二、马尔可夫预测法

为了运用马尔可夫预测法对事件发展过程中状态出现的概率进行预测，还需要再介绍一个名词：状态概率 $\pi_j(k)$ 。 $\pi_j(k)$ 表示事件在初始 ($k=0$) 时状态为已知的条件下，经过 k 次状态转移后，第 k 个时刻 (时期) 处于状态 E_j 的概率。根据概率的性质，显然有：

$$\sum_{j=1}^N \pi_j(k) = 1 \quad (6)$$

从初始状态开始，经过 k 次状态转移后到达状态 E_j 这一状态转移过程，可以看作是首先经过 $(k-1)$ 次状态转移后到达状态 E_i ($i=1, 2, \dots, n$)，然后再由 E_i 经过一次状态转移到达状态 E_j 。根据马尔可夫过程的无后效性及 Bayes 条件概率公式，有

$$\pi_j(k) = \sum_{i=1}^n \pi_i(k-1)P_{ij} \quad (j=1, 2, \dots, n) \quad (7)$$

若记行向量 $\pi(k)=[\pi_1(k), \pi_2(k), \dots, \pi_n(k)]$ ，则由(7)式可得逐次计算状态概率的递推公式：

$$\begin{cases} \pi(1) = \pi(0)P \\ \pi(2) = \pi(1)P = \pi(0)P^2 \\ \vdots \\ \pi(k) = \pi(k-1)P = \dots = \pi(0)P^k \end{cases} \quad (8)$$

(8)式中, $\pi(0)=[\pi_1(0), \pi_2(0), \dots, \pi_n(0)]$ 为初始状态概率向量。

(一)第 k 个时刻 (时期)的状态概率预测

由上述分析可知, 如果某一事件在第 0 个时刻 (或时期)的初始状态已知 (即 $\pi(0)$ 已知), 则利用递推公式 (8)式, 就可以求得它经过 k 次状态转移后, 在第 k 个时刻 (时期)处于各种可能的状态的概率 (即 $\pi(k)$), 从而得到该事件在第 k 个时刻 (时期)的状态概率预测。

在前例中, 如果将 1989 年的农业收成状态记为 $\pi(0)=[0, 1, 0]$ (因为 1989 年处于“平收”状态), 则将状态转移概率矩阵 (5)式及 $\pi(0)$ 代入递推公式 (8)式, 就可以求得 1990—2000 年可能出现的各种状态的概率 (见表 2-19)。

表 2-19 某地区 1990—2000 年农业收成状态概率预测值

年份	1990			1991			1992			
状态概率	E_1	E_2	E_3	E_1	E_2	E_3	E_1	E_2	E_3	
	0.5385	0.1528	0.3077	0.3024	0.4148	0.2837	0.3867	0.3334	0.2799	0
年份	1994			1995			1996			
状态概率	E_1	E_2	E_3	E_1	E_2	E_3	E_1	E_2	E_3	
	0.3677	0.3509	0.2799	0.3647	0.3532	0.2799	0.3656	0.3524	0.2799	0
年份	1998			1999			2000			
状态概率	E_1	E_2	E_3	E_1	E_2	E_3	E_1	E_2	E_3	
	0.3653	0.3525	0.2799	0.3653	0.3525	0.2799	0.3653	0.3525	0.2799	

(二)终极状态概率预测

经过无穷多次状态转移后所得到的状态概率称为终极状态概率, 或称平衡状态概率。如果记终极状态概率向量为 $\pi=[\pi_1, \pi_2, \dots, \pi_n]$, 则

$$\pi_i = \lim_{k \rightarrow \infty} \pi_i(k) \quad (i=1, 2, \dots, n) \quad (9)$$

即:

$$\begin{aligned} \pi &= [\lim_{k \rightarrow \infty} \pi_1(k), \lim_{k \rightarrow \infty} \pi_2(k), \dots, \lim_{k \rightarrow \infty} \pi_n(k)] \\ &= \lim_{k \rightarrow \infty} \pi(k) \end{aligned} \quad (10)$$

按照极限的定义可知:

$$\lim_{k \rightarrow \infty} \pi(k) = \lim_{k \rightarrow \infty} \pi(k+1) = \pi \quad (11)$$

将 (11)式代入马尔可夫预测模型的递推公式 (8)式得

$$\lim_{k \rightarrow \infty} \pi(k+1) = \lim_{k \rightarrow \infty} \pi(k) P$$

即:

$$\pi = \pi P \quad (12)$$

这样, 就得到了终极状态概率应满足的条件:

$$(1) \pi = \pi P$$

$$(2) 0 \leq \pi_i \leq 1 \quad (i=1, 2, \dots, n)$$

$$(3) \sum_{i=1}^n \pi_i = 1$$

以上条件 (2) 与 (3) 是状态概率的要求，其中，条件 (2) 表示，在无穷多次状态转移后，事件必处在 n 个状态中的任意一个；条件 (1) 就是用来计算终极状态概率的公式。终极状态概率是用来预测马尔可夫过程在遥远的未来会出现什么趋势的重要信息。

在前例关于某地区农业收成状态概率的预测中，设终极状态的概率为 $\pi = [\pi_1, \pi_2, \pi_3]$ ，则

$$[\pi_1, \pi_2, \pi_3] = [\pi_1, \pi_2, \pi_3] \begin{bmatrix} 0.2000 & 0.4667 & 0.3333 \\ 0.5385 & 0.1538 & 0.3077 \\ 0.3636 & 0.4545 & 0.1818 \end{bmatrix}$$

即

$$\begin{cases} \pi_1 = 0.2000 \pi_1 + 0.5385 \pi_2 + 0.3636 \pi_3 \\ \pi_2 = 0.4667 \pi_1 + 0.1538 \pi_2 + 0.4545 \pi_3 \\ \pi_3 = 0.3333 \pi_1 + 0.3077 \pi_2 + 0.1818 \pi_3 \end{cases} \quad (13)$$

求解方程组 (13) 式得： $\pi_1=0.3653$ ， $\pi_2=0.3525$ ， $\pi_3=0.2799$ 。这说明，该地区农业收成的变化，在无穷多次状态转移后，“丰收”和“平收”状态出现的概率都将大于“欠收”状态出现的概率。

在地理事件的预测中，被预测对象所经历的过程中各个阶段（或时点）的状态和状态之间的转移概率是最为关键的。马尔可夫预测的基本方法就是利用状态之间的转移概率矩阵预测事件发生的状态及其发展变化趋势。马尔可夫预测法的基本要求是状态转移概率矩阵必须具有一定的稳定性。因此，必须具有足够多的统计数据，才能保证预测的精度与准确性。换句话说，马尔可夫预测模型必须建立在大量的统计数据的基础之上。这一点也是运用马尔可夫预测方法预测地理事件的一个最为基本的条件。

第三章 线性规划方法

线性规划，是数学规划中发展较快、应用较广和比较成熟的一个分支。早在本世纪 30 年代末，就有人从运输问题开始研究应用线性规划的方法。自 1947 年丹泽 (G.B.Dantzing) 提出求解线性规划问题的一般方法——单纯形法之后，线性规划在理论上趋于成熟，在实际应用中日益广泛与深入。特别是随着电子计算机的发展和计算速度的不断提高，线性规划适用的领域更加广泛，从工程技术的优化设计到工业、农业、商业、交通运输规划及管理诸问题的研究中，它已成为必不可少的重要手段之一。本章，我们将结合有关实例，介绍和探讨线性规划在地理学研究中的应用问题。

第一节 线性规划及其单纯形求解方法

一、线性规划数学模型

(一) 线性规划之实例

线性规划研究的问题主要有两类：一是某项任务确定后，如何统筹安排，以最少的人力、物力和财力去完成该项任务；二是面对一定数量的人力、物力和财力资源，如何安排使用，使得完成的任务最多。实际上，这是一个问题的两个方面，它们都属于最优规划的范畴。以下，我们列举线性规划问题之若干实例，供读者研究。

1. 运输问题 假设某种物资（譬如煤炭、钢铁、石油等）有 m 个产地， n 个销地。第 i 产地的产量为 a_i ($i=1, 2, \dots, m$)，第 j 销地的需求量为 b_j ($j=1, 2, \dots, n$)，它们满足产销平衡条件 $\sum_{i=1}^m a_i = \sum_{j=1}^n b_j$ 。如果产地 i 到销地 j 的单位物资的运费为 c_{ij} ，试问如何安排该种物资调运计划，才能使总运费达到最小？

设 x_{ij} 表示由产地 i 供给销地 j 的物资数量，则上述问题可以表述为求一组实值变量 x_{ij} ($i=1, 2, \dots, m; j=1, 2, \dots, n$)，使其满足：

$$\begin{cases} \sum_{i=1}^m x_{ij} = b_j & (j=1, 2, \dots, n) \\ \sum_{j=1}^n x_{ij} = a_i & (i=1, 2, \dots, m) \\ x_{ij} \geq 0 & (i=1, 2, \dots, m; j=1, 2, \dots, n) \end{cases}$$

而且使：

$$z = \sum_{i=1}^m \sum_{j=1}^n c_{ij} x_{ij} \rightarrow \min$$

2. 资源利用问题 假设某地区拥有 m 种资源，其中，第 i 种资源在规划期内的限额为 b_i ($i=1, 2, \dots, m$)。这 m 种资源可用来生产 n 种产品，其中，生产单位数量的第 j 种产品需要消耗的第 i 种资源的数量为 a_{ij} ($i=1, 2, \dots,$

$m; j=1, 2, \dots, n)$, 第 j 种产品的单价为 c_j ($j=1, 2, \dots, n$)。试问如何安排这几种产品的生产计划, 才能使规划期内资源利用的总产值达到最大?

设第 j 种产品的生产数量为 x_j ($j=1, 2, \dots, n$), 则上述资源利用问题就是:

在约束条件

$$\begin{cases} \sum_{j=1}^n a_{ij}x_j \leq b_i & (i=1, 2, \dots, m) \\ x_j \geq 0 & (j=1, 2, \dots, n) \end{cases}$$

下, 求一组实数变量 x_j ($j=1, 2, \dots, n$), 使

$$Z = \sum_{j=1}^n c_j x_j \rightarrow \max$$

3. 合理下料问题 用某种原材料切割零件 $A_1, A_2, \dots, A_m$ 的毛坯, 现已设计出在一块原材料上有 $B_1, B_2, \dots, B_n$ 种不同的下料方式, 如用 B_j 下料方式可得 A_i 种零件 a_{ij} 个, 设 A_i 种零件的需要量为 b_i 个。试问应该怎样组织下料活动, 才能使得既满足需要, 又使用去的原材料最少?

设采用 B_j 方式下料的原材料数为 x_j , 则上述问题可表示为:

在约束条件

$$\begin{cases} \sum_{j=1}^n a_{ij}x_j \geq b_i & (i=1, 2, \dots, m) \\ x_j \geq 0 & (j=1, 2, \dots, n) \end{cases}$$

下, 求一组整数变量 x_j ($j=1, 2, \dots, n$), 使得

$$Z = \sum_{j=1}^n x_j \rightarrow \min$$

(二) 线性规划数学模型

线性规划应用之实例还有很多, 譬如生产布局问题, 连续投资问题, 等等, 不一一列举。从以上的例子可以看出, 尽管它们表示的形式不尽相同, 但它们都具有以下共同的特征: (1) 每一个问题都用一组未知变量 ($x_1, x_2, \dots, x_n$) 表示某一规划方案, 这组未知变量的一组定值代表一个具体的方案, 而且通常要求这些未知变量的取值是非负的。

(2) 每一个问题都有两个主要组成部分: 一是目标函数, 按照研究问题的不同, 常常要求目标函数取最大或最小值; 二是约束条件, 它定义了一种求解范围, 使问题的解必须在这一范围之内。

(3) 每一个问题的目标函数和约束条件都是线性的。

根据上述问题的三个基本特征, 我们可以抽象出线性规划问题的数学模型。它一般地可表示为:

在线性约束条件

$$\sum_{j=1}^n a_{ij}x_j \leq (\geq, =) b_i \quad (i=1, 2, \dots, m) \quad (1)$$

以及非负约束条件

$$x_j \geq 0 (j=1, 2, \dots, n) \quad (2)$$

下，求一组未知变量 x_j ($j=1, 2, \dots, n$) 的值，使

$$Z = \sum_{j=1}^n c_j x_j \rightarrow \min(\max) \quad (3)$$

若采用矩阵记号，上述线性规划模型的一般形式可进一步描述为：在约束条件

$$AX \leq (\geq, =) b \quad (1')$$

$$\text{以及} \quad x \geq 0 \quad (2')$$

下，求未知向量 $x=[x_1, x_2, \dots, x_n]^T$ ，使得

$$Z=CX \rightarrow \min(\max) \quad (3')$$

其中

$$b=[b_1, b_2, \dots, b_m]^T$$

$$c=[c_1, c_2, \dots, c_n]$$

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

二、线性规划的标准形式

(一)线性规划的标准形式

为了讨论与计算上的方便，常常需要将线性规划问题的数学模型转化为标准形式，即在约束条件：

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + \cdots + a_{1n}x_n = b_1 \\ a_{21}x_1 + a_{22}x_2 + \cdots + a_{2n}x_n = b_2 \\ \dots\dots\dots \\ a_{m1}x_1 + a_{m2}x_2 + \cdots + a_{mn}x_n = b_m \end{cases} \quad (4)$$

以及

$$x_j \geq 0 (j=1, 2, \dots, n) \quad (5)$$

下，求一组未知变量 x_j ($j=1, 2, \dots, n$) 的值，使

$$Z = \sum_{j=1}^N c_j x_j \quad (6)$$

达到最小值。

其缩写形式为：在约束条件

$$\sum_{j=1}^n a_{ij} x_j = b_i (i=1, 2, \dots, m) \quad (4')$$

以及

$$x_j \geq 0 (j=1, 2, \dots, n) \quad (5')$$

下，求一组未知变量 x_j ($j=1, 2, \dots, n$) 的值，使得：

$$Z = \sum_{j=1}^N c_j x_j \rightarrow \min \quad (6')$$

采用矩阵记号，上述标准形式还可以进一步简记为：在约束条件

$$AX=b \quad (4'')$$

以及

$$X \geq 0 \quad (5'')$$

下，求未知向量 X ，使得：

$$Z=CX \rightarrow \min \quad (6'')$$

在通常情况下， b 和 c 为已知常数向量； A 为已知常数矩阵，且 A 的秩为 m 。

上述标准形式的线性规划，常被记为如下更为紧凑的形式：

$$\begin{cases} \min z = \sum_{j=1}^n c_j x_j \\ \sum_{j=1}^n a_{ij} x_j = b_i \quad (i=1, 2, \dots, m) \\ x_j \geq 0 \quad (j=1, 2, \dots, n) \end{cases}$$

或

$$\begin{cases} \min Z = CX \\ AX = b \\ X \geq 0 \end{cases}$$

(二)化为标准形式的方法

具体的线性规划问题，其数学模型常常是各式各样的，它们不一定符合线性规划的标准形式的要求，为了将其转化为标准形式，常常需要对目标函数或约束条件采用一定的变换方法进行转换。

1. 目标函数化为标准形式的方法 如果其线性规划问题的目标函数为：

$$\max Z = CX$$

则令 $Z' = -Z$ ，显然：

$$-\max Z = \min (-Z) = \min Z'$$

所以，可将原问题的目标函数换为：

$$\min Z' = -CX$$

这就是标准形式的目标函数了。

2. 约束方程化为标准形式的方法 如果第 k 个约束方程为不等式，即

$$a_{k1}x_1 + a_{k2}x_2 + \dots + a_{kn}x_n \leq (\geq) b_k$$

则只需在原问题中引入松弛变量 $x_{n+k} \geq 0$ ，并将第 k 个方程改写为：

$$a_{k1}x_1 + a_{k2}x_2 + \dots + a_{kn}x_n + (-)x_{n+k} = b_k$$

而将其目标函数看作

$$z = \sum_{j=1}^n c_j x_j = \sum_{j=1}^n x_j c_j + 0 \cdot x_{n+k}$$

这样就将原问题转化为标准形式的线性规划模型了。

三、线性规划的解及其性质

(一)线性规划的解

1. 可行解与最优解 在线性规划问题中，称满足约束条件（即满足线性约

束和非负约束)的一组变量值 $x=[x_1, x_2, \dots, x_n]^T$ 为可行解。所有可行解组成的集合称为可行域。使目标函数最大或最小化的可行解称为最优解。

2. 基本解与基本可行解 在线性规划问题 (4'')—(6'') 式中, 如果我们把约束方程组的 $m \times n$ 阶系数矩阵 A 写成由 n 个列向量组成的分块矩阵

$$A=[p_1, p_2, \dots, p_n] \quad (7)$$

在 (7) 式中, $p_j=[a_{1j}, a_{2j}, \dots, a_{mj}]^T$ ($j=1, 2, \dots, n$)。则 p_j 是对应变量 x_j 的系数列向量。

如果 B 是 A 中的一个 $m \times m$ 阶的非奇异子阵, 则称 B 为该线性规划问题的一个基。不失一般性, 不妨设:

$$B = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1m} \\ a_{21} & a_{22} & \cdots & a_{2m} \\ M & M & & M \\ a_{m1} & a_{m2} & \cdots & a_{mm} \end{pmatrix} = [p_1, p_2, \dots, p_m] \quad (8)$$

则称 p_j ($j=1, 2, \dots, m$) 为基向量, 与基向量 p_j 相对应的变量 x_j ($j=1, 2, \dots, m$) 为基变量, 而其余的变量 x_j ($j=m+1, m+2, \dots, n$) 为非基变量。

如果 $X_B=[x_1, x_2, \dots, x_m]^T$ 是方程组

$$BX_B=b \quad (9)$$

的解, 则 $X=[x_1, x_2, \dots, x_m, 0, 0, \dots, 0]^T$ 为方程组 (4'') 式的一个解, 它称之为对应于基 B 的基本解。

满足非负约束条件的基本解, 称为基本可行解。对应于基本可行解的基称为可行基。

线性规划问题的以上几个解的关系, 可用图 3-1 来描述。

(二) 线性规划解的性质

1. 凸集和顶点 为了说明线性规划解的性质, 需要引入凸集和顶点的概念。

若连结 n 维点集 S 中的任意两点 $x^{(1)}$ 和 $x^{(2)}$ 之间的线段仍在 S 中, 则称 S 为凸集。例如, 三角形、平行四边形、正多边形、圆、球体、正多面体等都是凸集。而圆环、空心球等都不是凸集。

若凸集 S 中的点 $X^{(0)}$ 不能成为 S 中任何线段的内点, 则称 $X^{(0)}$ 为 S 的顶点或极点。例如, 三角形、平行四边形、正多边形、正多面体的顶点以及圆周上的点等都是极点。

2. 线性规划解的性质 可以证明, 线性规划问题的解具有以下性质:

(1) 线性规划问题的可行解集 (可行域) 为凸集。

(2) 可行解集 S 中的点 X 是顶点的充要条件为 X 是基本可行解。

(3) 若可行解集有界, 则线性规划问题的最优值一定可以在其顶点上达到。由于系数矩阵 A 中的基是有限的, 因此基本可行解也是有限的, 这就是说可行解集的顶点数目是有限的。所以, 如果线性规划问题有最优解, 就只需从其可行解集的有限个顶点中去寻找。

四、线性规划问题的求解方法——单纯形法

(一)单纯形表

在标准形式线性规划 (4'') - (6'') 句式中, 不妨设 $B=[p_1, p_2, \dots, p_m]$ 是一个基, 令 $N=[p_{m+1}, p_{m+2}, \dots, p_n]$, 则 $A=[B, N]$ 。记基变量为 $x_B=[x_1, x_2, \dots, x_m]^T$, 非基变量为 $x_N=[x_{m+1}, x_{m+2}, \dots, x_n]^T$, 则运用分块矩阵的运算法则可知, (4'') 式可以被进一步改写为

$$BX_B + NX_N = b \quad (10)$$

用 B^{-1} 左乘 (10) 式两端, 并作适当整理, 可得

$$X_B = B^{-1}b - B^{-1}NX_N \quad (11)$$

(11) 式就是用非基变量表示基变量的关系式。

相应地, 记 $C_B=[c_1, c_2, \dots, c_m]$, $C_N=[c_{m+1}, c_{m+2}, \dots, c_n]$, $C=[C_B, C_N]$, 则目标函数亦可以改写为

$$Z = C_B B^{-1}b + (C_N - C_B B^{-1}N)X_N \quad (12)$$

显然, $X_B = B^{-1}b$, $X_N = 0$ 构成了对应于基 B 的基本解, 其相应的目标函数值为 $Z = C_B B^{-1}b$ 。

如果 $B^{-1}b \geq 0$, 则 $X_B = B^{-1}b$, $X_N = 0$ 构成了一个基本可行解, B 是一个可行基。

如果 $C_B B^{-1}N - C_N \leq 0$, 则由 (12) 式可以看出, 对于一切可行解 X , 有 $Z = CX \geq C_B B^{-1}b$ 。这就是说, 对应于 B 的基本可行解为最优解, 这时, B 也被称为最优基。

由于

$$\begin{aligned} C_B B^{-1}A - C &= C_B B^{-1}[B, N] - [C_B, C_N] \\ &= [C_B, C_B B^{-1}N] - [C_B, C_N] \\ &= [0, (C_B B^{-1}N - C_N)] \end{aligned} \quad (13)$$

即 $C_B B^{-1}N - C_N \leq 0$ 与 $C_B B^{-1}A - C \leq 0$ 等价, 故可得以下最优性判定定理。

定理: 对于基 B , 若 $B^{-1}b \geq 0$, 且 $C_B B^{-1}A - C \leq 0$, 则对应于基 B 的基本解为最优解, B 为最优基。

由 (12) 和 (13) 式可得:

$$Z + (C_B B^{-1}A - C)X = C_B B^{-1}b \quad (14)$$

用 B^{-1} 左乘 (4'') 式两端得:

$$B^{-1}AX = B^{-1}b \quad (15)$$

联合 (14) 与 (15) 式可得:

$$\begin{bmatrix} 1 & C_B B^{-1}A - C \\ 0 & B^{-1}A \end{bmatrix} \begin{bmatrix} Z \\ X \end{bmatrix} = \begin{bmatrix} C_B B^{-1}b \\ B^{-1}b \end{bmatrix} \quad (16)$$

在 (16) 式中, 称系数矩阵

$$\begin{bmatrix} C_B B^{-1}b & 1 & C_B B^{-1}A - C \\ B^{-1}b & 0 & B^{-1}A \end{bmatrix}$$

或

$$\begin{bmatrix} C_B^{-1}b & C_B B^{-1}A - C \\ B^{-1}b & B^{-1}A \end{bmatrix}$$

为对应于基 B 的单纯形表，记作 $T(B)$ 。

如果记 $C_B B^{-1}b = b_{00}$, $C_B B^{-1}A - C = [b_{01}, b_{02}, \dots, b_{0n}]$, $B^{-1}b = [b_{10}, b_{20}, \dots, b_{m0}]^T$, 以及

$$B^{-1}A = \begin{bmatrix} b_{11} & b_{12} & \cdots & b_{1n} \\ b_{21} & b_{22} & \cdots & b_{2n} \\ \cdots & \cdots & \cdots & \cdots \\ b_{m1} & b_{m2} & \cdots & b_{mn} \end{bmatrix}$$

则：

$$T(B) = \begin{bmatrix} b_{00} & b_{01} & b_{02} & \cdots & b_{0n} \\ b_{10} & b_{11} & b_{12} & \cdots & b_{1n} \\ b_{20} & b_{21} & b_{22} & \cdots & b_{2n} \\ M & M & M & & M \\ b_{m0} & b_{m1} & b_{m2} & \cdots & b_{mn} \end{bmatrix}$$

(17)式表明，单纯形表就是把非基变量作为参数，表示基变量和目标函数时的系数矩阵。 b_{00} 就是对应于基 B 的基本解下的目标函数值； $b_{10}, b_{20}, \dots, b_{m0}$ 就是对应于基 B 的基本解的基变量值； $b_{01}, b_{02}, \dots, b_{0n}$ 为检验系数，如果这组数均非正，则这一基本可行解就是最优解。

(二)单纯形法的计算步骤

单纯形法求解线性规划问题的计算步骤如下：

第一步，找出初始可行基 $B = [p_{j1}, p_{j2}, \dots, p_{jm}]$ ，建立初始单纯形表。

第二步，判别、检查所有的检验系数 b_{0j} ($j=1, 2, \dots, n$)。

(1)如果所有的 $b_{0j} \leq 0$ ($j=1, 2, \dots, n$)，则由最优性判定定理知，已获最优解。

(2)若检验系数 b_{0j} ($j=1, 2, \dots, m$)中，有些为正数，但其中某一正的检验系数所对应的列向量的各分量均非正，则线性规划问题无解。

(3)若检验系数 b_{0j} ($j=1, 2, \dots, n$)中，有些为正数，且它们所对应的列向量中有正的分量，则进行换基迭代。

第三步，选主元。在所有 $b_{0j} > 0$ 的检验系数中选取最大的一个 b_{0s} ，其对应的非基变量为 x_s ，对应的列向量为 $p_s = [b_{1s}, b_{2s}, \dots, b_{ms}]^T$ 。如果

$$\theta = \min \left\{ \frac{b_{i0}}{b_{is}} \mid b_{is} > 0 \right\} = \frac{b_{r0}}{b_{rs}}$$

则确定 b_{rs} 为主元项。

第四步，在基 B 中调进 p_s ，换出 p_{jr} ，得到一个新的基 $B' = [p_{j1}, p_{j2}, \dots, p_{jr-1}, p_s, p_{jr+1}, \dots, p_{jm}]$ 。

第五步，在单纯形表上进行初等行变换，使第 s 列向量变为单位向量，又得一张新的单纯形表。

第六步，转入上述第二步。

例 1：用单纯形方法求解线性规划问题：

$$\begin{cases} x_1 + 3x_2 \leq 12 \\ 2x_1 + x_2 \leq 9 \\ x_1 \geq 0, x_2 \geq 0 \end{cases}$$

$$\max Z = 2x_1 + 3x_2$$

解：首先引入松弛变量 x_3, x_4 ，把原问题化为标准形式：

$$\begin{cases} x_1 + 3x_2 + x_3 = 12 \\ 2x_1 + x_2 + x_4 = 9 \\ x_1, x_2, x_3, x_4 \geq 0 \end{cases}$$

$$\min Z' = -2x_1 - 3x_2$$

在上述问题中，

$$A = \begin{bmatrix} 1 & 3 & 1 & 0 \\ 2 & 1 & 0 & 1 \end{bmatrix}, p_1 = \begin{bmatrix} 1 \\ 2 \end{bmatrix}, p_2 = \begin{bmatrix} 3 \\ 1 \end{bmatrix}, p_3 = \begin{bmatrix} 1 \\ 0 \end{bmatrix},$$

$$p_4 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}, b = \begin{bmatrix} 12 \\ 9 \end{bmatrix}, C = [-2, -3, 0, 0]$$

第一步，因为 $B_1 = [p_3, p_4]$ 为单位矩阵，且 $B_1^{-1}b = b > 0$ ，故 B_1 是一个可行基。由于 $C_B B_1^{-1}b = 0$ ， $C_B B_1^{-1}A - C = -C$ ， $B_1^{-1}A = A$ ，所以对应于 B_1 的初始单纯形表（表 3-1）如下：

表 3-1

		x_1	x_2	x_3	x_4
z'	0	2	3	0	0
x_2	12	1	3	1	0
x_4	9	2	1	0	1

第二步，判别。在初始单纯形表中， $b_{01}=2, b_{02}=3$ ，所以 B_1 不是最优基，要进行换基迭代。

第三步，选主元。由于 $\max\{2, 3\}=3$ ，所以取 $s=2$ ，其对应的非基变量为 x_2 ，对应的列向量为 $p_2 = \begin{pmatrix} 3 \\ 1 \end{pmatrix}$ 。 $\theta = \min\{\frac{12}{3}, \frac{9}{1}\} = 3$ ，所以 $r = 1$ 。因而主元项为 $b_{12}=3$ 。

第四步， p_2 调入基， p_3 退出基，得新一的基 $B_2 = (p_2, p_4) = \begin{pmatrix} 3 & 0 \\ 1 & 1 \end{pmatrix}$ 。

第五步，对初始单纯形表进行初等行变换，使 p_2 变为单位向量，可得 B_2 下的新单纯形表（表 3-2）。

表 3-2

		x_1	x_2	x_3	x_4
x'	-12	1	0	-1	0
x_2	4	$\frac{1}{3}$	1	$\frac{1}{3}$	0
x_4	5	$\frac{5}{3}$	0	$-\frac{1}{3}$	0

第六步，转入第二步。因为在对应于 B_2 的单纯形表中，检验系数有正数

$b_{01} = 1$ ，重复以上步骤可得对应于 $B_3 = (p_2, p_1) = \begin{pmatrix} 3 & 1 \\ 1 & 2 \end{pmatrix}$ 的单纯形表。

的单纯形表 (表 3-3)。因为在对应于 B_3 的单纯形表中，检验系数已经没有正数，所以 B_3 是最优基，其对应的基本最优解为： $x_1=3$ ， $x_2=3$ ， $x_3=0$ ， $x_4=0$ ，目标函数最小值为 $Z' = -15$ 。而对于原问题，其最优解为 $x_1=3$ ， $x_2=3$ ，目标函数最大值为 $Z=15$ 。

表 3-3

		x_1	x_2	x_3	x_4
x'	-15	0	0	$-\frac{4}{5}$	$-\frac{3}{5}$
x_2	3	0	1	$\frac{2}{5}$	$-\frac{1}{5}$
x_1	3	1	0	$-\frac{1}{5}$	$\frac{3}{5}$

五、线性规划应用实例

(一)问题的提出

1983 年底，甘肃省委从实际出发，在大量的调查和可行性分析的基础上，对中部干旱地区提出了“三年停止植被破坏，五年解决温饱问题”的近期目标，鼓励农民种草养畜，走农牧草相结合的道路，并先后在通渭县的申家山、会宁县的大岔社等地，由省科委和草原工作队技术指导，对种草养畜工作进行试点。1986 年初，中部干旱地区的农业生产结构发生了重大变化，草在种植业中占了较大的比重，养畜工作也初见成效，同时大部分地区的粮食生产已基本上得到了自给，农业生态环境逐步改善，农民依靠出售草籽而得到了较多的经济收入。但是，自 1986 年以来，草籽价格暴跌，而且国家不再统一收购草籽，这直接影响了农民的经济收入，有些人不愿再用耕地种草，在这种情况下，中部干旱地区的农业生产结构如何调整，怎样实现草畜转化？这是人们所关心的问题。为了解决这一问题，笔者曾以会宁县大岔社为例，运用线性规划方法建立了农业生产结构调整的规划模型，并由此探讨了甘肃中部干旱地区的农业生产结构问题。该规划模型的基期为 1986 年，报告期为 1990 年。实践证明，该模型在指导本地农业生产活动方面起到了一定的作用。以下，我们将这一研究成果作一些简单地介绍。

(二)模型的建立

1. 建模原则 在模型的建立过程中，我们提出了如下的建模原则。

(1)实现“草畜转化”，使农业生态环境继续得到改善，使土地资源的利用继续趋于合理化。

(2)经济效益与生态效益相结合，用经济效益刺激生态效益的提高，以生态效益保证经济效益的实现。

(3)着重考虑“三料”问题，同时也考虑农民的吃粮与花钱问题。

(4)模型不但要有理论意义，而且要有实用价值，因此，模型变量的选取要恰当，约束条件的提出要合理，模型参数的确定要准确，切合实际。

(5)考虑到剩余劳动力的存在，把“劳务输出”作为农业生产结构调整的一个重大决策问题考虑。

2. 模型结构

(1)模型变量：模型中的变量，分为外生变量与内生变量两类，外生变量是指由方针政策等因素决定的变量，内生变量是由模型本身所决定的变量。本模型的内生变量如下：

x_1 ：粮食作物面积（亩）； x_2 ：经济作物面积（亩）； x_3 ：耕地种草面积（亩）； x_4 ：荒山造林或种草面积； x_5 ：大畜存栏数（头）； x_6 ：生猪年末存栏数（头）； x_7 ：兔年末存栏数（只）； x_8 ：劳务输出的劳动力个数（个）。

(2)目标函数：使农、林、牧各业的年平均产值达到最大（以1986年现行价格计算）。

(3)约束条件：①粮食的生产与分配；粮食总产量-籽种-公购粮-牲畜饲料粮 $\geq$ 农业人口的总口粮。

②燃料。作物秸杆 $\times 30\%$ +荒山产柴量 $\times 80\%$ +衣壳 $\times 50\% \geq$ 全年燃料总量；

③饲料。作物秸杆 $\times 70\%$ +装子 $\times 50\%$ +耕地种草产量+可利用的野生草量 $\geq$ 各类畜的食草总量。（作物秸杆的30%作燃料；70%作饲草；装子的50%作燃料，50%作饲草；对于中部地区的林业应采取保护性政策，因此，用柴量可按产柴量的80%计）。

④肥料：人类 + 猪粪 + 畜粪 $\times \frac{2}{3}$ + 鸡羊粪 $\geq$ 粮食作物与经济作物所需

有机肥总量。（畜粪有 $\frac{2}{3}$ 用于土地，其余 $\frac{1}{3}$ 被用于烧炕）。

⑤耕地总面积限制。

⑥经济作物面积限制。

⑦耕地草面积限制。

⑧荒山面积限制。

⑨耕畜限制。

⑩水约束。

(11)劳动力限制。

(12)经济收入。

3. 数学模型 以上模型结构数量化，可得如下的数学模型：

求： $\max Z = 33x_1 + 50x_2 + 135x_5 + 40x_2 + 100x_6 + 5x_7 + 800x_8 + 9334$ 满足：

$$\left\{ \begin{array}{llll} 160x_1 & -120x_5 & -100x_6 & \geq 101070 \\ 99x_1 & -160x_4 & & \geq 196020 \\ 171x_2 + 500x_2 & -4380x_5 - 1095x_5 - 120x_7 & & \geq 0 \\ 250x_1 + 255x_2 & -2433x_5 - 730x_6 & & \leq 43150 \\ x_1 + x_2 + x_3 & & & \leq 1625 \\ x_2 & & & \geq 77.5 \\ x_3 & & & \geq 650 \\ x_4 & & & \leq 1020 \\ x_5 & & & \geq 54 \\ 3.95x_6 & & & \leq 340 \\ 3x_1 + 3x_2 + 3x_3 + x_4 & & & \leq 6750 \\ 25x_2 + 20x_3 + 80x_6 + 4.5x_7 + 80x_8 & & & > 38568 \\ x_i \geq 0 \ (i = 1, 2, \dots, 8) \end{array} \right.$$

上述模型求解得： $\max Z = 11519190$ （元）

$x_1 = 730.41$ （亩）， $x_2 = 77.50$ （亩）； $x_3 = 817.10$ （亩）； $x_4 = 773.19$ （亩）， $x_5 = 54.00$ （头）； $x_6 = 93.15$ （头）； $x_7 = 1264.47$ （只）， $x_8 = 25.59$ （个）。

（三）甘肃中部干旱地区农业生产结构的特征

通过对会宁县大盆社农业生产结构优化结果的分析，笔者认为甘肃中部干旱地区农业生产结构应具有下述特征。

（1）自给性的农业是产业构成的基础，粮食生产的自给是其它一切生产的必要前提。

（2）商品性的畜牧业是产业构成的主体，草畜转化是生态效益变成经济效益的有效途径。

（3）纽带性的草业是农牧相互联系、相互促进的中枢性产业。

（4）保护性林业是改善生态环境的重要措施，也是农村燃料的主要来源。

（5）剩余劳动力的存在，为农副产品的初级加工及农村第二、第三产业的发展提供了可能的条件。

第二节 线性规划的对偶理论

一、对偶问题的提出

对于本章第一节中所介绍的资源利用问题

$$\begin{cases} \max Z = \sum_{j=1}^n c_j x_j & (1) \\ \sum_{j=1}^n a_{ij} x_j \leq b_i \quad (i=1, 2, \dots, m) & (2) \\ x_j \geq 0 \quad (j=1, 2, \dots, n) & (3) \end{cases}$$

我们也可以从另一个角度去作这样的考虑，即资源拥有者为了实现一定的收入目标，他可以不用这 m 种资源去生产 n 种产品，而是将其所拥有的资源出售。这时，他就必须考虑给每一种资源如何定价的问题。

如果用 y_i ($i=1, 2, \dots, m$) 表示出售单位数量的第 i 种资源的价格。显然，资源拥有者在作出出售资源的决策时，必须考虑到：出售资源的收入不应该低于生产所获得的收入，这样就有：

$$\sum_{i=1}^m a_{ij} y_i \geq c_j \quad (j=1, 2, \dots, n)$$

显然 $y_i \geq 0$ ($i=1, 2, \dots, m$)

如果资源拥有者将所有资源出售，则他得到的总收入为

$$W = \sum_{i=1}^m b_i y_i$$

对于资源拥有者来说，他当然希望 W 愈大愈好；但对于资源接受者来说，他当然希望 W 愈小愈好。所以，资源拥有者出售每一种资源的最低估价，应该通过求解线性规划问题而得到。

$$\begin{cases} \min W = \sum_{i=1}^m b_i y_i & (4) \\ \sum_{i=1}^m a_{ij} y_i \geq c_j \quad (j=1, 2, \dots, n) & (5) \\ y_i \geq 0 \quad (i=1, 2, \dots, m) & (6) \end{cases}$$

上述讨论表明，对于同样一个资源利用问题，从两个不同的角度去考虑，可以得到两个不同但又相互联系的线性规划模型，这就是线性规划的对偶问题。其中，(1) — (3) 称为原线性规划问题，(4) — (6) 称为原线性规划问题的对偶线性规划问题。

一般地，称线性规划问题

$$(I) \begin{cases} \max Z = CX \\ AX \leq b \\ X \geq 0 \end{cases}$$

与线性规划问题

$$(II) \begin{cases} \min W = Yb \\ YA \geq C \\ Y \leq 0 \end{cases}$$

为相互对偶的线性规划问题。这就是说，如果将 (I) 看作原问题，则 (II) 为其对偶问题；反之，如果将 (II) 看作原问题，则 (I) 为其对偶问题。

(I) 与 (II) 是线性规划原问题与对偶问题的标准形式，它们之间的变换关系是对称形式 (如表 3-4 所示)。因此，根据原问题的系数矩阵就能很容易地写出它的对偶问题。

表 3-4

$y_i \backslash x_j$	$x_1 \quad x_2 \quad \cdots \quad x_n$	原关系	minW
y_1	$a_{11} \quad a_{12} \quad \cdots \quad a_{1n}$	$\leq$	b_1
y_2	$a_{21} \quad a_{22} \quad \cdots \quad a_{2n}$	$\leq$	b_2
M	M M M	M	M
y_m	$a_{m1} \quad a_{m2} \quad \cdots \quad a_{mn}$	$\leq$	b_m
对偶关系	$\geq \quad \geq \quad \cdots \quad \geq$	maxZ=minW	
maxZ	$C_1 \quad C_2 \quad \cdots \quad C_n$		

当原问题的约束条件中含有等式约束方程时，原问题与其对偶问题之间的变换关系就是非对称形式了。这时，对偶问题的求法可按以下步骤进行：

(1) 首先将原问题的约束条件中的每一个等式约束方程都用两个不等式约束方程代替，从而使原问题的约束条件中的所有约束方程都变为同号不等式约束。

(2) 按对称形式变换关系 (表 6-4) 写出它的对偶问题。

譬如 对于线性规划问题

$$\begin{cases} \max Z = \sum_{j=1}^n c_j x_j \\ \sum_{j=1}^n a_{ij} x_j = b_i \quad (i = 1, 2, \cdots, m) \\ x_j \geq 0 \quad (j = 1, 2, \cdots, n) \end{cases}$$

可以按下述步骤求出其对偶问题：

第一步，将约束条件中的每一个约束方程 (等式约束) 都分解为两个不等式约束方程。这样，上述线性规划问题就可表示为

$$\begin{cases} \max Z = \sum_{j=1}^n c_j x_j & (7) \end{cases}$$

$$\begin{cases} \sum_{j=1}^n a_{ij} x_j \leq b & (8) \end{cases}$$

$$\begin{cases} \sum_{j=1}^n (-a_{ij}) x_j \leq -b & (9) \end{cases}$$

$$\begin{cases} x_j \geq 0 (j = 1, 2, \dots, n) & (10) \end{cases}$$

第二步，设 y_i' 和 y_i'' ($i = 1, 2, \dots, m$) 分别代表对应于 (8) 式和 (9) 式的对偶变量，则可以按对称形式变换关系 (表 6-4) 写出它的对偶问题

$$\begin{cases} \min W = \sum_{i=1}^m b_i y_i' + \sum_{i=1}^m (-b_i y_i'') \\ \sum_{i=1}^m a_{ij} y_i' + \sum_{i=1}^m (-a_{ij} y_i'') \geq c_j (j = 1, 2, \dots, m) \\ y_i' \geq 0, y_i'' \geq 0 (i = 1, 2, \dots, m) \end{cases}$$

上述线性规划问题的各式，经过整理后得到：

$$\begin{cases} \min W = \sum_{i=1}^m b_i (y_i' - y_i'') \\ \sum_{i=1}^m a_{ij} (y_i' - y_i'') \geq c_j (j = 1, 2, \dots, n) \\ y_i' \geq 0, y_i'' \geq 0 (i = 1, 2, \dots, m) \end{cases}$$

令 $y_i = y_i' - y_i''$ ，由于 $y_i' \geq 0, y_i'' \geq 0$ ，可见 y_i 不受正、负限制，将 y_i 代入上述线性规划问题，便可得到原线性规划问题的对偶问题

$$\begin{cases} \min W = \sum_{i=1}^m b_i y_i \\ \sum_{i=1}^m a_{ij} y_i \geq c_j (j = 1, 2, \dots, n) \\ y_i (i = 1, 2, \dots, m) \text{无约束} \end{cases}$$

二、原问题与对偶问题的关系

综上所述，线性规划原问题与对偶问题之间的形式变换关系可以由表 3-5 予以概述。

利用表 3-5 所描述的形式变换关系，我们可以写出任何一个线性规划问题的对偶问题。譬如，对于线性规划问题

$$\left\{ \begin{array}{l} \min Z = 2x_1 + 3x_2 - 5x_3 + x_4 \\ x_1 + x_2 - 3x_3 + x_4 \geq 5 \\ 2x_1 + 2x_3 - x_4 \leq 4 \\ x_2 + x_3 + x_4 = 6 \\ x_1 \leq 0; x_2, x_3 \geq 0; x_4 \text{ 无约束} \end{array} \right.$$

其对偶问题为

$$\left\{ \begin{array}{l} \max Z' = 5y_1 + 4y_2 + 6y_3 \\ y_1 + 2y_2 \geq 2 \\ y_1 + y_3 \leq 3 \\ -3y_1 + 2y_2 + y_3 \leq -5 \\ y_1 - 2y_2 + y_3 = 1 \\ y_1 \leq 0, y_2 \geq 0, y_3 \text{ 无约束} \end{array} \right.$$

表 3-5

原问题（或对偶问题）		对偶问题（或原问题）	
目标函数 maxZ		目标函数 minW	
变	个数 n	约	个数 n
量	≥ 0	束	$\geq$
	≤ 0	方	$\leq$
	无约束	程	$=$
约	个数 m	变	m
束	$\leq$		≥ 0
方	$\geq$	量	≤ 0
程	$=$		无约束
约束方程右端项		目标函数中变量的系数	
目标函数中变量的系数		约束方程右端项	

可以证明，原线性规划问题的对偶问题具有以下基本性质：

(1) 对称性。即对偶问题的对偶是原问题。

(2) 弱对偶性。即若 $\bar{X}$ 是原问题的可行解， $\bar{Y}$ 是对偶问题的可行解，则存在如下关系：

$$c\bar{X} \leq \bar{Y}b$$

(3) 无界性。若原问题（对偶问题）为无界解，则其对偶问题（原问题）无可行解。

(4) 对偶定理：若原问题有最优解，那么对偶问题也有最优解，而且它们的最优目标值相等。

(5) 松紧定理：若 X^* 与 Y^* 分别为原问题与对偶问题的可行解，则它们为最优解的充要条件为 $cX^* = Y^*b$

(6) 设原问题是

$$\begin{cases} \max Z = cX \\ AX + X_s = b \\ X \geq 0, X_s \geq 0 \end{cases}$$

它的对偶问题是

$$\begin{cases} \min W = Yb \\ YA - Y_s = c \\ Y \geq 0, Y_s \geq 0 \end{cases}$$

则原问题单纯形表的检验数行对应其对偶问题的一个基解，其对应关系如下表（表 3-6）。

表 3-6

X_B	X_N	X_s
0	$C_N - C_B B^{-1} N$	$-C_B B^{-1}$
Y_{s1}	$-Y_{s2}$	$-Y$

（7）互补松弛性。若 $\hat{X}$ 和 $\hat{Y}$ 分别是原问题和对偶问题的可行解。

那么 $\hat{Y} X_s = 0$ 和 $Y_s \hat{X} = 0$ ，当且仅当 $\hat{X}$ 和 $\hat{Y}$ 为最优解。

三、对偶单纯形法

由线性规划原问题与对偶问题的解之间的关系可知，在单纯形表中进行迭代时，在 b 列中得到的是原问题的基可行解，而在检验数行中得到的是对偶问题的基解。通过逐步迭代，当在检验数行中得到对偶问题的解也是基可行解时，则它就是最优解，这时原问题与对偶问题都是最优解。

根据对偶问题的对称性，我们也可以这样考虑问题：若保持对偶问题的解是基可行解，即 $C_j - C_B B^{-1} P_j \leq 0$ ，而原问题在非可行解的基础上，通过逐步迭代达到基可行解，这样也就得到了最优解。这种思路就是对偶单纯形法求解线性规划的思想。这种方法的优点是原问题的初始解不一定是基可行解，可从非基可行解开始迭代。对偶单纯形法的基本原理如下：

对于原问题

$$\begin{cases} \max Z = cX \\ AX = b \\ X \geq 0 \end{cases}$$

设 B 是一个基，不失一般性，令 $B = [P_1, P_2, \dots, P_m]$ ，它的应基变量为 $X_B = [x_1, x_2, \dots, x_m]$ 。当非基变量都为零时，可以得到 $X_B = B^{-1}b$ 。若在 $B^{-1}b$ 中至少有一个负分量，设 $(B^{-1}b)_i < 0$ ，并且在单纯形表的检验数行中的检验数都为非正，即对偶问题保持可行解，它的各分量是

（1）对应基变量 $x_1, x_2, \dots, x_m$ 的检验数是

$$\sigma_i = C_i - Z_i = C_i - C_B B^{-1} P_i = 0, i = 1, 2, \dots, m$$

（2）对应非基变量 $x_{m+1}, \dots, x_n$ 的检验数是

$$\sigma_j = C_j - Z_j = C_j - C_B B^{-1} P_j \leq 0, j=m+1, \dots, n$$

每次迭代是将基变量中的负分量 x_l 取出，去替换非基变量中的 x_k ，经基变换，所有检验数仍保持非正。从原问题来看，经过每次迭代，原问题由非可行解往可行解靠近，当原问题得到可行解时，便得到了最优解。

对偶单纯形法的计算步骤如下：

(1)根据线性规划问题，列出初始单纯形表，检查 b 列中的各分量，若都为非负，且检验数都为非正，则已得到最优解，停止计算。若 b 列中至少有一个负分量，检验数保持非正，则需要进行以下计算。

(2)确定换出变量。按照法则

$$\min_i \{ (B^{-1}b)_i \mid (B^{-1}b)_i < 0 \} = (B^{-1}b)_l$$

确定对应的基变量 x_l 为换出变量。

(3)确定换入变量。在单纯形表中检查 x_l 所在行的各系数 a_{lj} ($j=1, 2, \dots, n$)，若所有 $a_{lj} \geq 0$ ，则无可行解，停止计算。若存在 $a_{lk} < 0$ ($j=1, 2, \dots, n$)，计算

$$\theta = \min_j \left\{ \frac{C_j - Z_j}{a_{lj}} \mid a_{lj} < 0 \right\} = \frac{C_k - Z_k}{a_{lk}}$$

按 θ 规则所对应的列的非基变量 x_k 为换入变量，这样以保持得到的对偶问题的解仍为可行解。

(4)以 a_{lk} 为主元素，按原单纯形法在表中进行迭代运算，得出新的单纯形表。

(5)重复 (1) — (4) 的步骤，直至求得最优解。

例 2：试用对偶单纯形法求解如下线性规划问题

$$\begin{cases} \min W = 2x_1 + 3x_2 + 4x_3 \\ x_1 + 2x_2 + x_3 \geq 3 \\ 2x_1 - x_2 + 3x_3 \geq 4 \\ x_1, x_2, x_3 \geq 0 \end{cases}$$

解：首先将这一问题化成下列型式，以便得到对偶问题的初始可行基。

$$\begin{cases} \max Z = -2x_1 - 3x_2 - 4x_3 \\ -x_1 - 2x_2 - x_3 + x_4 = -3 \\ -2x_1 + x_2 - 3x_3 + x_5 = -4 \\ x_j \geq 0, j=1, 2, \dots, 5 \end{cases}$$

该问题的初始单纯形表，如表 3-7 所示。

表 3-7

$C_j \rightarrow$			-2	-3	-4	0	0
C_B	X_B	b	x_1	x_2	x_3	x_4	x_5
0	x_4	-3	-1	-2	-1	1	0
0	x_5	-4	[-2]	1	-3	0	1
$C_j - Z_j$			-2	-3	-4	0	0

从表 3-7 看到，检验数行对应的对偶问题的解是可行的。因为 b 列各分量均为负，所以需要进行迭代运算。

换出变量的确定：按照前述对偶单纯形法计算步骤 (2)，计算可得

$$\min_i \{ (B^{-1}b)_i \mid (B^{-1}b)_i < 0 \} \\ = \min \{-3, -4\} = -4$$

故 x_5 为换出变量。

换入变量的确定：按照前述对偶单纯形法计算步骤 (3)，计算可得

$$\theta = \min_j \left\{ \frac{C_j - Z_j}{a_{lj}} \mid a_{lj} < 0 \right\} = \min \left\{ \frac{-2}{-2}, \frac{-4}{-3} \right\} \\ = \frac{-2}{-2} = 1$$

故 x_1 为换入变量。换入、换出变量所在列、行的交叉处“2”为主元项。

按单纯形法计算步骤进行迭代运算，得表 3-8。

表 3-8

$c_j \rightarrow$			-2	-3	-4	0	0
c_B	x_B	b	x_1	x_2	x_3	x_4	x_5
0	x_4	-1	0	[-5/2]	1/2	1	-1/2
-2	x_1	2	1	-1/2	3/2	0	-1/2
$c_j - Z_j$			0	-4	-1	0	-1

由表 3-8 看出，对偶问题仍是可行解，而 b 列中仍有负分量，故还需进行迭代运算。重复上述迭代步骤，得到表 3-9。

表 3-9

$c_j \rightarrow$			-2	-3	-4	0	0
c_B	x_B	b	x_1	x_2	x_3	x_4	x_5
-3	x_2	2/5	0	1	-1/5	-2/5	1/5
-2	x_1	11/5	1	0	7/5	01/5	-2/5
$c_j - Z_j$			0	0	-3/5	-8/5	-1/5

在表 3-9 中，b 列各分量全为非负，检验数全为非正，故问题的最优解为 $X^* = [11/5, 2/5, 0, 0, 0]^T$ ，而其对偶问题的最优解为 $Y^* = [8/5, 1/5]$ 。

第三节 运输问题的求解方法——表上作业法

一、产销平衡表与单位运价表

运输问题，作为一类特殊的线性规划问题，我们已在本章第一节中介绍其数学模型，对于这类问题，除用线性规划数学模型进行描写之外，还可用产销平衡表和单位运价表进行描述。

假设某种物资有 m 个生产地点 A_i ($i=1, 2, \dots, m$)，其产量（供应量）分别为 a_i ($i=1, 2, \dots, m$)，有 n 个销地 B_j ($j=1, 2, \dots, n$)，其销量（需求量）分别为 b_j ($j=1, 2, \dots, n$)。从 A_i 到 B_j 运输单位物资的运价（单价）为 c_{ij} 。将这些数据汇总可以得到产销平衡表（表 3-10）和单位运价表（表 3-11），将表 3-10 和表 3-11 合并，可得表 3-12。

表 3-10 产销平衡表

产地 \ 销地	1	2	...	n	产 量
1					a_1
2					a_2
M					M
m					a_m
销量	b_1	b_2	...	b_n	

表 3-11 单位运价表

产地 \ 销地	1	2	...	n	
1	c_{11}	c_{12}	...	c_{1n}	
2	c_{21}	c_{22}	...	c_{2n}	
M	M	M			
m	c_{m1}	c_{m2}	...	c_{mn}	

表 3-12

产地 \ 销地	1	2	...	n	产 量
1	c_{11}	c_{12}	...	c_{1n}	a_1
2	c_{21}	c_{22}	...	c_{2n}	a_2
M	M	M			M
m	c_{m1}	c_{m2}	...	c_{mn}	a_m
销 量	b_1	b_2	...	b_n	

二、表上作业法

由于运输问题是一类特殊的线性规划问题，其约束方程组的系数矩阵具有特殊的结构，因而可以找到一种比单纯形法更为简便的求解方法，即表上作业法。

表上作业法是单纯形法在求解运输问题时的一种简化方法，其实质仍是单纯形法，只是其具体计算过程和使用的有关术语有所不同。这一方法的计算步骤可以归纳如下：

- (1)找出初始基可行解，即在产销平衡表上给出 $m+n-1$ 个数字格。
- (2)求各非基变量的检验数，即在表上计算空格的检验数，判别是否达到最优解。如果已是最优解，则停止计算，否则转入下一步。
- (3)确定换入变量和换出变量，找出新的基可行解（即在表上用闭回路法进行调整）。
- (4)重复 (2) — (3)，直到求出最优解为止。

下面，我们通过具体例子来说明表上作业法的计算步骤。

例 3：假设某种物资共有三个产地，其日产量分别是： A_1 为 7 吨， A_2 为 4 吨， A_3 为 9 吨，该种物资的四个销售地，其日销售量分别是： B_1 为 3 吨， B_2 为 6 吨， B_3 为 5 吨， B_4 为 6 吨；各产地到销售地的单位物资的运价如表 3-13 所示。试问在满足各销售点需要量的前提下，如何调运该种物资，才能使总运费达到最小？

表 3-13 某物资运输的单位运价表

销地 产地	B_1	B_2	B_3	B_4
A_1	3	11	3	10
A_2	1	9	2	8
A_3	7	4	10	5

解：首先列出这一问题的产销平衡表，见表 3-14。

表 3-14 某物资运输的产销平衡表

销地 产地	B_1	B_2	B_3	B_4	产量
A_1					7
A_2					4
A_3					9
销 量					

(一)确定初始基可行解

确定初始基可行解的方法很多，一般而言，人们所希望的方法是既简便，又尽可能地接近最优解。在用表上作业法求解运输问题时，确定初始基可行解的最常用的方法有最小元素法和伏格尔（Vogel）法。

1.最小元素法 这一方法的基本思想是就近供应，即从单位运价表中最小的运价开始确定供销关系，然后考虑运价次小的，一直到给出初始基可行解为止。以上例 3 为例，这一方法可由以下几个步骤来完成。

第一步：从表 3-13 中找出最小运价为 1，表示应先将 A_2 的产品供应 B_1 。

因为 $a_2=4>b_1=3$ ，故 A_2 除满足 B_1 的全部需要外，还可多余 1 吨物资。在表 3-14 中 (A_2, B_1) 的交叉格处填上 3，得表 3-15。将表 3-13 中的 B_1 列运价划去，得表 3-16。

第二步：在表 3-16 未划去的元素中再找出最小运价 2，确定 A_2 多余的 1 吨物资供应 B_3 ，得表 3-17。将表 3-16 中 A_2 行运价划去，得表 3-18。

表 3-15

销地 产地	B_1	B_2	B_3	B_4	产量
A_1					7
A_2	3				4
A_3					9
销 量	3	6	5	6	

表 3-16

销地 产地	B_1	B_2	B_3	B_4
A_1		11	3	10
A_2		9	2	8
A_3		4	10	5

表 3-17

销地 产地	B_1	B_2	B_3	B_4	产量
A_1					7
A_2	3		1		4
A_3					9
销 量	3	6	5	6	

表 3-18

销地 产地	B_1	B_2	B_3	B_4
A_1		11	3	10
A_2				
A_3		4	10	5

第三步：在表 3-18 未划去的元素中再找出最小运价 3，这表示应将 A_1 的产品供应给 B_3 ，因为 $b_3=5$ ，而 A_2 已供应给 B_3 1 吨的产品，故 A_1 应供应 B_3 4 吨产品。在表 3-17 中 (A_1, B_3) 的位置上填上 4。将表 3-18 中 B_3 列的运价划去。……这样一直下去，直到单位运价表上的所有元素被划去为止。最后在产销平衡表上得到一个调运方案，即初始基可行解，见表 3-19。

表 3—19

销地 \ 产地	B ₁	B ₂	B ₃	B ₄	产量
A ₁			4	3	7
A ₂	3		1		4
A ₃		6		3	9
销 量	3	6	5	6	

2. 伏格尔法

最小元素法的缺点是，为了节省一处的费用，有时造成在其它处要多花几倍的运费。伏格尔法考虑到，一产地的产品假如不能按最小运费就近供应，就考虑次小运费，这就有一个差额，差额越大，说明不能按最小运费调运时，运费增加越多。因而对差额最大处，就应当采用最小运费调运。伏格尔法的步骤是：

第一步：在表 3-13 中分别计算出各行、各列的最小运费和次最小运费的差额，并填入该表的最右列和最下行，见表 3-20。

表 3—20

销地 \ 产地	B ₁	B ₂	B ₃	B ₄	行差额
A ₁	3	11	3	10	0
A ₂	1	9	2	8	1
A ₃	7	4	10	5	1
列差额	2	5	1	3	

第二步：从行或列差额中选出最大者，选择它所在行或列中的最小元素。在表 3-20 中，B₂ 列是最大差额所在列，B₂ 列中最小元素为 4，因此可确定 A₃ 的产品应首先供应 B₂ 的需要，由此可得表 3-21。将单位运价表中 B₂ 列的数字划去，得表 3-22。

第三步：对表 3-22 中未划去的元素再分别计算出各行、各列的最小运费和次最小运费的差额，并填入该表的最右列和最下行。重复第一、二步，直到给出初始基可行解为止。用此方法得到的初始基可行解列于表 3-23。

可以看出，伏格尔法与最小元素法除在确定供求关系的原则上不同外，其余步骤均相同。伏格尔法给出的初始基可行解比最小元素法给出的初始基可行解更接近最优解。本例中用伏格尔法给出的初始基可行解就是最优解。

表 3-21

销地 \ 产地	B ₁	B ₂	B ₃	B ₄	产量
A ₁					7
A ₂					4
A ₃					9
销 量	3	6	5	6	

表 3-22

产地 \ 销地	B ₁	B ₂	B ₃	B ₄
A ₁	3		3	10
A ₂	1		2	8
A ₃	7		10	5

表 3-23

产地 \ 销地	B ₁	B ₂	B ₃	B ₄	产量
A ₁			5	2	7
A ₂	3			1	4
A ₃				3	9
销 量	3	6	5	6	

(二)最优解的判别

要判别一个基可行解是否最优解，就需要计算非基变量（空格）的检验数 $C_{ij} - C_B B^{-1} P_{ij}$ ($i, j \in N$)。因为运输问题的目标函数是求最小值，故当所有的 $C_{ij} - C_B B^{-1} P_{ij} \geq 0$ 时，为最优解。下面，我们将介绍两种求空格检验数的方法。

1. 闭回路法在给出调运方案的计算表上，如表 3-19，从每一空格出发找一条闭回路，它是以某一空格为起点，用水平或垂直线向前划，每碰到一数字格转 90° 后，继续前进，直到回到起始空格为止。闭回路如图 3-2 的 (a)，(b)，(c) 等所示。

可以证明，从每一个空格出发一定存在并且可以找到唯一的闭回路。因为 $m+n-1$ 个数字格（基变量）对应的系数向量是一个基，任一空格（非基变量）对应的系数向量是这个基的线性组合。譬如， P_{ij} ($i, j \in N$) 可表示为：

$$\begin{aligned}
 P_{ij} &= e_i + e_{m+j} \\
 &= e_i + e_{m+k} - e_{m+k} + e_l - e_l + e_{m+s} - e_{m+s} + e_u - e_u + e_{m+j} \\
 &= (e_i + e_{m+k}) - (e_l + e_{m+k}) + (e_l + e_{m+s}) - (e_u + e_{m+s}) + (e_u + e_{m+j}) \\
 &= P_{ik} - P_{lk} + P_{ls} - P_{us} + P_{uj}
 \end{aligned} \tag{1}$$

在 (1) 式中， $P_{ik}, P_{lk}, P_{ls}, P_{uj} \in B$ ，而这些向量构成了闭回路（见图 3-3）。

闭回路法计算检验数的经济解释为：在已给出初始基可行解的表 3-19 中，可从任一空格出发，如 (A_1, B_1) ，若让 A_1 的产品调运 1 吨给 B_1 ，为了保持产销平衡，就要依次进行调整：在 (A_1, B_3) 处减少 1 吨，在 (A_2, B_3) 处增加 1 吨， (A_2, B_1) 处减少 1 吨，即构成了以 (A_1, B_1) 空格为起点，其它为数字格的闭回路，如表 3-24 中的虚线所示。在这个表中，闭回路各顶点所在格的右上角数字是单位运价。

调整的方案使运费增加

$$(+1) \times 3 + (-1) \times 3 + (+1) \times 2 + (-1) \times 1 = 1 \text{ (元)}$$

这表明若作这样的运量调整，将会使运费增加 1 元。将“1”填 (A_1, B_1) 格中，这就是检验数。按照上述办法，可找出所有空格的检验数，见表 3-25。当检验数还存在负数时，说明原方案不是最优解，还需要对原方案进行改进。

表 3-25

空格	闭回路	检验数
(A_1, B_1)	$(A_1, B_1) - (A_1, B_3) - (A_2, B_3) - (A_2, B_1) - (A_1, B_1)$	1
(A_1, B_2)	$(A_1, B_2) - (A_1, B_4) - (A_3, B_4) - (A_3, B_2) - (A_1, B_2)$	2
(A_2, B_2)	$(A_2, B_2) - (A_2, B_3) - (A_1, B_3) - (A_1, B_4) - (A_3, B_4) - (A_3, B_2) - (A_2, B_2)$	
(A_2, B_4)	$(A_2, B_4) - (A_2, B_3) - (A_3, B_3) - (A_1, B_4) - (A_2, B_4)$	
(A_3, B_1)	$(A_3, B_1) - (A_3, B_4) - (A_1, B_4) - (A_1, B_3) - (A_2, B_3) - (A_2, B_1) - (A_3, B_1)$	
(A_3, B_3)	$(A_3, B_3) - (A_3, B_4) - (A_1, B_4) - (A_1, B_3) - (A_3, B_3)$	

2. 位势法

用闭回路法求检验数时，需给每一个空格找一条闭回路。当产销地很多时，这种计算就变得更繁。下面我们介绍一种较为简便的求检验数方法——位势法。

设 $u_1, u_2, \dots, u_m; v_1, v_2, \dots, v_n$ 是对应运输问题的 $m+n$ 个约束条件的对偶变量。 B 是含有一个人工变量 x_a 的 $(m+n) \times (m+n)$ 阶初始基矩阵。人工变量 x_a 在目标函数中的系数 $c_a=0$ ，由线性规划的对偶理论可知

$$C_B B^{-1} = [u_1, u_2, \dots, u_m; v_1, v_2, \dots, v_n] \quad (2)$$

而每一个决策变量 x_{ij} 的系数向量 $P_{ij} = e_i + e_{m+j}$ ，所以 $C_B B^{-1} P_{ij} = u_i + v_j$ 。于是检验数为

$$\sigma_{ij} = C_{ij} - C_B B^{-1} P_{ij} = C_{ij} - (u_i + v_j) \quad (3)$$

由单纯形法知所有基变量的检验数等于 0，即

$$C_{ij} - (u_i + v_j) = 0 \quad i, j \in B \quad (4)$$

譬如，在例 3 由最小元素法得到的初始解中， $x_{24}, x_{34}, x_{21}, x_{32}, x_{13}, x_{14}$ 是基变量。 x_a 为人工变量，这时对应的检验数是

基变量	检验数	
x_a	$c_a - u_1 = 0$	因为 $c_a = 0$ ，所以 $u_1 = 0$
x_{24}	$c_{24} - (u_2 + v_4) = 0$	即 $8 - (u_2 + v_4) = 0$
x_{34}	$c_{34} - (u_3 + v_4) = 0$	即 $5 - (u_3 + v_4) = 0$
x_{21}	$c_{21} - (u_2 + v_1) = 0$	即 $1 - (u_2 + v_1) = 0$
x_{32}	$c_{32} - (u_3 + v_2) = 0$	即 $4 - (u_3 + v_2) = 0$
x_{13}	$c_{13} - (u_1 + v_3) = 0$	即 $3 - (u_1 + v_3) = 0$
x_{14}	$c_{14} - (u_1 + v_4) = 0$	即 $10 - (u_1 + v_4) = 0$

从以上 7 个方程中可求得： $u_1 = 0, u_2 = -1, u_3 = -5, v_1 = 2, v_2 = 9, v_3 = 3, v_4 = 10$

对于非基变量的检验数

$$\sigma_{ij} = c_{ij} - (u_i + v_j) \quad i, j \in N \quad (5)$$

可以由已知的 u_i, v_j 求得。所有这些计算均可以在表中进行，以下我们以例 3 说明之。

第一步：按最小元素法给出表 3-19 的初始基可行解，作表 3-26。在对应表 3-19 的数字格处填入单位运价，见表 3—26。

表 3-26

销地 \ 产地	B ₁	B ₂	B ₃	B ₄
A ₁			3	10
A ₂	1		2	5
A ₃		4		

第二步：在表 3-26 中增加一行一列，在列中填入 u_i ，在行中填入 v_j ，得表 3-27。

首先令 $u_1=0$ ，然后按 $u_i+v_j=c_{ij} \quad (i, j \in B)$ 相继确定 u_i, v_j 。由表 3-27 可见，当 $u_1=0$ 时，由 $u_1+v_3=3$ 可得 $v_3=3$ ，由 $u_1+v_4=10$ 可得 $v_4=10$ ；进而再由 $u_3+v_4=5$ 可得 $u_3=-5$ ，以此类推可确定所有 u_i 和 v_j 的数值。

第三步：按 $\sigma_{ij}=c_{ij} - (u_i+v_j) \quad (i, j \in N)$ 计算所有空格的检验数。譬如， $\sigma_{11}=c_{11} - (u_1+v_1)=3 - (0+2)=1$ ， $\sigma_{12}=c_{12} - (u_1+v_2)=11 - (0+9)=2$ 。这些计算可直接在表 3-27 上进行。为了方便，特设计计算表，如表 3-28 所示。

表 3-27

销地 \ 产地	B ₁	B ₂	B ₃	B ₄	u_i
A ₁			3	10	0
A ₂	1		2		-1
A ₃		4		5	-5
v_j	2	9	3	10	

表 3-28

销地 \ 产地	B ₁	B ₂	B ₃	B ₄	u_i
A ₁	1 3	2 11	0 3	0 10	0
A ₂	0 1	1 9	0 2	-1 8	-1
A ₃	10 7	0 4	12 10	0 5	-5
v_j	2	9	3	10	

(三)改进的方法——闭回路调整法

在计算表中，如果在空格处出现负检验数，则表明当前解不是最优解。若有两个和两个以上的负检验数时，一般选其中最小的负检验数，以它对应的空格为调入格，即以它对应的非基变量为换入变量。在表 3-28 中， (A_2, B_4) 为调入格，以此格为出发点，作一闭回路，如表 3-29 所示。

表 3-29

销地 产地	B_1	B_2	B_3 B_4	产 量
A_1			$4(+1) \cdots 3(-1)$	7
A_2	3		$1(-1) \cdots (+1)$	4
A_3		6		9
销 量	3	6	5 6	

(A_2, B_4) 格的调入量 θ 是选择闭回路上具有 (-1) 的数字格中的最小者，即 $\theta = \min(1, 3) = 1$ ，然后，按闭回路上的正、负号，加、减此值得到调整方案，如表 3-30 所示。

表 3-30

销地 产地	B_1	B_2	B_3	B_4	产 量
A_1	3		5	2	7
A_2				1	4
A_3		6		3	9
销 量	3	6	5	6	

对表 3-30 给出的解，再用闭回路法或位势法求各空格的检验数，得表 3-31。在表 3-31 中，因为所有检验数都非负，故表 3-30 所给出的解是最优解，这时得到的总运费（最小运费）为 85（元）。

销地 产地	B_1	B_2	B_3	B_4
A_1	0	2		
A_2		2	1	
A_3	9		12	

三、产销不平衡的运输问题的求解方法

前面所介绍的求解运输问题的表上作业法，是以产销平衡，即

$$\sum_{i=1}^m a_i = \sum_{j=1}^n b_j \quad (6)$$

为前提的。但是，在实际问题中，产销往往是不平衡的。为了求解，需要把

产销不平衡的问题化成产销平衡的问题。

当产大于销

$$\sum_{i=1}^m a_i > \sum_{j=1}^n b_j \quad (7)$$

时，运输问题的数学模型为：求 x_{ij} ($i=1, 2, \dots, m$; $j=1, 2, \dots, n$)

$$\text{使：} \quad \min Z = \sum_{i=1}^m \sum_{j=1}^n c_{ij} x_{ij} \quad (8)$$

且满足：

$$\begin{cases} \sum_{j=1}^n x_{ij} \leq a_i \quad (i=1, 2, \dots, m) & (9) \\ \sum_{i=1}^m x_{ij} = b_j \quad (j=1, 2, \dots, n) & (10) \\ x_{ij} \geq 0 \quad (i=1, 2, \dots, m, j=1, 2, \dots, n) & (11) \end{cases}$$

由于总的产量大于销量，所以就要考虑多余的物资在哪一个产地就地贮存的问题。设 $x_{i, n+1}$ 是产地 A_i 的贮存量，于是有：

$$\sum_{j=1}^n x_{ij} + x_{i, n+1} = \sum_{j=1}^{n+1} x_{ij} = a_i \quad (i=1, 2, \dots, m) \quad (12)$$

$$\sum_{i=1}^m x_{ij} = b_j \quad (j=1, 2, \dots, n) \quad (13)$$

$$\sum_{i=1}^m x_{i, n+1} = \sum_{i=1}^m a_i - \sum_{j=1}^n b_j = b_{n+1} \quad (14)$$

当 $i=1, 2, \dots, m, j=1, 2, \dots, n$ 时，令 $c'_{ij}=c_{ij}$ ；当 $i=1, 2, \dots, m, j=n+1$ 时，令 $c'_{ij}=0$ ，将其代入 (8) 式，并将 (9) 式进行改写，产销不平衡的运输问题 (8)- (11) 式就可以改写成：求 x_{ij} ($i=1, 2, \dots, m$; $j=1, 2, \dots, n, n+1$)

$$\begin{aligned} \text{使} \quad \min Z' &= \sum_{i=1}^m \sum_{j=1}^{n+1} c'_{ij} x_{ij} = \sum_{i=1}^m c'_{i, n+1} x_{i, n+1} \\ &= \sum_{i=1}^m \sum_{j=1}^n c_{ij} x_{ij} \end{aligned} \quad (8')$$

且满足：

$$\begin{cases} \sum_{j=1}^{n+1} x_{ij} = a_i \quad (i=1, 2, \dots, m) & (9') \\ \sum_{i=1}^m x_{ij} = b_j \quad (j=1, 2, \dots, n, n+1) & (10') \\ x_{ij} \geq 0 \quad (i=1, 2, \dots, m; j=1, 2, \dots, n, n+1) & (11') \end{cases}$$

在这个模型中，由于

$$\sum_{i=1}^m a_i = \sum_{j=1}^n b_j + b_{n+1} = \sum_{j=1}^{n+1} b_j \quad (12)$$

成立，所以这是一个产销平衡的运输问题。

当产大于销时，只要增加一个假想的销地 $j=n+1$ （实际上是贮存），让该销地的需求量为 $(\sum_{i=1}^m a_i - \sum_{j=1}^n b_j)$ ，并在单位运价表中令从各产地到假想的销地的单位运价为 $c'_{i, n+1}=0$ ($i=1, 2, \dots, m$)，就可以将其转化为一个产销平衡的运输问题。同理，当销大于产时，只要在产销平衡表中增加一个假想的产地 $i = m + 1$ ，让该地的产量为 $(\sum_{j=1}^n b_j - \sum_{i=1}^m a_i)$ ，并在单位运

价中令从该假想的产地到各销费地的运价为 $c'_{m+1, j}=0$ ，就可以将其转化为一个产销平衡的运输问题。

对于产销不平衡的运输问题，当用上述方法将其转化为产销平衡的运输问题后，就可以用表上作业法对其求解，对此，我们不拟再作过多地赘述。

第四章 多目标规划方法

在地理学研究中，对于许多规划问题，常常需要考虑多个目标，如经济效益目标，生态效益目标，社会效益目标，等等。为了满足这类问题研究之需要，本章拟结合有关实例，对多目标规划方法及其在地理学研究中的应用问题作一些介绍和探讨。

第一节 多目标规划及其求解技术简介

一、多目标规划及其非劣解

(一)多目标规划数学模型

传统的单目标规划问题，一般可用如下数学模型描述：

$$\max(\min) Z=f(X) \quad (1)$$

$$\Phi(X) \leq G \quad (2)$$

在上式中， $X=[x_1, x_2, \dots, x_n]^T$ 为规划决策变量向量； $Z=f(X)$ 是多元标量函数； $\Phi(X)$ 是 m 维函数向量； G 是 m 维的常数向量， m 是约束条件个数。

对于多目标规划问题，其数学模型也可以类似地描写为如下形式：

$$\max(\min) Z=F(X) \quad (3)$$

$$\Phi(X) \leq G \quad (4)$$

在上式中， $X=[x_1, x_2, \dots, x_n]^T$ 为规划决策变量向量； $Z=F(X)$ 是 K 维函数向量， K 是目标函数的个数； $\Phi(X)$ 是 m 维函数向量； G 是 m 维常数向量， m 是约束方程的个数。

对于线性多目标规划问题，(3) — (4) 式可以进一步写作：

$$\max(\min) Z=AX \quad (5)$$

$$BX \leq b \quad (6)$$

在 (5) 和 (6) 式， A 为 $K \times m$ 阶矩阵； B 为 $m \times n$ 阶矩阵， b 为 m 维的向量； X 为 n 维决策变量向量。

(二)多目标规划的非劣解

对于上述多目标规划问题，求解就意味着需要作出如下的复合选择：

(1) 每一个目标函数取什么值，原问题可以得到最满意的解决？

(2) 每一个决策变量取什么值，原问题可以得到最满意的解决？

在单目标规划问题中，各种方案的目标值之间是可以比较的，因此各种方案总是可以分出优劣的。但在多目标规划中，问题就变得比较复杂。例如，当规划问题是要求所有的目标都取最大值时，一个目标值的增大就有可能导致另一个目标值的减小。因此，多目标规划问题的求解就不可能象在单目标规划中那样，只追求一个目标的最优（最大或最小）化，而置其它目标于不顾。

在多目标规划问题的求解中，非劣解是一个十分重要的概念，对于这一概念，我们可用图 4-1 说明。在图 4-1 中，就方案①和②来说，①的目标值 f_2 比②大，但其目标值 f_1 比②小，因此无法确定这两个方案的优与劣。在各个方案之间，显然：③比②好，④比①好，⑦比③好，⑤比④好。而对于方案⑤、⑥、⑦，它们之间无法确定优劣，而且又没有比它们更好的其它方案，它们就被称之为多目标规划问题的非劣解或有效解，其余方案都称为劣解。所有非劣解构成的集合称为非劣解集。

二、多目标规划求解技术简介

多目标规划问题的求解，就是要在非劣解集中寻求一个最为满意的规划方案。然而，非劣解集中往往包含有许多非劣解，究竟哪一个最为满意呢？为了解决这一问题，常常需要将多目标规划问题转化为单目标规划问题去处理。实现这种转化，有以下几种建模方法可供借鉴。

(一)效用最优化模型

效用最优化模型建立的依据是基于这样一种假设：规划问题的各个目标函数可以通过一定的方式进行求和运算。这种方法将一系列的目标函数与效用函数建立相关关系，各目标之间通过效用函数协调，从而使多目标规划问题转化为传统的单目标规划问题：

$$\max Z = \psi(X) \quad (7)$$

$$\Phi(X) \leq G \quad (8)$$

(7)式中， ψ 是与各目标函数相关的效用函数的和函数。

在用效用函数作为规划目标时，需要确定一组权值 λ_i 来反映原问题中各目标函数在总体目标中的权重，即：

$$\max \psi = \sum_{i=1}^k \lambda_i \psi_i \quad (9)$$

$$\Phi_i(x_1, x_2, \dots, x_n) \leq g_i \quad (i=1, 2, \dots, m) \quad (10)$$

在(9)式中，诸 λ_i 应满足：

$$\sum_{i=1}^k \lambda_i = 1 \quad (11)$$

若采用向量与矩阵记号，则上述模型可以进一步改写为：

$$\max \psi = \lambda^T \psi \quad (12)$$

$$\Phi(X) \leq G \quad (13)$$

(二)罚款模型

如果对每一个目标函数，规划决策者都能提出一个所期望的值（或称满意值） f_i^* ，那么，就可以通过比较实际值 f_i 与期望值 f_i^* 之间的偏差来选择问题来选择问题的解。罚款模型的数学表达式如下：

$$\min Z = \sum_{i=1}^k a_i (f_i - f_i^*)^2 \quad (14)$$

$$\Phi(X) \leq G \quad (17)$$

或写成矩阵形式：

$$\min Z = (F - F^*)^T A (F - F^*) \quad (16)$$

$$\Phi(X) \leq G \quad (17)$$

在上式中， α_i 是与第 i 个目标函数相关的权重； A 是由诸 α_i ($i=1, 2, \dots, K$)组成的 $m \times m$ 阶对角矩阵。

(三)目标规划模型

目标规划模型与罚款模型类似，它也需要预先确定各个目标的期望值 f_i^* 。目标规划模型为数学形式为：

$$\min Z = \sum_{i=1}^k (f_i^+ + f_i^-) \quad (18)$$

$$\Phi_i(x_1, x_2, \dots, x_n) \leq g_i \quad (i=1, 2, \dots, m) \quad (19)$$

$$f_i + f_i^- - f_i^+ = f_i^* \quad (i=1, 2, \dots, m) \quad (20)$$

在上式中， f_i^+ 和 f_i^- 分别表示与 f_i 相应的、与 f_i^* 相比的目标超过值和不足值。

采用矩阵形式表示，则 (18)- (20)式可以进一步简记为：

$$\min Z = v^T (F^+ + F^-) \quad (21)$$

$$\Phi(X) \leq G \quad (22)$$

$$F + F^- - F^+ = F^* \quad (23)$$

在(21)式中， v 表示各元素均为 1 的 K 维列向量。

(四)约束模型

约束模型的立论依据是：如果规划问题的某一目标可以给出一个可供选择的范围，则该目标就可以作为约束条件而被排除出目标组，进入约束条件组中。

假如，除了第一个目标外，其余目标都可以提出一个可供选的范围，则按上述思路，该多目标规划问题就可以转化为单目标规划问题：

$$\max(\min Z) = f_1(x_1, x_2, \dots, x_n) \quad (24)$$

$$\Phi_i(x_1, x_2, \dots, x_n) \leq g_i \quad (i=1, 2, \dots, m) \quad (25)$$

$$f_j^{\min} \leq f_j \leq f_j^{\max} \quad (j=2, 3, \dots, K) \quad (26)$$

采用矩阵记号，上述模型可以进一步改写为如下形式：

$$\max(\min) Z = f_1(X) \quad (27)$$

$$\Phi(X) \leq G \quad (28)$$

$$F_1^{\min} \leq F_1 \leq F_1^{\max} \quad (29)$$

第二节 目标规划方法

目标规划方法，是在线性规划的基础上逐步发展起来的一种多目标规划方法。这一方法是由美国学者查恩斯 (A.Charnes) 和库伯 (W.W.Cooper) 于 1961 年首先提出来的。后来，查斯基莱恩 (U.Jaashelainen) 和李 (Sang.Lee) 等人在查恩斯和库伯研究工作的基础上，给出了求解目标规划问题的一般性方法。

一、目标规划数学模型

目标规划的基本思想是，给定若干目标以及实现这些目标的优先顺序，在有限的资源条件下，使总的偏离目标值的偏差最小。

为了具体说明目标规划与线性规划在处理问题的方法上的区别，我们首先通过下面的例子来介绍目标规划的有关概念及数学模型。

例：设某企业利用某种原材料和现有设备可生产甲、乙两种产品，其中，甲、乙两种产品的单价分别为 8 元和 10 元；生产单位甲、乙两种产品需要消耗的原材料分别为 2 个单位和 1 个单位，需要占用的设备分别为 1 台时和 2 台时；原材料拥有量为 11 个单位；可利用的设备总台时为 10 台时。试问：该企业领导者应如何确定其生产方案？

如果决策者所追求的唯一目标是使总产值达到最大，则这个企业的生产方案可以由如下的线性规划模型给出：求 x_1, x_2 使

$$\max Z = 8x_1 + 10x_2 \quad (1)$$

且满足：

$$\begin{cases} 2x_1 + x_2 \leq 11 & (2) \\ x_1 + 2x_2 \leq 10 & (3) \\ x_1, x_2 \geq 0 & (4) \end{cases}$$

在 (1) - (4) 式中， x_1 和 x_2 为决策变量， Z 为目标函数值。将上述问题化为标准后，用单纯形法求解可得最佳决策方案为 $x_1^* = 4$ ， $x_2^* = 3$ ， $Z^* = 62$ (元)。

但是，在实际决策时，企业领导者必须考虑市场等一系列其它条件，如：

(1) 根据市场信息，甲种产品的需求量有下降的趋势，故甲种产品的产量不应大于乙种产品的产量。

(2) 超过计划供应的原材料，需用高价采购，这就会使生产成本增加。

(3) 应尽可能地充分利用设备的有效台时，但不希望加班。

(4) 应尽可能达到并超过计划产值指标 56 元。这样，该企业生产方案的确定，便成为一个多目标决策问题，这一问题可用目标规划方法进行求解。

下面，我们引入与建立目标规划数学模型有关的一些概念。

1. 偏差变量 在目标规划数学模型中，除了决策变量外，还需要引入正、负偏差变量 d^+ 、 d^- 。其中，正偏差变量 d^+ 表示决策值超过目标值的部分，负偏差变量 d^- 表示决策值未达到目标值的部分。因为决策值不可能既超过目标值同时又未达到目标值，故恒有 $d^+ \times d^- = 0$ 成立。

2. 绝对约束和目标约束 绝对约束，是指必须严格满足的等式约束和不等式约束，譬如，线性规划问题的所有约束条件都是绝对约束，不能满足这些约束条件的解称为非可行解，所以它们是硬约束。

目标约束是目标规划所特有的，我们可以将约束方程右端项看作是追求的目标值，在达到此目标值时允许发生正或负偏差，因此在这些约束条件中加入正、负偏差变量，它们是软约束。线性规划问题的目标函数，在给定目标值和加入正、负偏差变量后可变换为目标约束，也可以根据问题的需要将绝对约束变换为目标约束。譬如，线性问题 (1) - (4) 式中的目标函数 $Z=8x_1+10x_2$ 可变换为目标约束 $8x_1+10x_2+d_1-d_1=56$ ，绝对约束 $2x_1+x_2\leq 11$ 可变换为 $2x_1+x_2+d_2-d_2=11$ 。

3. 优先因子 (优先等级) 与权系数 一个规划问题常常有若干个目标，决策者对这些目标的考虑。是有主次或轻重缓急的不同。凡要求第一位达到的目标赋予优先因子 P_1 ，次位的目标赋予优先因子 P_2 ，……，并规定 $P_k \gg P_{k+1}$

($k=1, 2, \dots, K$) 表示 P_k 比 P_{k+1} 有更大的优先权。这就是说，首先保证 P_1 级目标的实现，这时可以不考虑次级目标；而 P_2 级目标是在实现 P_1 级目标的基础上考虑的；以此类推。若要区别具有相同优先因子 P_k 的目标的差别，这时可分别赋予它们不同的权系数 k_l ($l=1, 2, \dots, L$)。这些都由决策者按具体情况而定。

4. 目标规划的目标函数 目标规划的目标函数 (准则函数) 是按各目标约束的正、负偏差变量和赋予相应的优先因子而构造的。当每一目标值确定后，决策者的要求是尽可能缩小偏离目标值。因此，目标规划的目标函数只能是：

$$\min Z = f(d^+, d^-) \quad (5)$$

(5) 式的基本形式有三种：

(1) 要求恰好达到目标值，即正、负偏差变量都要尽可能地小。这时，有

$$\min Z = f(d^+ + d^-) \quad (6)$$

(2) 要求不超过目标值，即允许达不到目标值，就是正偏差变量要尽可能地小。这时，有：

$$\min Z = f(d^+) \quad (7)$$

(3) 要求超过目标值，即超过量不限，但必须使负偏差变量要尽可能地小。这时，有：

$$\min Z = f(d^-) \quad (8)$$

对于每一个具体的目标规划问题，可根据决策者的要求和赋予各目标的优先因子来构造目标函数。

对例 1 所描述的生产方案的决策问题，如果决策者在原材料供应受严格限制的基础上考虑：首先是甲种产品的产量不超过乙种产品的产量；其次是充分利用设备的有效台时，不加班；再次是产值不小于 56 元。并分别赋予这三个目标 P_1, P_2, P_3 优先因子。则，这一决策问题的目标规划模型就是：

$$\min Z = P_1 d_1^+ + P_2 (d_2^- + d_2^+) + P_3 d_3^- \quad (9)$$

$$2x_1 + x_2 \leq 11 \quad (10)$$

$$x_1 - x_2 + d_1^- - d_1^+ = 0 \quad (11)$$

$$x_1 + 2x_2 + d_2^- - d_2^+ = 10 \quad (12)$$

$$8x_1 + 10x_2 + d_3^- - d_3^+ = 56 \quad (13)$$

$$x_1, x_2, d_i^-, d_i^+ \geq 0 \quad (i=1, 2, 3) \quad (14)$$

目标规划的一般性数学模型如下：

假定有 L 个目标， K 个优先级 ($K \leq L$)， n 个变量。在同一优先级 P_k 中不同目标的正、负偏差变量的权系数分别为 ω_{kl}^+ 、 ω_{kl}^- ，则多目标规划问题可以表示为：

$$\min Z = \sum_{k=1}^K P_k \sum_{l=1}^L (\omega_{kl}^- d_l^- + \omega_{kl}^+ d_l^+) \quad (15)$$

$$\sum_{j=1}^n c_j^{(l)} x_j + d_l^- - d_l^+ = g_l \quad (l = 1, 2, \dots, L) \quad (16)$$

$$\sum_{j=1}^n a_{ij} x_j \leq (=, \geq) b_i \quad (i = 1, 2, \dots, m) \quad (17)$$

$$x_j \geq 0 \quad (j=1, 2, \dots, n) \quad (18)$$

$$d_l^+, d_l^- \geq 0 \quad (l = 1, 2, \dots, L) \quad (19)$$

在以上各式中， ω_{kl}^+ 、 ω_{kl}^- 分别为赋予 P_k 优先因子的第 l 个目标的正、负偏差变量的权系数， g_l 为第 l 个目标的预期值， x_j 为决策变量， d_l^+ 、 d_l^- 分别为第 l 个目标的正、负偏差变量，(15) 式为目标函数，(16) 式为目标约束，(17) 式为绝对约束，(18) 式和 (19) 式为非负约束， $c_j^{(l)}$ 、 a_{ij} 、 b_i 分别为目标约束、绝对约束中决策变量的系数及约束值。其中， $i=1, 2, \dots, m$ ； $j=1, 2, \dots, n$ ； $k=1, 2, \dots, K$ ； $l=1, 2, \dots, L$ 。

下面我们来介绍目标规划的求解方法

二、求解目标规划的单纯形方法

目标规划的数学模型结构，与线性规划的数学模型结构没有本质的区别，所以可用单纯形法求解目标规划问题。但考虑到目标规划模型的特点，在用单纯形法求解目标规划时，应作以下规定：

(1) 因为目标规划问题的目标函数都是求最小值，所以 $c_j - Z_j \geq 0$

($j=1, 2, \dots, n$) 为最优判别准则。

(2) 因为非基变量的检验数中含有不同等级的优先因子，即：

$C_j - Z_j = \sum_{k=1}^K \alpha_{kj} P_k$ ($j=1, 2, \dots, n$)，而且 $P_1 \gg P_2 \gg \dots \gg P_K$ ，所以从每个检验数的整体来看，检验数的正、负首先决定于 P_1 的系数 α_{1j} 的正、负，若 $\alpha_{1j}=0$ ，则检验数的正、负就决定于 P_2 的系数 α_{2j} 的正、负，下面可依此类推。

求解目标规划问题的单纯形法计算步骤如下：

(1) 建立初始单纯形表，在表中将检验数行按优先因子个数分别排成 K 行，置 $k=1$ 。

(2) 检查该行中是否存在负数，且对应的前 $k-1$ 行的系数是零。若有，取其中最小者对应的变量为换入变量，转 (3)。若无负数，则转 (5)。

(3) 按最小比值规则 (θ 规则) 确定换出变量，当存在两个和两个以上相同的最小比值时，选取具有较高优先级别的变量为换出变量。

(4) 按单纯形法进行基变换运算，建立新的计算表，返回 (2)。

(5) 当 $k=K$ 时，计算结束，表中的解即为满意解。否则置 $k=k+1$ ，返回 (2)。

例 2：试用单纯形法求解上节例 1 所描述的目标规划问题 (9) - (14) 式。

解：首先将这一问题化为如下的标准形式：

$$\begin{aligned} \min Z &= P_1 d_1^+ + P_2 (d_2^- + d_2^+) + P_3 d_3^- \\ 2x_1 + x_2 + x_3 &= 11 \\ x_1 - x_2 + d_1^- - d_1^+ &= 0 \\ x_1 + 2x_2 + d_2^- - d_2^+ &= 10 \\ 8x_1 + 10x_2 + d_3^- - d_3^+ &= 56 \\ x_i, d_i^-, d_i^+ &\geq 0 \quad (i=1, 2, 3) \end{aligned}$$

(1) 取 x_3, d_1^-, d_2^-, d_3^- 为初始基变量，列出初始单纯形表，见表 4-1。

C_j			O	O	O	O	P_1	P_2	P_2	P_3	O	θ
C_B	X_B	b	x_1	x_2	x_3	d_1^-	d_1^+	d_2^-	d_2^+	d_3^-	d_3^+	
0	x_3	11	2	1	1	0	0	0	0	0	0	$\frac{10}{2}$
0	d_1^-	0	1	-1	0	1	-1	0	0	0	0	
P_2	d_2^-	10	1	[2]	0	0	0	0	-1	0	0	
P_3	d_3^-	56	8	10	0	0	0	0	0	1	-1	
$E_j - Z_j$	P_1		0	0	0	0	1	0	0	0	0	
	P_2		-1	-2	0	0	0	0	2	0	0	
	P_3		-8	-10	0	0	0	0	0	0	1	

(2) 取 $k=1$ ，检查检验数的 P_1 行，因该行无负检验数，故转 (5)。

(5) 因为 $k=1 < K=3$ ，置 $k=k+1=2$ ，返回 (2)。

(2) 检查发现检验数 P_2 行中有 -1, -2，因为有 $\min \{-1, -2\} = -2$ ，所以 x_2 为换入变量，转入 (3)。

(3) 在表 4-1 中，按 θ 规则计算最小比值：

$$\theta = \min \left\{ \frac{11}{1}, \frac{10}{5}, \frac{56}{10} \right\} = \frac{10}{2}$$

所以 d_2^- 为换出变量，转入 (4)。

(4) 进行换基运算，得表 4-2。以此类推，直至得到最终单纯形表为止，如表 4-3 所示。由表 4-3 可知， $x_1^* = 2, x_2^* = 4$ 为满意解。检查表 4-3 中的检验数行，发现非基变量 d_3^+ 的检验数为 0，这表明该问题存在多重解。在表 4-3 中，以非基变量 d_3^+ 为换入变量， d_1^- 为换出变量，经迭代得到表 4-4。从表 4-4 可以看出， $x_1^* = 10/3, x_2^* = 10/3$ 也是该问题的满意解。事实上，上述两组满意解的线性组合均是该问题的满意解。

表 4-2

c_j			O	O	O	O	P_1	P_2	P_2	P_3	O	θ
C_B	X_B	b	x_1	x_2	x_3	d_1^-	d_1^+	d_2^-	d_2^+	d_3^-	d_3^+	
0	x_3	6	3/2	0	1	0	0	-1/2	1/2	0	0	6/3
0	d_1^-	5	3/2	0	0	1	-1	1/2	-1/2	0	0	
0	x_2	5	1/2	1	0	0	0	1/2	-1/2	0	0	
P_3	d_3^-	6	[3]	0	0	0	0	-5	-5	1	-1	
c_j-Z_j	P_1		0	0	0	0	1	0	0	0	0	
	P_2		0	0	0	0	0	1	1	0	0	
	P_3		-3	0	0	0	0	5	-5	0	1	

表 4-3

c_j			O	O	O	O	P_1	P_2	P_2	P_3	O	θ
C_B	X_B	b	x_1	x_2	x_3	d_1^-	d_1^+	d_2^-	d_2^+	d_3^-	d_3^+	
0	x_3	3	0	0	1	0	0	2	-2	-1/2	1/2	
0	d_1^-	2	0	0	0	1	-1	3	-3	-1/2	1/2	
0	x_2	4	0	1	0	0	0	4/3	-4/3	-1/6	1/6	
0	x_1	2	1	0	0	0	0	-5/3	5/3	1/3	-1/3	
c_j-Z_j	P_1		0	0	0	0	1	0	0	0	0	
	P_2		0	0	0	0	0	1	1	0	0	
	P_3		0	0	0	0	0	0	0	1	0	

表 4-4

c_j			O	O	O	O	P_1	P_2	P_2	P_3	O	θ
C_B	X_B	b	x_1	x_2	x_3	d_1^-	d_1^+	d_2^-	d_2^+	d_3^-	d_3^+	
0	x_3	1	0	0	1	-1	1	-1	-1	0	0	
0	d_3^+	4	0	0	0	2	-2	6	-6	-1	1	
0	x_2	10/3	0	1	0	-1/3	1/3	1/3	-1/3	0	0	
0	x_1	10/3	1	0	0	2/3	-2/3	1/3	-1/3	0	0	
c_j-Z_j	P_1		0	0	0	0	1	0	0	0	0	
	P_2		0	0	0	0	0	1	1	0	0	
	P_3		0	0	0	0	0	0	0	1	0	

第三节 多目标规划应用实例

作为多目标规划方法在地理学研究中的应用实例，本节拟主要介绍绿洲型城市优化选址模型和工业结构优化模型。

一、绿洲型城市优化选址模型

笔者曾与罗格平同志合作，从研究新疆冲积扇型绿洲城市形成的环境地质基础入手，分析了影响绿洲型城市选址的环境地质要素，运用多目标线性规划方法建立了绿洲型城市优化选址模型，对绿洲型城市优化选址问题在理论上作了探讨，并以新疆奎屯绿洲为例，对模型进行了验证评价。下面，我们将这一模型作一些简单地介绍。

(一) 绿洲型城市环境地质基础及影响城市选址的环境地质要素

新疆绿洲型城市，按其所处的地理位置可分为以下几种类型：冲积扇型、洪积扇型、冲洪积平原型、河谷型、冲积平原型、湖岸平原型等六个类型，其中冲积扇型城市分布最为普遍，约占 39%，是新疆绿洲型城市的典型代表，冲积扇型城市的优化选址问题，是我们主要的研究对象。

冲积扇型城市在新疆绿洲型城市中最为多见，在天山南、北麓和昆仑山北麓，冲积扇发育较广，并且许多扇形地都具有相当大的规模。新疆冲积扇的发育与晚第三纪以来的中国西部环境的演变有着密切的关系。晚第三纪以来，昆仑山、天山、阿尔泰山剧烈上升，使来自海洋的湿润气流受到阻隔，逐渐形成了新疆干旱气候。在干旱气候环境下，植被逐渐变得稀疏，内陆湖泊变小乃至消亡，自然景观以荒漠为主。

由于山体抬升，在干旱环境下，风化产生大量的碎屑物质，这为冲积扇的形成提供了丰富的物质来源。

冲积扇的物质组成，扇顶部分为粗粒相堆积带，堆积厚度较大，主要由卵砾石组成，孔隙大、透水性强；冲积扇的中部为粗细粒过渡带；下部和边缘为细粒相堆积带，主要由砂土、亚粘土、粘土组成。冲积扇的扇面坡度自扇顶向扇缘具有明显的倾斜，在扇顶部坡角较大，而到了扇缘带倾角很小，只有 $1-2^{\circ}$ 。

由于城市是区域经济发展和地方行政管理中心。因此，优越的地理位置、便利的交通、充足的水源、潜在的经济环境及优良的自然环境等条件应是良好城市区位的基本特征。对于冲积扇上的绿洲型城市而言，其区位选择与环境地质因素有着极为密切的关系。因此，环境地质因素是影响绿洲型城市选址的主要因素，主要有如下几方面。

(1) 土地条件在干旱绿洲地区，宜农性土地是很有限的，也是很宝贵的土地资源，一般情况是不允许作为其它用地的。通常，土质差的非宜农性土地的城市建设征用费低，土质好的宜农性土地征用费高。因此，绿洲城市的区位应选在土质能满足城市建设要求，而土地征用费又较低廉的部位。显然，土地条件是限制绿洲城市选址的重要因子。

(2) 水文地质条件

① 地下水资源量：水资源是影响城市区位的重要因素是城市选址首选因素之一。对于绿洲型城市，地下水是城市主要的供水源，地下水位及其富水性直接关系到其是否可作为城市选址，因为地下水的开发利用潜力（如地下水天然补给量、允许开采量等）直接关系到城市的发展规模和方向（在本模型中，地下水的富水性，我们用单位涌水量来表示）。因此，我们认为地下水资

源是影响绿洲型城市优化选址重要的环境地质要素。

②地下水水质：地下水对绿洲型城市优化选址的影响，除了地下水量外，还有地下水水质问题。水质的好坏，直接影响到城市居民的身体健康和工业用水及其生产产品的质量、成本等。作为城市最优区位的地下水水质的有关毒理学指标如氟化物、氰化物、砷、汞、酚、铬等及其它表征水质状况的指标，如硬度、 SO_4^{+2} 、 Cl^- 、 NO_2^- 、 NO_3^- 、Cu、Pb、pv胺基、化学耗氧量、氨等经过简单的净化处理后，均应符合国家生活饮用水卫生标准和工业用水水质标准。

(3)工程地质条件：现代化城市，以高层建筑和高层建筑群为其外观特征。因此，城市选址必须考虑工程地质条件，对于冲积扇上的绿洲型城市而言，由于扇体的各个部位的物质组成不尽相同，因而其地基承载力也不尽相同，扇体各个部位的稳定性也有差异。因此，绿洲型城市的选址。应充分考虑现代高层建筑对地基承载力的要求，其次地下水位不应过高(埋藏小于 3m)否则会对建筑基础不利。

另外，地下水位也不应过低(埋深大于 100m)。否则，用水费用(主要包括打井费用、配套设备费用和电费等)过高，这样也将提高城市建设费用。

(4)环境地质灾害：地震、滑坡、泥石流、洪水等环境地质灾害在地处干旱内陆的新疆时有发生，特别是在春夏季节，泥石流、洪水灾害较为多见。因此，绿洲型城市的选址还应考虑环境地质灾害因素。城市选址应尽量避免环境地质灾害的多发地段，特别应避开受泥石流、洪水可能影响的地段。

(二)绿洲型城市优化选址模型

冲积扇型城市是新疆绿洲型城市的典型代表，本节旨在回答在冲积扇的什么部位建立城市最为合理或已建城镇合理发展规划问题。因此，冲积扇扇形地就是城市选址所允许的空间范围。也是我们优化选址模型的求解区域。通过前面关于绿洲型城市形成的环境地质基础的分析，在数学的几何形态上，可以将这个优化选址模型的求解区域近似地看成是一个规则的扇形体 Ω 。对扇面上的任意一点 M，可用柱面坐标表示，即用点 M 在平面上的投影的极坐标 (γ, θ) 以及点 M 到平面的距离 z (立标)描述。位于扇面上的不同点 M (γ, θ, z) ，其环境地质条件是有差异的，任何一个环境地质要素均可看成是坐标点即 γ, θ, z 的函数。

1. 绿洲型城市优化选址模型

(1)决策变量的确定及含义：

对于建在冲积扇上的绿洲型城市，其优化选址是用上述坐标中 M (γ, θ, z) 点表示。因此选定的决策变量只能是：

γ ：极径，表示扇面点 M 到出山口的平面距离；

θ ：极角，表示线 OM 在平面的投影与 OM_0 的夹角 (OM_0 为极径的始边)；

z：立标，表示点 M 到基准平面的距离深度，此平面可认为是扇缘部位等高线所在平面。

上述各变量所表示的柱面坐标如图 4-2 所示。

(2)目标函数：建立优化选址模型主要的一步是确定模型应追求的目标，从前面分析的影响绿洲型城市优化选址的因子看，我们认为城市建设选址应是在尽量少占耕地或不占耕地(即土地征用费低)的前提下，选择地下水资源丰富且取水便利、水质优良、工程地质条件好的区位，另外，最优区位还应是自然灾害危及区域外。因此，绿洲型城市优化选址所要求的目标具有

多重性。经研究分析，我们拟定“城市土地征用费”和“城市用水费用”最低作为本模型追求的目标：即追求“城市土地征用费最小”和“城市用水费用最低廉”。

①土地征用费：一般来说，城市建设征地的费用由土地的农业适宜性决定，即越适宜于农业利用的土地，其征用费就越高，反之，就越低。位于冲积扇上不同部位的土地宜农性差异是显著的，在扇面的上部（ γ 较小和 z 较大的部位）。由于地表物质组成较粗，地下水位深，土地宜农性较差，因而征用费也较低；在扇面的中部，土地的宜农性仍然较差，因而征用费较低；在扇面的中下部（ γ 较大， z 较小的部位），土质良好，地下水位亦较浅，土地的宜农性好，土地征用费亦高；在扇缘部位（ γ 最大， z 接近 0），土壤质量高，地下水埋藏适宜。只要采取有效的排水措施，能够避免土地沼泽化和盐碱化，因此农业适应性也较高（奎屯垦区几十年的农业开发证明了这一点）。因而土地征用费也相对较高，但由于排水较困难，且工程地质条件差，不宜建设城市。如果记 p 为单位土地面积的征用费，则 p 应是点 $M(\gamma, \theta, z)$ 的函数。那么，对于追求“征用费最少”这一目标，其目标函数可表示为：

$$\min p = p(\gamma, \theta, z) \quad (1)$$

②城市用水费用：这里城市用水费用主要指的是打井费用和配套设备费用及抽水费用。在冲积扇的下部，地下水位浅，用水费用低。在冲积扇的中上部，地下水位深，用水费用高。记 W 为单位城市土地面积上的用水费用，则所追求的“用水费用最低”这一目标函数可表示为：

$$\min W = W(\gamma, \theta, z) \quad (2)$$

(3)约束条件

①水量约束：前已分析了水资源是影响城市区位的主要因素。绿洲型城市应选在地下水丰富的区域，采用“地下水单位涌水量”来表征地下水的富水性，地下水单位涌水量大的区域，一般来说，含水层较厚，富水性好。设 $Q_{\min}$ 为满足城市用水最小的地下水单位涌水量（在奎屯绿洲，经讨论 $Q_{\min} = 800 \text{ m}^3/\text{m} \cdot \text{d}$ 较为合适）。

用 Q 表示地下水单位涌水量，则其受约为：

$$Q = Q(\gamma, \theta, z) \geq Q_{\min} \quad (3)$$

②水质约束

a. 酸碱性要求 pH 满足：

$$6.5 \leq \text{pH}(\gamma, \theta, z) \leq 8.5 \quad (4)$$

b. 五毒元素。选取酚、氰、砷、汞、铬等五项主要有毒污染物（离子）作为评价因子，采用 N.L. 内梅罗指数法表示，其约束形式为：

$$\rho = \rho(\gamma, \theta, z) \leq 0.5 \quad (5)$$

上式中， ρ 为内梅罗指数，选用生活饮用水卫生标准，若 $\rho \leq 0.5$ ，则认为各毒物均未超标。

c. 其它影响水质的因子。除了五毒元素外，还有其它影响水质的因子（或指标），主要有

矿化度、硬度、 SO_4^{+2} 、 Cl^- 、 NO_2^- 、 NO_3^- 、Cu、Pb、胺基、化学耗氧量和胺氧量和胺等十一种，其评价方法采用如下水质指数公式：

$$PI_w = \sum_{i=1}^n \frac{c_i}{S_i} \quad (6)$$

(6)式中, PI_w 表示地下水水质指数; c_i 表示某种污染物的检出含量 (mg/l); s_i 表示某种污染物的评价标准 (mg/l)。

地下水水质指数的约束条件为:

$$PI_w = PI_w(r, \theta, z) = \sum_{i=1}^n \frac{c_i(r, \theta, z)}{s_i} \leq 3 \quad (7)$$

上式 $PI_w \leq 3$, 表示地下水污染程度较轻, 一般可以作为生活饮用水, 处理简单、经济、水质完全符合国家颁布的生活饮用水标准。

③工程地质条件约束

a. 地下水位约束。地下水位埋深小于 3m, 对城市建筑施工不利; 大于 100m 则导致城市取水困难, 因此对地下水位埋深 H 要求:

$$3 \text{ (m)} \leq H(\gamma, \theta, z) \leq 100 \text{ (m)} \quad (8)$$

b. 地基承载力约束。对于不同的楼层建筑, 要求的地基承载力条件不同, 一般设 F_{min} 为城市建筑施工所要求的最低地基承载力, 则地基承载力 F 应满足:

$$F(\gamma, \theta, z) \geq F_{min} \quad (9)$$

2. 模型的分析评价可以看出, 以上的优化选址模型仅仅是借助于数学语言, 对绿洲型城市的优化选址问题作了一般性的理论描述。模型中的目标函数以及所有约束条件中所涉及的环境地质要素均是坐标点 $M(\gamma, \theta, z)$ 的函数。要将上述描述性的模型变成可利用数学方法和计算机求解的计算模型, 则需要知道每一个函数的具体形式。也就是说, 要将上述一般性的理论模型应用于某一个具体的冲积扇空间, 使其成为可以求解的实用模型, 则需要通过分析扇面上各个不同部位的环境地质资料, 以确定具体的模型形式。

下面我们以奎屯绿洲为例, 评价此模型的科学性。

在奎屯绿洲我们收集了较为详实的有关建立优化选址模型所需的具体资料, 分析这些资料可知, 在奎屯冲积扇上, 上述目标函数和约束条件的变化均有线性或近于线性的变化趋势。利用这些资料, 采用多元线性回归的方法, 将上述目标函数及约束条件的抽象数学表达式转化为具体的线性函数表达式, 其线性表达式分别通过了 $\alpha=0.01$ 或 0.05 的信度水平检验 (见表 4-5)。

在表 4-5 中, 无水质约束、无五毒元素约束及地基承载力约束, 这是因为这方面资料不足, 1988、1989 年在奎屯市区 (奎屯绿洲的下部) 取水样化验, 结果表明, 五毒元素均未超标, 有的样品中还未检出, 因此我们认为在奎屯绿洲, 五毒元素这一约束条件不会影响城市优化选址。另外, 地基承载力调查在冲积扇的中下部能够满足城市建设需要。因此这一约束条件对城市选址也可不考虑。

我们将采用目标加权法, 求解上述建立的多目标线性规划模型。这种方法的核心是把多目标模型转化为单目标模型, 适合于用传统的单纯形法在计算机上求解。这种方法的关键在于找到目标函数合理的加权系数。

表 4-5 表征目标函数及约束条件的函数表达式

函数名称		多元线性回归函数表达式	复相关系数	F 值	显著水平
目标	P	$P = -9307674 + 225878r - 106402\theta - 4476.2z$	0.9844	28.89	**
函数	W	$W = 504805 - 7800r + 599\theta - 151.5z$	0.9999	5999	**
约束	Q	$Q = 3006.6 - 25.6r - 17.68\theta - 0.97z$	0.9511	6.33	*
条件	PIW	$PIW = 0.46 + 0.0148r + 0.0010\theta - 0.0013z$	0.9796	15.91	**
	PH	$PH = 63.7 + 0.256r - 0.103\theta + 0.001z$	0.9996	456.57	**
	H	$H = 1009.6 - 15.6r + 1.98\theta - 0.30z$	0.9986	59.7	**

*通过 $\alpha = 0.05$ 信度水平的 F 检验；

**通过 $\alpha = 0.01$ 信度水平的 F 检验。

经讨论，两目标的权重分别取 0.6 和 0.4 较为合理。这样就将多目标转化为单目标，其表达式为

$$Y = 0.6P + 0.4W = -5383882.4 + 222758r - 106006\theta - 4536.6z$$

约束条件应满足：

$$Q \geq 800$$

(在奎屯绿洲， $Q \geq 800 \text{ m}^3/\text{m} \cdot \text{d}$ 的地段，被认为能够满足城市正常的供水)

$$\text{即 } 2566r + 17.68\theta - 0.97z \leq 2406.6$$

$$PIW \leq 2.18$$

($PI_w < 2.18$ 是依据 $PI_w < 3$ 换算出的，因为在奎屯绿洲只有矿化度等

8项指标的数据，缺三项，这样即有 $\frac{8}{11} \times 3 = 2.18$)

$$\text{即 } 0.0148r + 0.0010\theta - 0.0013z \leq 1.72$$

$$3 \leq H \leq 100$$

$$\text{即 } 15.6r - 1.98\theta + 0.30z \leq 1006.6$$

$$15.6r - 1.98\theta + 0.30z \geq 909.6$$

$$65 \leq pH \leq 85 \text{ (将 pH 值扩大 10 倍)}$$

$$\text{即 } 0.256r - 0.103\theta + 0.001z \geq 1.3$$

$$0.256r - 0.103\theta + 0.001z \leq 21.3$$

线性规划求解结果为

$$Y = 4756442.22$$

$$r = 57.7$$

$$\theta = 20.4$$

$$z = 121.3$$

求解结果指出城市最优区位在 (57.7, 20.4, 121.3)，位于冲积扇的中部略偏下，大致在奎屯兵站以北 2km，奎屯市应以此点为中心向四周发展 (最好是向东西方向延伸)。由此看来，奎屯市以前的区位是不合理的，应向南 (即

冲积扇的中部)发展。奎屯市政府作出的未来城市向南伸展的决策符合上述优化选址模型的计算结论。

二、工业结构优化模型

工业结构，是指工业经济系统内各部门或行业之间的生产联系和比例关系。工业结构与区域经济的发展有着极其密切的关系。一方面，工业结构随着区域经济的发展而变化，它是区域经济发展的表现和产物，反映着区域经济发展的方向、速度和水平；另一方面，工业结构又对区域经济的发展有着极为深刻的影响。合理的工业结构对区域经济的发展会产生巨大的推动力，不合理的工业结构则是区域经济发展的严重障碍。因此，工业结构的优化对于提高经济效益，促进区域经济持续、稳定、协调发展具有十分重要的意义。

工业结构优化的目的，是要使工业生产系统形成这样的结构，即在这样的结构下，能源、资源、资金、劳动力等生产要素得到合理的分配，并能以最少的消耗、良好的环境条件获得最大的经济效益。

对于工业结构优化问题，其优化的目标，不仅要考虑产值指标，同时还要考虑利润、资金投入、能源与原材料消耗等指标，而且这些目标的重要程度（即优先顺序）也不相同。因此，工业结构优化问题，是一个多因素、多目标优化问题。解决这一问题，传统的单目标规划（如线性规划）方法已经不能奏效了，因为这些方法只适用于解决单目标优化问题。鉴于这种情况，祝俊明同志曾运用目标规划方法建立了上海市工业结构优化模型，并且以此模型为依据，对上海市 2000 年的工业结构方案作了模拟分析^①。以下，我们将这一研究成果作一些简单地介绍。

(一)模型概况

1.决策变量：本模型共包括 32 个决策变量，它们分别代表 32 个工业行业的产值，其对应关系如表 4-6 所示。

2.约束方程：本模型共有 77 个约束方程，其中包括 11 个目标约束、32 个平衡约束、27 个上限约束和 7 个下限约束。

(1)目标约束 目标约束方程共 11 个。包括工业总产值、工业净产值、出口产品产值、利润税金、资金总额、固定资产总额、电能消耗量、煤碳消耗量、燃油消耗量、钢材消耗量和木材消耗量。方程如下：

①工业总产值

$$\sum_{j=1}^{32} a_{1j}x_j + d_1^- - d_1^+ = b_1$$
$$a_{1j} = 1 \quad j = 1, 2, \dots, 32$$

②工业净产值

$$\sum_{j=1}^{32} a_{2j}x_j + d_2^- = d_2^+ = b_2$$

表 4-6 模型中的决策变量决策变量

^① 祝俊明，上海市工业结构优化模型的设计与应用，全国人口与环境学术讨论会，兰州，1993 年月。

决 策 变 量	对 应 的 工 业 行 业	单 位
x ₁	建筑材料及其它非金属矿采选业	亿 元
x ₂	自来水生产和供应业	亿 元
x ₃	食品制造业	亿 元
x ₄	饮料制造业	亿 元
x ₅	烟草加工业	亿 元
x ₆	饲料工业	亿 元
x ₇	纺织业	亿 元
x ₈	缝纫业	亿 元
x ₉	皮革、毛皮及其制品业	亿 元
x ₁₀	木材加工及竹、藤、棕、草制品业	亿 元
x ₁₁	家具制造业	亿 元
x ₁₂	造纸及纸制品业	亿 元
x ₁₃	印刷业	亿 元
x ₁₄	文教体育用品及工艺美术制造业	亿 元
x ₁₅	电力、蒸汽、热水生产和供应业	亿 元
x ₁₆	石油加工业	亿 元
x ₁₇	炼焦、煤气及煤制品业	亿 元
x ₁₈	化学工业	亿 元
x ₁₉	医药工业	亿 元
x ₂₀	化学纤维工业	亿 元
x ₂₁	橡胶制品业	亿 元
x ₂₂	塑料制品业	亿 元
x ₂₃	建材工业	亿 元
x ₂₄	黑色金属冶炼及压延加工业	亿 元
x ₂₅	有色金属冶炼及压延加工业	亿 元
x ₂₆	金属制品业	亿 元
x ₂₇	机械工业	亿 元
x ₂₈	交通运输设备制造业	亿 元
x ₂₉	电器机械及器材制造业	亿 元
x ₃₀	电子及通信设备制造业	亿 元
x ₃₁	仪器仪表及其它计量工具制造业	亿 元
x ₃₂	其它工业	亿 元

$$a_{2j} = \frac{j_{\text{行业净产值}}}{j_{\text{行业总产值}}}$$

③出口产品产值

$$\sum_{j=1}^{32} a_{3j} x_j + d_3^- - d_3^+ = b_3$$

$$a_{3j} = \frac{j_{\text{行业出口产品产值}}}{j_{\text{行业总产值}}}$$

$$j = 1, 2, \dots, 32$$

④利润税金

$$\sum_{j=1}^{32} a_{4j} x_j + d_4^- - d_4^+ = b_4$$

$$a_{4j} = \frac{j_{\text{行业利润税金}}}{j_{\text{行业总产值}}}$$

$$j = 1, 2, \dots, 32$$

⑤资金总额

$$\sum_{j=1}^{32} a_{5j} x_j + d_5^- - d_5^+ = b_5$$

$$a_{5j} = \frac{j_{\text{行业资金总额}}}{j_{\text{行业总产值}}}$$

$$j = 1, 2, \dots, 32$$

⑥流动资金

$$\sum_{j=1}^{32} a_{6j} x_j + d_6^- - d_6^+ = b_6$$

$$a_{6j} = \frac{j_{\text{行业流动资金}}}{j_{\text{行业总产值}}}$$

$$j = 1, 2, \dots, 32$$

⑦电能消耗量

$$\sum_{j=1}^{32} a_{7j} x_j + d_7^- - d_7^+ = b_7$$

$$a_{7j} = \frac{j_{\text{行业耗电量}}}{j_{\text{行业总产值}}}$$

$$j = 1, 2, \dots, 32$$

⑧煤碳消耗量

$$\sum_{j=1}^{32} a_{8j} x_j + d_8^- - d_8^+ = b_8$$

$$a_{8j} = \frac{j_{\text{行业耗煤量}}}{j_{\text{行业总产值}}}$$

$$j = 1, 2, \dots, 32$$

⑨燃油消耗量

$$\sum_{j=1}^{32} a_{9j} x_j + d_9^- - d_9^+ = b_9$$

$$a_{9j} = \frac{j_{\text{行业耗油量}}}{j_{\text{行业总产值}}}$$

$$j = 1, 2, \dots, 32$$

⑩ 钢材消耗量

$$\sum_{j=1}^{32} a_{10j} x_j + d_{10}^- - d_{10}^+ = b_{10}$$

$$a_{10j} = \frac{j_{\text{行业耗钢材量}}}{j_{\text{行业总产值}}}$$

$$j = 1, 2, \dots, 32$$

(11) 木材消耗量

$$\sum_{j=1}^{32} a_{11j} x_j + d_{11}^- - d_{11}^+ = b_{11}$$

$$a_{11j} = \frac{j_{\text{行业耗木材量}}}{j_{\text{行业总产值}}}$$

$$j = 1, 2, \dots, 32$$

诸 α_{ij} ($i=1, 2, \dots, 11; j=1, 2, \dots, 32$)。可根据上海市 1989 年工业统计资料计算得到。 b_i ($i=1, 2, \dots, 11$) 以 1987 年和 1989 年为基础, 根据上海市国民经济发展规划的速度得到 2000 年的目标值。按照规划, 工业总产值、工业净产值、出口产品产值、利润税金、资金总额和流动资金的年增长率分别为 6%、5.7%、10.4%、4%、6.5% 和 6.5%, 而对于能源、原材料的消耗增长率分别为电力 6.3%、煤碳 4%、燃油 3%、钢材 6%、木材 3%。

(2) 平衡约束平衡约束是考虑各工业行业之间存在着客观的内在联系, 各行业的发展必须保持合理的比例关系, 才能使整个系统持续稳定和协调发展。

根据上海市 1987 年投入产出表把 86 个工业生产部门归并为 32 个产业 (与表 4-6 对应), 将建筑业等物质生产部门和非物质生产部门如交通运输业等, 以及农业部门消耗都归到最终使用中去, 同时考虑调入和调出, 进口和出口, 按一定发展速度把各行业最终使用折合到 2000 年, 这样就得到了 32 个平衡方程:

$$\sum_{j=1}^{32} c_{ij} x_j = B_j \quad i = 1, 2, \dots, 32$$

$$\text{其中, } c_{ij} = \delta_{ij} - d_{ij} \quad \delta_{ij} = \begin{cases} 0 & i \neq j \\ 1 & i = j \end{cases}$$

d_{ij} 为投入产出直接消耗系数, 表示 j 部门生产单位产品所直接消耗 i 部门产品的产值。 B_j 为 j 部门的最终需求。

(3) 上下限约束根据上海市 1978—1989 各行业的总产值计算出各行业的平均增长率, 若平均增长率小于 6%, 则按 2 倍增长率以 1989 年为基数计

算上限；若平均增长率大于 6%，则按照 12%的增长率计算上限。下限则是根据 1989 年基数，考虑到 2000 年的需求增长确定的。上下限约束的目的是保证工业行业的发展持续稳定，又能够满足需求。

3. 目标函数：本模型目标函数共分 4 个优先等级，所有的系统约束的偏差变量都放入第一优先级，而目标约束的偏差变量则根据目标的优先顺序而定。目标函数是偏差变量的线性函数，本模型中对所有偏差变量的权系数取值为 0 或 1。

(二)模拟方案选择与运算结果分析

1. 方案选择：经过反复调试、修改模型、调整参数之后，得到了一个基本模型，然后只改变目标优先次序或少数目标值，形成了七个不同方案。

方案一是把总产值、能源和原材料消耗限制、资金限制都放入第一优先级，工业净产值、出口产值、利税分别放入第二、三、四优先级。也就是要首先满足总产值达到 2607 亿元，工业电力消耗量不超过 445 亿千瓦时，工业煤碳消耗量不超过 3436 万吨，工业燃油消耗量不超过 649 万吨，工业木材消耗量不超过 141 万立方米，工业资金总额不超过 1679 亿元。

方案二是把出口产值的优先级从第三级升为第一级，即提高出口创汇的重要程度。

方案三是把利润税金的优先级提高升为第一级，并把能源和原材料的优先级降为第二级，即把利润税金放到首位考虑，并稍微放松对能源和原材料消耗的限制。

方案四则是把资金总额和固定资产的限制优先级提高，把出口产值和利润税金的优先级降低。

方案五则是把工业总产值和出口产值放入第一优先级，而资金总额、固定资产、能源和

表 4-7 各方案的各行业产值（2000 年）（单位：亿元）

方案 产值 变量	一	二	三	四	五	六	七
x ₁	0.00	0.00	0.09	0.00	0.00	0.00	0.00
x ₂	2.61	3.59	4.25	4.24	4.24	4.24	4.24
x ₃	69.34	68.91	69.33	68.69	68.91	68.89	69.39
x ₄	46.23	45.23	46.19	44.27	45.24	45.16	46.35
x ₅	37.40	36.76	37.35	37.07	36.76	36.56	37.41
x ₆	31.09	30.98	31.11	30.99	30.98	31.35	31.13
x ₇	209.67	209.67	209.67	277.99	209.67	209.67	209.67
x ₈	128.48	185.64	128.45	152.91	185.64	185.64	128.48
x ₉	59.26	59.26	59.26	59.26	59.26	59.26	59.26
x ₁₀	0.00	0.00	0.00	0.00	0.00	0.00	0.00
x ₁₁	1.51	2.59	1.47	1.56	2.60	2.97	1.45
x ₁₂	84.54	72.36	82.43	71.95	72.37	72.63	84.59
x ₁₃	67.06	61.60	66.16	61.60	61.60	61.60	67.10
x ₁₄	108.21	103.78	108.09	110.26	103.68	101.52	108.24
x ₁₅	19.54	21.61	22.00	22.00	22.00	22.00	22.00
x ₁₆	40.95	40.62	40.97	40.81	40.62	40.42	40.98
x ₁₇	11.08	10.35	9.91	9.66	10.16	10.04	5.85
x ₁₈	79.06	78.24	79.12	78.12	78.24	77.91	79.15
x ₁₉	80.97	80.36	81.00	80.30	80.37	102.07	81.07
x ₂₀	41.77	41.45	41.77	42.04	41.45	41.35	41.81
x ₂₁	83.70	83.36	83.71	90.17	83.36	83.20	83.74
x ₂₂	71.15	70.41	71.18	71.82	70.41	70.41	71.24
x ₂₃	65.48	64.46	65.51	65.36	64.46	64.08	65.55
x ₂₄	150.00	150.00	150.00	150.00	150.00	150.00	150.00
x ₂₅	64.88	45.53	64.88	64.88	45.01	33.62	64.88
x ₂₆	137.79	134.36	137.81	63.10	134.29	132.39	137.87
x ₂₇	318.58	317.41	318.61	316.64	317.39	316.68	318.62
x ₂₈	159.70	158.66	159.73	157.72	158.64	157.99	159.75
x ₂₉	172.54	167.65	172.53	170.69	167.53	164.79	172.58
x ₃₀	156.69	155.35	156.69	155.13	155.32	154.56	156.73
x ₃₁	70.56	70.56	70.56	70.56	70.56	70.56	70.56
x ₃₂	37.42	36.53	37.45	37.48	36.52	35.96	37.48

原材料都放入第二优先级。即重视工业产值和出口创汇，并稍微放松对资金、能源和原材料消耗的限制。

方案六是在方案五的基础上把纺织业的出口系数从 0.3366 提高到 0.5，缝纫业的出口系数从 0.5732 提高到 0.6。

方案七则是把工业产值、利税放入第一优先级，并把电力和煤碳供应量稍微减少。即把电力供应从 445 亿千瓦时降到 430 亿千瓦时，煤碳从 3436 万吨减为 3000 万吨。2. 运算结果分析：各方案的运算结果见表 4-7 和表 4-8，与 1989 年各行业产值比较，每种方案各行业产值的年增长率见表 4-9。根据 2000 年的目标值和正、负偏差，得 11 个目标的实际值，见表 4-10。再根据 1989 年几个主要指标和 1987 年工业对能源和原材料的消耗量数据计算的各目标的平均增长率见表 4-11。

表 4-8 各方案目标值的正负偏差（2000 年）

方案 目标值 偏差变量	一	二	三	四	五	六
d1 ± 亿元	0.000	0.000	0.000	0.000	0.000	0.000
d2 ± 亿元	81.482	81.644	81.667	76.843	81.586	81.587
d3 ± 亿元	-164.405	-137.328	-164.602	-140.625	-137.423	-135.443
d4 ± 亿元	69.211	67.514	69.345	64.944	67.445	64.208
d5 ± 亿元	-167.484	-163.753	-158.435	-171.786	-160.241	-161.580
d6 ± 亿元	180.924	-179.804	-171.630	-178.153	-176.150	-176.097
d7 ± 亿度	0.000	0.000	6.961	1.294	1.958	0.614
d8 ± 万吨	0.000	0.000	0.000	0.000	0.000	0.000
d9 ± 万吨	-20.228	0.000	8.792	4.631	5.875	8.957
d10 ± 万吨	198.849	189.493	199.027	116.655	189.347	187.955
d11 ± 万米 ³	0.000	0.000	0.000	0.000	0.000	0.000

注：表中负数表示负偏差，正数表示正偏差。

结果表明，在方案一的目标约束下，优化得到的结果其各行业产值分布上基本合理，增长速度基本适当，整个系统发展比较平衡。在此工业结构下，工业总产值可达 2607 亿元，平均每年增长 6%，工业净产值 707 亿元，平均每年增长 6.02%，出口产品产值可达 382 亿元，平均每年增长 6.87%，工业完成利税可达 419.6 亿元，平均每年递增 5.72%，资金和原材料限制基本得到满足。

在方案二下，由于提高了出口产值的优先级，出口产值有所增加，可达 409.4 亿元，平均每年递增 7.54%，这主要是因为一些出口系数较大的行业如缝纫业的产值增加了。

方案三由于放松了对资金和原材料消耗的限制，提高了利润税金的优先级，对资金和原材料的消耗有所增加，利润可达 419.8 亿元，平均年递增 5.72%。增加的资金消耗主要用于自来水生产和供应业、化纤工业、建材工业、石油工业、机械工业和交通工具制造业等行业中去了。

方案四由于降低了利润税金和出口新产品产值的优先级，出口产品产值

比方案二有所下降，为 406 亿元，平均每年递增 7.46%。利润税金也比方案三有所下降，为 415 亿元，年均增长

表 4-9 各行业产值的年平均增长率（%）

方案 增长率 行业	一	二	三	四	五	六	
x ₁	-0.5	-0.5	2.29	-0.5	-0.5	-0.5	
x ₂	3.75	6.80	8.46	8.43	8.43	8.43	
x ₃	2.69	2.63	2.68	2.60	2.62	2.62	
x ₄	12.83	12.61	12.82	12.39	12.61	12.59	1
x ₅	4.93	4.77	4.92	4.84	4.76	4.71	
x ₆	9.12	9.09	9.13	9.09	9.09	9.21	
x ₇	0.60	0.60	0.60	3.22	0.60	0.60	
x ₈	9.44	13.16	9.43	11.18	13.16	13.16	1
x ₉	13.16	13.16	13.16	13.16	13.16	3.16	1
x ₁₀	-0.5	-0.5	-0.5	-0.5	-0.5	-0.5	.
x ₁₁	-10.5	-5.99	-10.7	-10.2	-5.96	-4.82	.
x ₁₂	11.48	9.91	11.22	9.86	9.92	9.95	1
x ₁₃	14.04	13.16	13.90	13.16	13.16	13.16	1
x ₁₄	9.93	9.52	9.92	10.12	9.51	9.30	
x ₁₅	12.91	3.85	4.02	4.02	4.02	4.02	
x ₁₆	3.06	2.98	3.06	3.03	2.98	2.94	
x ₁₇	2.76	2.13	1.72	1.49	1.96	1.85	.
x ₁₈	2.10	2.00	2.11	2.00	2.45	1.97	
x ₁₉	10.80	10.73	10.81	10.72	10.73	13.16	1
x ₂₀	-1.11	-1.18	-1.11	-1.05	-1.18	-1.20	.
x ₂₁	9.87	9.83	9.87	10.62	9.83	9.81	
x ₂₂	7.58	7.48	7.59	7.68	7.48	7.44	
x ₂₃	5.22	5.07	5.22	5.20	5.07	5.01	
x ₂₄	0.20	0.20	0.20	0.20	0.20	0.20	
x ₂₅	5.68	2.33	5.68	5.68	2.22	-0.45	
x ₂₆	7.01	6.77	7.01	-0.32	6.76	6.62	
x ₂₇	4.14	4.11	4.14	4.08	4.11	4.09	
x ₂₈	9.36	9.29	9.36	9.23	9.29	9.25	
x ₂₉	4.91	4.63	4.91	4.80	4.63	4.47	
x ₃₀	7.08	7.00	7.08	6.99	7.00	6.95	
x ₃₁	13.16	13.16	13.16	13.16	13.16	13.16	1
x ₃₂	7.19	6.95	7.19	7.20	6.95	6.80	

率为 5.62%。

方案五放松了资金总额和固定资产限制，把出口产值放到第一优先级。结果表明，出口产值有明显提高。可达 409.35 亿元，平均年增长率为 7.53 %。工业资金总额和固定资产分别达到 1519.5 亿元和 793.56 亿元，年平均增长率分别为 5.53%和 4.58%。增加的资金消耗方案主要用于自来水生产和供应业、缝纫业、家具制造等行业。

方案六把纺织业和缝纫业的出口系数提高，结果表明纺织业的产值并没有增加，缝纫业和医药工业产值有大幅度提高，出口产值有所增加，达到 411.3 亿元，年平均增长率为 7.58%。

方案七考虑到 2000 年环境污染状况，把电力和煤碳供应量减少。结果表明，电力供应量

表 4-10 11 个目标的实际值

方 案 目标值 项 目	一	二	三	四	五	六
总产值(亿元)	2607.275	2607.275	2607.275	2607.275	2607.275	2607.27
净产值(亿元)	707.626	707.788	707.811	702.987	707.730	707.722
出口产值 (亿元)	382.369	409.446	382.172	406.149	409.351	411.331
利税(亿元)	419.622	417.925	419.756	415.355	417.656	418.619
资金总额 (亿元)	1512.363	1516.904	1521.412	1508.061	1519.606	1518.26
固定资产 (亿元)	788.784	789.000	798.078	791.555	793.558	793.611
电力消耗 (亿度)	445.000	445.000	451.961	446.249	446.958	445.614
煤碳消耗 (万吨)	3436.000	3436.000	3436.000	3436.000	3436.000	3436.00
燃油消耗 (万吨)	619.772	649.000	657.792	653.631	654.875	657.957
钢材消耗 (万吨)	686.849	677.493	687.027	604.655	677.347	675.951
木材消耗 (亿米 ³)	141.000	141.000	141.000	141.000	141.000	141.00

表 4-11 11 个目标的年平均增长率 (%)

方案 目标值 项目	一	二	三	四	五	六	七
总产值	6.00	6.00	6.00	6.00	6.00	6.00	6.00
净产值	6.02	6.03	6.03	5.96	6.03	6.03	6.04
出口产品值	6.87	7.54	6.86	7.46	7.53	7.58	7.79
利税	5.72	5.68	5.72	5.62	5.68	5.70	5.74
资金总额	5.49	5.51	5.55	5.46	5.53	5.53	5.53
固定资产净值	4.52	4.53	4.63	4.55	4.58	4.58	4.61
电力消耗	6.34	6.34	6.47	6.37	6.38	6.36	6.45
煤碳消耗	4.00	4.00	4.00	4.00	4.00	4.00	2.92
燃油消耗	2.64	3.00	3.11	3.06	3.08	3.11	2.91
钢材消耗	8.82	8.71	8.83	7.76	8.71	8.69	8.83
木材消耗	2.97	2.97	2.97	2.97	2.97	2.97	2.97

减少的目标不能实现，而煤碳消耗量可有所降低，在保持工业总产值不变的情况下，一些耗煤较多的行业如炼焦、煤气及煤制品业、建材工业等的产值有所降低，而一些耗煤较少的行业如医药工业、电气机械及器材制造业、电子工业、化纤工业产值则有所增加。工业净产值、出口产品产值、利税和资金消耗都有所增加。说明在煤碳供应最短缺的情况下，应多发展轻纺工业。

第五章 随机型决策方法

随机型决策方法，是一种处理随机型决策问题的决策技术。这种方法是地理决策研究中必不可少的方法之一。本章拟结合有关实例，介绍和探讨随机型决策方法在地理决策研究中的应用问题。

第一节 随机型决策问题

一、决策的基本概念

决策，是人类活动的基本组成部分之一，任何工作都离不开决策。那么，究竟什么是决策？许多人常常认为，所谓决策就是“制定政策”或“确定方案”，这作为对决策概念的狭义理解，无疑是正确的。但是，决策作为管理科学的一个特定术语，其含义要广泛得多。一般来说，凡是根据预定的目标作出的任何行动决定，都可以称之为决策。

为了使读者更进一步加深对决策概念的理解，以下我们来介绍关于决策的几个名词。

1. 决策问题在实际生活和生产中，对于一个需要处理的事件，面临几种客观条件，又有几种方案可供选择，这就构成了一个决策问题。

2. 自然状态在决策问题中，决策者所面临的每一种客观条件就称之为一个自然状态，简称状态或条件，有时也称为状态变量。

2. 自然状态在决策问题中，决策者所面临的每一种客观条件就称之为一个自然状态，简称状态或条件，有时也称为状态变量。

3. 行动方案 在决策问题中，那些可供选择的方案就称之为行动方案，简称方案或策略，有时也称方案变量或决策变量。

4. 状态概率 状态概率，是指在决策问题中各种自然状态出现的可能性大小的概率。

5. 益损值 益损值，是指采取某种行动方案在不同的自然状态下所获得的报酬。

6. 最佳决策方案 最佳决策方案，就是依照某种决策准则，使决策目标取最优值（譬如，最大值或最小值）的那个（些）行动方案。

例 1 某地区种植业生产中，可供选择的农作物有水稻、小麦、大豆、燕麦等四种。该地区每年可能发生的天气类型共有五种，即极旱年、旱年、平年、湿润年、极湿年。根据多年来的气候记录资料和农业生产统计资料计算得到的各种天气类型发生的概率，以及在最佳投入-产出条件下，每一种天气类型所对应的各种农作物收益（元/亩）如表 5-1 所示。试问，该地区生产经营者应该怎样选择其农作物种植种类？

该例所描述的问题就是一个决策问题。在这一决策问题中，共有五个自然状态——天气类型，即“极旱年”、“旱年”、“平年”、“湿润年”、“极湿年”，其自然状态概率分别为 0.1, 0.2, 0.4, 0.2, 0.1；共有四个行动方案，即“水稻”、“小麦”、“大豆”或“燕麦”；在每一种自然状态下，各个行动方案的益损值就是在每一种天气类型下各种农作物的收益值。

表 5-1 不同的天气类型所发生的概率及农作物收益

天气类型		极旱年	旱年	平年	湿润年	极湿年
发生概率		0.1	0.2	0.4	0.2	0.1
农作物收益	水稻	100	126	180	200	220
	小麦	250	210	170	120	80
	大豆	120	170	230	170	110
	燕麦	118	130	170	190	210

二、随机型决策问题

根据人们对决策问题的自然状态的认识程度不同，我们可以把决策问题划分为两种基本类型，即确定型决策问题和随机型决策问题。

所谓确定型决策问题，是指决策者已经完全确切地知道将发生什么样的自然状态，从而可以在既定的自然状态下选择最佳行动方案的一类决策问题。换句话说，对于确定型决策问题而言，只存在一个唯一确定的自然状态。

确定型决策问题看起来似乎很简单，但在实际工作中，决策者所面临的方案数目可能是很大的，其最佳决策方案的选择往往需要采用各种规划方法（如线性规划、目标规划、动态规划等方法）才能实现。

所谓随机型决策问题，是指决策者所面临的各種自然状态是随机出现的一类决策问题。

一个随机型决策问题，必须具备以下几个条件：

- (1)存在着决策者希望达到的明确目标；
- (2)存在着不依决策者的主观意志为转移的两个以上的自然状态；
- (3)存在着两个以上的可供选择的行动方案；
- (4)不同行动方案在不同自然状态下的益损值可以计算出来；

随机型决策问题，又可以进一步分为风险型决策问题和非确定型决策问题。在风险型决策问题中，虽然未来自然状态的发生是随机的，但是每一种自然状态发生的概率是已知的或者可以预先估计的。显然，例 1 所描述的决策问题就是一个典型的風險型决策问题。

在非确定型决策问题中，不仅未来自然状态的发生是随机的，而且各种自然状态发生的概率也是未知的和无法预先估计的。假如，在例 1 所描述的决策问题中，如果各种天气类型发生的概率未知，而且也无法预先估计，则该决策问题就变成了非确定型决策问题。

第二节 风险型决策方法

许多地理决策问题，常常需要在自然、经济、技术、市场等各种因素共存的环境下作出决策。而在这些因素中，有许多是决策者所不能控制和完全了解的。所以，风险型决策问题在地理决策研究中占有十分重要的地位。

对于风险型决策问题，其常用的决策方法主要有最大可能法、期望值法、矩阵法、灵敏度分析法、效用分析法等。虽然各种方法有其一定的应用场合，但是有时几种方法可以同时应用于同一决策问题上，并且可能由于随机性或决策准则的不同而得到不同的结果。因此，在实际决策时，可以应用不同方法分别进行计算，然后进行综合分析，以便减少决策的风险性。

一、最大可能法

我们知道，在某些情况下，确定型决策要比风险型决策容易些。那么，在什么条件下才能把风险型决策问题转化成确定型决策问题呢？概率论告诉我们，一个事件的概率越大，其发生的可能性就越大。根据这一原理，我们可以在风险型决策问题中选择一个概率最大的自然状态，而把其它概率较小的自然状态忽略，然后通过比较各行动方案在该自然状态下的益损值进行决策，这就是最大可能决策法。由此可见，最大可能法实际上就是在“大概率事件可看成必然事件，小概率事件可看成不可能事件”这样的假设条件下，将风险型决策问题转化成确定型决策问题的决策方法。

应用最大可能法进行决策，必须满足这样的条件：在一组自然状态中，某一自然状态出现的概率比其它自然状态出现的概率大很多，而且它们的益损值差别不很大。

若满足上述条件，则应用最大可能法进行决策，效果较好；若不满足上述条件，则决策效果不理想，甚至还可能引起重大失误。

容易知道，例 1 所描述的问题满足应用最大可能法决策的条件。以下，我们应用最大可能法对这一问题进行决策。

由表 5-1 可知，“极旱年”、“旱年”、“平年”、“湿润年”、“极湿年”各自然状态出现的概率分别为 0.1, 0.2, 0.4, 0.2, 0.1，显然，“平年”状态出现的可能性最大，按照最大可能法，可将“平年”状态的发生看成是必然事件。而在“平年”状态下，各行动方案的收益分别是：水稻为 180 元/亩，小麦为 170 元/亩，大豆为 230 元/亩，燕麦为 170 元/亩，显然，大豆的收益最大。所以，该地区生产经营者应该选择的最佳农作物种类是大豆。

二、期望值决策法及其矩阵运算

(一)期望值决策法

在概率论中，一个离散型随机变量的数学期望值，为这个随机变量的各个取值与其相应的概率的乘积之和。它代表了这个随机变量在概率意义上的平均取值。

期望值决策法，是通过计算各行动方案的期望益损值，并以此期望值为依据，选择平均收益最大或者平均损失最小的行动方案作为最佳决策方案的一种风险型决策方法。这一方法的决策过程如下：

(1)把每一个行动方案看成是一个随机变量，而它在不同自然状态下的益损值就是该随机变量的取值；

(2)把每一个行动方案在不同的自然状态下的益损值与其对应的状态概率相乘，再相加，则得到了该行动方案在概率意义上的平均益损值；

(3)选择平均收益最大或平均损失最小的行动方案作为最佳决策方案。
 以下，我们应用期望值决策法对例 1 所描述的问题进行决策。(1)在表 5-1 中，将行动方案“水稻”、“小麦”、“大豆”、“燕麦”分别看作随机变量 A_1, A_2, A_3, A_4 ；将自然状态“极旱年”、“旱年”、“平年”、“湿润年”、“极湿年”分别记作 $\theta_1, \theta_2, \theta_3, \theta_4, \theta_5$ ；对于每一个随机变量 A_i ($i=1, 2, 3, 4$)，将其在各个自然状态 θ_j ($j=1, 2, 3, 4, 5$)下的收益值 α_{ij} 看作是随机变量的取值。

(2)计算各个行动方案的期望收益值：

$$E(A_1) = 100 \times 0.1 + 126 \times 0.2 + 180 \times 0.4 + 200 \times 0.2 + 220 \times 0.1 \\ = 169.2 \text{ (元/亩)}$$

$$E(A_2) = 250 \times 0.1 + 210 \times 0.2 + 170 \times 0.4 + 120 \times 0.2 + 80 \times 0.1 \\ = 167 \text{ (元/亩)}$$

$$E(A_3) = 120 \times 0.1 + 170 \times 0.2 + 230 \times 0.4 + 170 \times 0.2 + 110 \times 0.1 \\ = 183 \text{ (元/亩)}$$

$$E(A_4) = 118 \times 0.1 + 130 \times 0.2 + 170 \times 0.4 + 190 \times 0.2 + 210 \times 0.1 \\ = 164.8 \text{ (元/亩)}$$

(3)选择最佳决策方案。

$$\because E(A_3) = \max E(A_i) = 183 \text{ (元/亩)}$$

故种植大豆为最佳决策方案。

综合上述决策过程，可以得出如下的决策表（见表 5-2）。

表 5-2 风险型决策问题的决策表

自然状态 状态概率 收益值 行动方案	极旱年(θ_1)	旱年(θ_2)	平年(θ_3)	湿润年(θ_4)	极湿年(θ_5)	期望收益值 $E(A_i)$
	0.1	0.2	0.4	0.2	0.1	
水稻(A_1)	100	126	180	200	220	170
小麦(A_2)	250	210	170	120	80	167
大豆(A_3)	120	170	230	170	110	183
燕麦(A_4)	118	130	170	190	210	165

利用期望值决策方法进行决策，虽然具有一定的准确程度，但它并不能防止在个别偶然情况下会出现较大的偏差，它掩盖了偶然情况下的损失值，所以在这一点上是具有一定风险的。然而，从统计的角度来看，以期望值作为选择最佳决策方案的依据还是比较合理的。

(二)期望值决策法的矩阵运算

在期望值决策法中，各个行动方案的期望益损值的计算，可以采用矩阵运算形式来进行。

假设某风险型决策问题，有 m 个行动方案 $A_1, A_2, \dots, A_m$ ；有 n 个自然状态 $\theta_1, \theta_2, \dots, \theta_n$ ，其状态概率分别为 $P_1, P_2, \dots, P_n$ 。如果在自然状态 θ_j 下，采取行动方案 A_i 的益损值为 α_{ij} ($i=1, 2, \dots, m; j=1, 2, \dots, n$)，则行动方案 A_i 的期望益损值为

$$E(A_i) = \sum_{j=1}^n a_{ij} P_j \quad (i = 1, 2, \dots, m) \quad (1)$$

如果引入下述向量

$$A = \begin{pmatrix} A_1 \\ A_2 \\ M \\ A_m \end{pmatrix}, \quad E(A) = \begin{pmatrix} E(A_1) \\ E(A_2) \\ M \\ E(A_m) \end{pmatrix}, \quad P = \begin{pmatrix} P_1 \\ P_2 \\ M \\ P_n \end{pmatrix},$$

及矩阵

$$B = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \Lambda & \Lambda & \Lambda & \Lambda & \Lambda & \Lambda & \Lambda & \Lambda \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

则公式 (1) 可以改写成如下矩阵运算形式：

$$E(A) = BP \quad (2)$$

在上例中，显然

$$A = \begin{pmatrix} A_1 \\ A_2 \\ A_3 \\ A_4 \end{pmatrix} \quad P = \begin{pmatrix} 0.1 \\ 0.2 \\ 0.4 \\ 0.2 \\ 0.1 \end{pmatrix} \quad B = \begin{pmatrix} 100 & 126 & 180 & 200 & 220 \\ 250 & 210 & 170 & 120 & 80 \\ 120 & 170 & 230 & 170 & 110 \\ 118 & 130 & 170 & 190 & 210 \end{pmatrix}$$

运用矩阵运算法则，经乘积运算可得

$$E(A) = \begin{pmatrix} E(A_1) \\ E(A_2) \\ E(A_3) \\ E(A_4) \end{pmatrix} = BP$$

$$= \begin{pmatrix} 100 & 126 & 180 & 200 & 220 \\ 250 & 210 & 170 & 120 & 80 \\ 120 & 170 & 230 & 170 & 110 \\ 118 & 130 & 170 & 190 & 210 \end{pmatrix} \begin{pmatrix} 0.1 \\ 0.2 \\ 0.4 \\ 0.2 \\ 0.1 \end{pmatrix} = \begin{pmatrix} 169.2 \\ 167 \\ 183 \\ 164.8 \end{pmatrix}$$

由于

$$E(A_3) = \max_i E(A_i) = 183 \text{ (元 / 亩)}$$

故该地区生产经营者应该选择种植大豆这一行动方案作为最佳决策方案。

三、树型决策法

树型决策法，是风险型决策问题研究中经常采取的决策方法。

决策树，是树型决策法的基本结构模型，它由决策点、方案分枝、状态结点、概率分枝和结果点等要素构成（见图 5-1）。

在图 5-1 中，小方框代表决策点；由决策点引出的每一条分支线段都代

表一个可能的决策方案，称为方案分枝；方案分枝末端的圆圈叫做状态结点；由状态结点引出的分枝代表可能的自然状态，叫做概率分枝；在概率分枝末端的小三角代表结果点。

在树型决策法中，决策的依据是各个方案在不同自然状态下的期望值，其决策的原则一般是，选择在各自然状态下的期望值最大（或最小）的方案作为最佳决策方案。进一步来说，如果决策目标是收益，则取期望值最大的方案作为最佳决策方案；如果决策目标是代价、成本或损失，则取期望值最小的方案作为最佳决策方案。

在运用树型决策方法进行风险型决策分析时，其逻辑顺序是从树根到树杆，再到树枝，最后向树梢逐渐展开；而各个方案在不同自然状态下的期望值的计算过程恰好与分析问题的逻辑顺序相反，它一般是从每一个树梢开始，经树枝、树杆，逐渐向树根进行。

用树型决策方法分析风险型决策问题的一般步骤如下：

(1) 画出决策树。即把某个决策问题未来发展情况的可能性和可能结果，由决策点逐级展开为方案分枝、状态结点和概率分枝等。

(2) 计算期望值。在决策树中，由树梢开始，依次计算各个方案在不同自然状态下的期望值，作为决策方案选择的依据。

(3) 剪枝。将各个方案在不同自然状态下的期望值分别标注在其对应的状态结点上，进行比较优选，将优胜者填入决策点，用“||”号剪掉舍弃方案，保留被选取的决策方案。

(一) 单级风险型决策

所谓单级风险型决策，就是指在整个决策过程中，决策者只需要作出一次决策方案选择的一类风险型决策。

例 2：某企业为了生产一种新产品，共有三个可行方案可供决策者选择：一是改造原有生产线；二是从国外引进生产线；三是与国内有关企业协作进行生产。市场需求状况大致有高、中、低三种可能，据估计，其发生的概率分别为 0.3，0.5，0.2。在不同的市场需求状况下，各种方案的效益值见表 5-3。试问该企业管理者究竟应该选择哪种决策方案。

该问题是一个典型的单级风险型决策问题，以下我们用树型决策方法对这一问题进行求解。

表 5-3 某企业在不同方案下生产某种新产品的效益 (万元)

效益值 需求状况 (概率) 行动方案	高需求 $\theta_1(0.3)$	中等需求 $\theta_2(0.5)$	低需求 $\theta_3(0.2)$
改造生产线 A_1	200	100	20
引进生产线 A_2	220	120	60
协作生产 A_3	180	100	80

(1) 首先画出该问题的决策树图形 (如图 5-2 所示)。

(2) 计算各个行动方案的期望效益值。

a. 状态结点 V_1 的期望效益值为：

$$E_V = 200 \times 0.3 + 100 \times 0.5 + 20 \times 0.2 = 114 \text{ (万元)}$$

b. 状态结点 V_2 的期望效益值为：

$$EV_2 = 220 \times 0.3 + 120 \times 0.5 + 60 \times 0.2 = 138 \text{ (万元)}$$

c. 状态结点 V_3 的期望效益值为：

$$EV_3 = 180 \times 0.3 + 100 \times 0.5 + 80 \times 0.2 = 120 \text{ (万元)}$$

(3) 剪枝。因为 $EV_2 > EV_1$, $EV_2 > EV_3$, 故剪掉状态结点 V_1 和 V_3 所对应的方案分枝, 保留状态结点 V_2 所对应的方案分枝。即该问题的最佳决策方案应该是：从国外引进生产线。

(二) 多级风险型决策

所谓多级风险型决策, 就是指在整个决策过程中, 决策者需要作出多次决策方案选择的一种风险型决策。

例 3：某企业, 由于生产工艺较落后, 产品成本高, 在价格保持中等水平的情况下无利可图, 在价格低落时就要亏损, 只有在价格较高时才能盈利。鉴于这种情况, 企业管理者有意改进其生产工艺, 用新的工艺代替原来的旧工艺。现在取得新工艺有两种途径：一是自行研制, 但其成功的可能性概率是 0.6；二是购买专利, 估计谈判成功的可能性概率是 0.8。如果研制成功或谈判成功, 生产规模都将考虑两种方案：一是产量不变；二是增加产量。如果研制或谈判都失败, 则仍采用原工艺进行生产, 并保持原生产规模不变。据市场预测, 该企业产品今后跌价的可能性概率是 0.1, 保持中等水平的可能性概率是 0.5, 涨价的可能性概率是 0.4。各个方案在不同价格状态下的效益值列于表 5-4。面临这种情况, 该企业管理者应如何做出决策？

表 5-4 某企业在不同生产方案下的效益值 (万元)

效益值 价格状态 (概率)	行 动 方 案	按原工艺生产	改进工艺成功			
			购买专利成功 (0.8)		自行研制成功 (0.6)	
			产量不变	增加产量	产量不变	增加产量
价格低落 (0.1)		-100	-200	-300	-200	-300
价格中等 (0.5)		0	50	50	0	-250
价格上涨 (0.4)		100	150	250	200	600

此问题是一个典型的多级风险型决策问题, 下面我们用树型决策方法对这一问题进行分析。

(1) 首先画出决策树, 如图 5-3 所示。

(2) 计算期望效益值, 并进行剪枝：

a. 状态结点 V_7 的期望效益值为 $EV_7 = (-200) \times 0.1 + 50 \times 0.5 + 150 \times 0.4 = 65$ (万元)；状态结点 V_3 的期望效益值为 $EV_3 = (-300) \times 0.1 + 50 \times 0.5 + 250 \times 0.4 = 95$ (万元)。由于 $EV_8 > EV_7$, 故剪掉状态结点 V_7 对应的方案分枝, 并将 EV_8 的数值填入决策点 V_4 , 即令 $EV_4 = EV_8 = 95$ (万元)。

b. 状态结点 V_3 的期望效益值为 $EV_3 = (-100) \times 0.1 + 0 \times 0.5 + 100 \times 0.4 = 30$ (万元), 故状态结点 V_1 的期望效益值为 $EV_1 = 30 \times 0.2 + 95 \times 0.8 = 82$ (万元)。

c. 状态结点 V_9 的期望效益值为 $EV_9 = (-200) \times 0.1 + 0 \times 0.5 + 200 \times 0.4 = 60$ (万元)；状态结点 V_{10} 的期望效益值为 $EV_{10} = (-300) \times 0.1 + (-250) \times 0.5 + 600 \times 0.4 = 85$ (万元)。由于 $EV_{10} > EV_9$, 故剪掉状态结点 V_9 对应的方案分枝, 将 EV_{10} 的数值填入决策点 V_5 , 即令 $EV_5 = EV_{10} = 85$ (万元)。

d. 状态结点 V_6 的期望效益值为： $EV_6 = (-100) \times 0.1 + 0 \times 0.5 + 100 \times 0.4 = 30$ (万元)，故状态结点 V_2 的期望效益值为： $EV_2 = 30 \times 0.4 + 85 \times 0.6 = 63$ (万元)。

e. 由于 $EV_1 > EV_2$ ，故剪掉状态结点 V_2 对应的方案分枝，将 EV_1 的数值填入决策点 V ，即令 $EV = EV_1 = 82$ (万元)。

综合上述期望效益值计算与剪枝过程可知，该问题的决策方案应是：首先采用购买专利方案进行工艺改造，当购买专利改造工艺成功后，再采用扩大生产规模（即增加产量）方案。

四、灵敏度分析法

对于风险型决策问题，其各个行动方案的期望益损值是在对状态概率预测的基础上求得的。由于状态概率的预测会受到许多不可控因素的影响，因而基于状态概率预测结果的期望益损值也不可能同实际完全一致，会产生一定的误差。这样，就必须对可能产生的数据变动是否会影响最佳决策方案的选择进行分析，这就是灵敏度分析。

例 4：某企业拟扩大其产品产量，现有两种行动方案可供选择：一是新建生产线；一是改造原生产线。该企业决策者经过研究分析，编制的决策表如表 5-5 所示。

由于市场情况极其复杂，它受许多不可控因素的影响，因而其销售状态的概率可能有所变动，所以必须对最佳决策方案进行灵敏度分析。

(1) 应用最大期望值准则确定最佳决策方案。由表 5-5 可知， $E(A_1) = \max\{E(A_1), E(A_2)\} = 290$ 万元，故新建生产线 (A_1) 方案为最佳方案。

表 5-5 某企业扩大产品产量决策表

益损值 行动方案	自然状态 状态概率	市场销售状态		期望益损值(万元) $E(A_i)$
		适销(θ_1)	滞销(θ_2)	
		0.7	0.3	
新建生产线(A_1)		500	-200	290
改造原生产线(A_2)		300	-100	180

(2) 灵敏度分析。当考虑市场销售状态中适销的概率由 0.7 变为 0.3 时，则两个行动方案的期望益损值的变化为： $E(A_1) = 10$ 万元， $E(A_2) = 20$ 万元。

由以上计算可知，当适销状态的概率从 0.7 变到 0.3 时，最佳决策方案由 A_1 变为 A_2 ，即由新建生产线变为改造原生产线。不难想象，在 0.7 与 0.3 之间一定存在一点 P ，当适销状态的概率等于 P 时，新建生产线方案与改造原生产线方案的期望益损值相等。这个概率 P 称为转移概率，它可以通过下式求得：

$$500P + (1-P)(-200) = 300P + (1-P)(-100)$$

解之得 $P = 0.33$ 。所以，当 $P > 0.33$ 时，新建生产线 (A_1) 为最佳决策方案；当 $P < 0.33$ 时，改造原生产线方案 (A_2) 为最佳决策方案。

五、效用分析法

以上所讨论的风险型决策问题，是以期望益损值作为决策准则的，这有时不一定符合决策的实际情况。因为决策是决策者自己作出的，决策者个人

的主观因素不能不对决策过程产生影响。如果完全采用期望损益值作为决策的准则，就会把决策过程变成机械地计算期望值的过程，而排除了决策者的作用，这当然是不合理的。

假定某决策问题有甲、乙两种行动方案可供决策者选择，其中，甲方案有 0.5 的概率可得 200 元，有 0.5 的概率损失 100 元；而乙方案将以 1.0 的概率稳得 25 元。试问决策者愿意接受哪一个方案？

对于这一问题，尽管甲方案的期望收益值 50 元大于乙方案的期望收益值 25 元，但对于大多数决策者来说，宁愿选择乙方案，这是因为他们不愿意承担遭受 100 元损失的风险。当然也会有决策者宁愿冒损失 100 元的风险而选择甲方案。

这一事例说明，面对同一决策问题，不同的决策者对相同的利益和损失的反应不同。即使对于相同的决策者，在不同的时期和情况下，这种反应也不相同。这就是决策过程中决策者的主观价值概念，即效用值概念。

效用理论应用于决策过程，其主要步骤可概括如下：

(1)画出效用曲线。以损益值为横坐标，以效用值为纵坐标。规定最大损益值的效用值为 1，最小损益值的效用值为 0，其余数值可以采用向决策者逐一提问的方式确定。不同的效用曲线反映着决策者的不同类型（见图 5-4）。在图 5-4 中，曲线 A 是保守型决策者的效用曲线，它反映的是一种不求大利，避免风险，谨慎小心的决策者；曲线 C 是风险型决策者的效用曲线，它反映的是一种谋求大利，不惧风险的进取型决策者；曲线 B 是中间型决策者的效用曲线。

(2)按效用值进行决策。当效用曲线画出后，就可以按效用期望值进行决策。其作法是：先在效用曲线上找出每一个行动方案，在不同自然状态下的损益值的效用值；然后再计算各个行动方案的效用期望值；最后选择效用期望值最大的行动方案作为最佳决策方案。

第三节 非确定型决策方法

对于非确定型决策问题，不仅其未来自然状态的发生是随机的，而且各自然状态发生的概率也是未知的和无法事先确定的。这类问题的决策，主要取决于决策者的素质、经验和决策风格等，没有一个完全固定的模式可循，对于同一个决策问题，不同的决策者可能会采用不同的处理方法。以下，我们介绍几种比较常用的处理非确定型决策问题的决策方法，以供决策者或决策分析者参考。

一、乐观法

乐观法，又叫最大最大准则法，其决策的原则是“大中取大”。其特点是决策者持最乐观的态度，决策时不放弃任何一个获得最好结果的机会，愿意以承担一定风险的代价去获得最大的利益。

假定某非确定型决策问题有 m 个行动方案 $A_1, A_2, \dots, A_m$ ；有 n 个自然状态 $\theta_1, \theta_2, \dots, \theta_n$ 。如果行动方案 A_i ($i=1, 2, \dots, m$) 在自然状态 θ_j ($j=1, 2, \dots, n$) 下的效益值为 $V(A_i, \theta_j)$ ，则乐观法的决策步骤如下：

- (1) 求出每一个行动方案在各个自然状态下的最大效益值 $\max_j \{V(A_i, \theta_j)\}$ ；
- (2) 求出各个行动方案的最大效益值的最大值 $\max_i \max_j \{V(A_i, \theta_j)\}$ ；
- (3) 选择最佳决策方案。如果

$$V(A_{i^*}, \theta_{j^*}) = \max_i \max_j \{V(A_i, \theta_j)\}$$

则： A_{i^*} 为最佳决策方案。

例 5：对于例 1 所描述的风险型决策问题，假如各自然状态发生的概率未知且无法预先估计，则这一问题就变成了表 5-6 所描述的非确定型决策问题。试问：应怎样选择最佳决策方案？

表 5-6 非确定型决策问题

收益 (元/亩) 自然状态 行动方案	天 气 类 型				
	极旱年(θ_1)	旱年(θ_2)	平年(θ_3)	湿润年(θ_4)	极湿年(θ_5)
种水稻(A_1)	100	126	180	200	220
种小麦(A_2)	250	210	170	120	80
种大豆(A_3)	120	170	230	170	110
种燕麦(A_4)	118	134	170	190	210

下面，我们用乐观法对该问题进行决策。

- (1) 求每个行动方案在各个自然状态下的最大效益值：

$$\max_j \{V(A_1, \theta_j)\} = \max\{100, 126, 180, 200, 220\}$$

$$= 220 \text{ (元/亩)} = V(A_1, \theta_5)$$

$$\max_j \{V(A_2, \theta_j)\} = \max\{250, 210, 170, 120, 80\}$$

$$= 250 \text{ (元/亩)} = V(A_2, \theta_1)$$

$$\begin{aligned}\max_j \{V(A_3, \theta_j)\} &= \max \{120, 170, 230, 170, 110\} \\ &= 230 \text{ (元/亩)} = V(A_3, \theta_3) \\ \max_j \{V(A_4, \theta_j)\} &= \max \{118, 130, 170, 190, 210\} \\ &= 219 \text{ (元/亩)} = V(A_4, \theta_5)\end{aligned}$$

(2)求各个行动方案的最大效益值的最大值：

$$\begin{aligned}\max_i \max_j &= \max \{220, 250, 230, 210\} \\ &= 250 \text{ (元/亩)} = V(A_2, \theta_1)\end{aligned}$$

(3)选择最佳决策方案。因为

$$\max_i \max_j \{V(A_i, \theta_j)\} = V(A_2, \theta_1)$$

故：种小麦 (A_2)为最佳决策方案。

二、悲观法

悲观法，又叫最大最小准则法或瓦尔德 (Wald) 准则法，其决策的原则是“小中取大”。这种方法的特点是决策者持最悲观的态度，他总是把事情估计得很不利。悲观法决策的步骤如下：

(1) 求出每一个行动方案在各个自然状态下的最低效益值 $\min_j \{V(A_i, \theta_j)\}$ ；

(2) 求出各个行动方案的最低效益值的最大值 $\max_i \min_j \{V(A_i, \theta_j)\}$ ；

(3)选择最佳决策方案。如果

$$V(A_{i^*}, \theta_{j^*}) = \max_i \min_j \{V(A_i, \theta_j)\}$$

则： A_{i^*} 为最佳决策方案。

下面，我们用悲观法对例 5 所描述的非确定型决策问题进行求解。

(1)求每一个行动方案在各个自然状态下的最低效益值：

$$\begin{aligned}\min_j \{V(A_1, \theta_j)\} &= \min \{100, 126, 180, 200, 220\} \\ &= 100 \text{ (元/亩)} = W(A_1, \theta_1) \\ \min_j \{V(A_2, \theta_j)\} &= \min \{250, 210, 170, 120, 80\} \\ &= 80 \text{ (元/亩)} = V(A_2, \theta_5) \\ \min_j \{V(A_3, \theta_j)\} &= \min \{120, 170, 230, 170, 110\} \\ &= 110 \text{ (元/亩)} = V(A_3, \theta_5) \\ \min_j \{V(A_4, \theta_j)\} &= \min \{118, 130, 170, 190, 210\} \\ &= 118 \text{ (元/亩)} = V(A_4, \theta_1)\end{aligned}$$

(2)求各个行动方案的最低效益值的最大值：

$$\begin{aligned}\max_i \min_j \{V(A_i, \theta_j)\} &= \max \{100, 80, 110, 118\} \\ &= 118 \text{ (元/亩)} = V(A_4, \theta_1)\end{aligned}$$

(3)选择最佳决策方案。因为：

$$\max_i \min_j \{V(A_i, \theta_j)\} = V(A_4, \theta_1)$$

故：种燕麦 (A_4) 为最佳决策方案。

三、折衷法

乐观法按照最好的可能性选择最佳决策方案，悲观法按照最坏的可能性选择最佳决策方案。这两种方法在决策过程中所损失的信息过多，而且决策结果也有很大的片面性。采用折衷法进行决策，在一定程度上可以克服这些缺点。

折衷法的特点是既不非常乐观，也不非常悲观，而是通过一个系数 α ($0 \leq \alpha \leq 1$) 来表示决策者对客观条件估计的乐观程度。折衷法决策的步骤如下：

- (1) 计算每一个行动方案在各个自然状态下的最大效益值 $\max_j \{V(A_i, \theta_j)\}$ ；
- (2) 计算每一个行动方案在各个自然状态下的最低效益值 $\min_j \{V(A_i, \theta_j)\}$ ；
- (3) 计算每一个行动方案的折衷效益值：

$$V_i = \alpha \max_j \{V(A_i, \theta_j)\} + (1 - \alpha) \min_j \{V(A_i, \theta_j)\}$$
- (4) 计算各个行动方案的折衷效益值的最大值 $\max_i \{V_i\}$ ；
- (5) 选择最佳决策方案。如果

$$V_{i^*} = \max_i \{V_i\}$$

则： A_{i^*} 即为最佳决策方案。

以下，我们用折衷法对例 5 所描述的非确定型决策问题进行求解。

- (1) 计算每一个行动方案在各个自然状态下的最大收益值：

$$\begin{aligned} \max_j \{V(A_1, \theta_j)\} &= \max \{100, 126, 180, 200, 220\} \\ &= 220 \text{ (元 / 亩)} = V(A_1, \theta_5) \end{aligned}$$

$$\begin{aligned} \max_j \{V(A_2, \theta_j)\} &= \max \{250, 210, 170, 120, 80\} \\ &= 250 \text{ (元 / 亩)} = V(A_2, \theta_1) \end{aligned}$$

$$\begin{aligned} \max_j \{V(A_3, \theta_j)\} &= \max \{120, 170, 230, 170, 110\} \\ &= 230 \text{ (元 / 亩)} = V(A_3, \theta_3) \end{aligned}$$

$$\begin{aligned} \max_j \{V(A_4, \theta_j)\} &= \max \{118, 130, 170, 190, 210\} \\ &= 210 \text{ (元 / 亩)} = V(A_4, \theta_5) \end{aligned}$$

- (2) 计算每一个行动方案在各个自然状态下的最低效益值：

$$\begin{aligned} \min_j \{V(A_1, \theta_j)\} &= \min \{100, 126, 180, 200, 220\} \\ &= 100 \text{ (元 / 亩)} = V(A_1, \theta_1) \end{aligned}$$

$$\begin{aligned} \min_j \{V(A_2, \theta_j)\} &= \min \{250, 210, 170, 120, 80\} \\ &= 80 \text{ (元 / 亩)} = V(A_2, \theta_5) \end{aligned}$$

$$\begin{aligned}\min_j \{V(A_3, \theta_j)\} &= \min \{120, 170, 230, 170, 110\} \\ &= 110 \text{ (元/亩)} = V(A_3, \theta_5) \\ \min_j \{V(A_4, \theta_j)\} &= \min \{118, 130, 170, 190, 210\} \\ &= 118 \text{ (元/亩)} = V(A_4, \theta_1)\end{aligned}$$

(3) 计算每一个行动方案的折衷效益值 (譬如取 $\alpha=0.5$):

$$\begin{aligned}V_1 &= \alpha V(A_1, \theta_5) + (1-\alpha)V(A_1, \theta_1) \\ &= 0.5 \times 220 + 0.5 \times 100 = 160 \text{ (元/亩)} \\ V_2 &= \alpha V(A_2, \theta_1) + (1-\alpha)V(A_2, \theta_5) \\ &= 0.5 \times 250 + 0.5 \times 80 = 165 \text{ (元/亩)} \\ V_3 &= \alpha V(A_3, \theta_3) + (1-\alpha)V(A_3, \theta_5) \\ &= 0.5 \times 230 + 0.5 \times 110 = 170 \text{ (元/亩)} \\ V_4 &= \alpha V(A_4, \theta_5) + (1-\alpha)V(A_4, \theta_1) \\ &= 0.5 \times 210 + 0.5 \times 118 = 164 \text{ (元/亩)}\end{aligned}$$

(4) 计算各个行动方案的折衷效益值的最大值:

$$\begin{aligned}\max_i \{V_i\} &= \max \{160, 165, 170, 164\} \\ &= 170 \text{ (元/亩)} = V_3\end{aligned}$$

(5) 选择最佳决策方案。因为

$$\max_i \{V_i\} = V_3$$

故种大豆 (A_3) 为最佳决策方案。

四、等可能性法

在非确定型决策问题中, 由于未来自然状态发生的概率未知, 故在决策过程中, 决策者或决策分析者就有理由认为每个自然状态发生的可能性概率是相等的。等可能性法就是基于这种假设的一种决策方法。

等可能性法求解非确定型决策问题的作法是, 首先假设各个自然状态发生的概率相等, 即 $P_1 = P_2 = \dots = P_n = \frac{1}{n}$, 然后在这一假设下计算各个行动方案的期望益损值, 通过比较各个行动方案的期望益损值, 选择最佳决策方案。

对于例 5 所描述的非确定型决策问题, 用等可能性法决策的过程如下:

(1) 首先假设“极旱年”、“旱年”、“平年”、“湿润年”、“极湿年”各天气类型发生的概率相等, 即 $P_1 = P_2 = P_3 = P_4 = P_5 = \frac{1}{5}$;

(2) 在上述假设下计算各个行动方案的期望效益值:

$$\begin{aligned}E(A_1) &= \frac{1}{5} \times 100 + \frac{1}{5} \times 126 + \frac{1}{5} \times 180 + \frac{1}{5} \times 200 + \frac{1}{5} \times 220 \\ &= 165.2 \text{ (元/亩)} \\ E(A_2) &= \frac{1}{5} \times 250 + \frac{1}{5} \times 210 + \frac{1}{5} \times 170 + \frac{1}{5} \times 120 + \frac{1}{5} \times 80 \\ &= 166 \text{ (元/亩)} \\ E(A_3) &= \frac{1}{5} \times 120 + \frac{1}{5} \times 170 + \frac{1}{5} \times 230 + \frac{1}{5} \times 170 + \frac{1}{5} \times 110 \\ &= 160 \text{ (元/亩)}\end{aligned}$$

$$E(A_4) = \frac{1}{5} \times 118 + \frac{1}{5} \times 130 + \frac{1}{5} \times 170 + \frac{1}{5} \times 190 + \frac{1}{5} \times 210$$

$$= 163.6 (\text{元} / \text{亩})$$

(3)选择最佳决策方案。由于

$$\max_i E(A_i) = 166 (\text{元} / \text{亩}) = E(A_2)$$

所以，应该选择种植小麦 (A_2)为最佳决策方案。

五、后悔值法

在实际决策过程中，当某一自然状态出现时，就能很容易地知道哪个行动方案的效益最大或损失最小。如果决策者当时在决策时没有采用这一方案而采用了其它方案，则事后他就会感到后悔，遗憾当时没有选准效益最大或损失最小的方案。对于非确定型决策问题，为了避免事后遗憾太大，可采用后悔值法进行决策。

后悔值是后悔值法决策的主要依据。所谓后悔值是指在其自然状态下的最大效益值与各方案的效益值之差。后悔值法也称最小最大后增值法。这种方法的决策步骤如下：

(1)求出每一个自然状态下各个行动方案的最大效益值：

$$\max_i \{V(A_i, \theta_j)\} = V(A_{i^*}, \theta_j)$$

(2)对于每一个自然状态下的各个行动方案，计算其后悔值：

$$V_{ij} = V(A_{i^*}, \theta_j) - V(A_i, \theta_j)$$

(3) 对于每一个行动方案，计算其最大后悔值 $\max_j \{V_{ij}\}$ ；

(4) 求出各个行动方案的最大后悔值的最小值 $\min_i \max_j \{V_{ij}\}$ ；

(5)选择最佳决策方案。如果

$$\min_i \max_j \{V_{ij}\} = V_{i^*j^*}$$

则 A_{i^*} 为最佳决策方案。

下面，我们用后悔值法对例 5 所描述的非确定型决策问题求解。

(1)计算每一个自然状态下各个行动方案的最大效益值：

$$\max_i \{V(A_i, \theta_1)\} = \max\{100, 250, 120, 118\}$$

$$= 250 (\text{元} / \text{亩}) = V(A_2, \theta_1)$$

$$\max_i \{V(A_i, \theta_2)\} = \max\{126, 210, 170, 130\}$$

$$= 210 (\text{元} / \text{亩}) = V(A_2, \theta_2)$$

$$\max_i \{V(A_i, \theta_3)\} = \max\{180, 170, 230, 170\}$$

$$= 230 (\text{元} / \text{亩}) = V(A_3, \theta_3)$$

$$\max_i \{V(A_i, \theta_4)\} = \max\{200, 120, 170, 190\}$$

$$= 200 (\text{元} / \text{亩}) = V(A_1, \theta_4)$$

$$\max_i \{V(A_i, \theta_5)\} = \max\{220, 80, 110, 210\}$$

$$= 220 (\text{元} / \text{亩}) = V(A_1, \theta_5)$$

(2)对于每一个自然状态下的各个行动方案，计算其后悔值： $V_{11}=250-$

$100=150$ (元/亩); $V_{21}=250-250=0$ (元/亩); $V_{31}=250-120=130$ (元/亩);
 $V_{41}=250-118=132$ (元/亩); $V_{12}=210-126=84$ (元/亩); $V_{22}=210-210=0$;
 $V_{32}=210-170=40$ (元/亩); $V_{42}=210-130=80$ (元/亩); $V_{23}=230-180=50$ (元/亩);
 $V_{23}=230-170=60$ (元/亩); $V_{33}=230-230=0$; $V_{43}=230-170=60$ (元/亩);
 $V_{14}=200-200=0$; $V_{24}=200-120=80$ (元/亩); $V_{34}=200-170=30$ (元/亩);
 $V_{44}=200-190=10$ (元/亩); $V_{15}=220-220=0$; $V_{25}=220-80=140$ (元/亩);
 $V_{35}=220-110=110$ (元/亩); $V_{45}=220-210=10$ (元/亩)

(3)对于每一个行动方案,求其最大后悔值:

$$\begin{aligned}
 \max_j \{V_{1j}\} &= \max \{150, 84, 50, 0, 0\} \\
 &= 150 \text{ (元/亩)} = V_{11} \\
 \max_j \{V_{2j}\} &= \max \{0, 0, 60, 80, 140\} \\
 &= 140 \text{ (元/亩)} = V_{25} \\
 \max_j \{V_{3j}\} &= \max \{130, 40, 0, 30, 110\} \\
 &= 130 \text{ (元/亩)} = V_{31} \\
 \max_j \{V_{4j}\} &= \max \{132, 80, 60, 10, 10\} \\
 &= 132 \text{ (元/亩)} = V_{41}
 \end{aligned}$$

(4)求各个行动方案的最大后悔值的最小值:

$$\begin{aligned}
 \min_i \max_j \{V_{ij}\} &= \min \{150, 140, 130, 132\} \\
 &= 130 \text{ (元/亩)} = V_{31}
 \end{aligned}$$

(5)选择最佳决策方案。由于

$$\min_i \max_j \{V_{ij}\} = V_{31}$$

以,应所该选择种植大豆 (A3)为最佳决策方案。

第六章 AHP 决策分析法

美国运筹学家 A.L.Saaty 于本世纪 70 年代提出的层次分析法 (Analytical Hierarchy Process, 简称 AHP 方法), 是一种定性与定量相结合的决策分析方法。它是一种将决策者对复杂系统的决策思维过程模型化、数量化的过程。应用这种方法, 决策者通过将复杂问题分解为若干层次和若干因素, 在各因素之间进行简单的比较和计算, 就可以得出不同方案的权重, 为最佳方案的选择提供依据。这种方法的特点是:

(1) 思路简单明瞭, 它将决策者的思维过程条理化、数量化, 便于计算, 容易被人们所接受。

(2) 所需要的定量数据较少, 但对问题的本质, 包含的因素及其内在关系分析得清楚。

(3) 可用于复杂的非结构化的问题, 以及多目标、多准则、多时段等各种类型问题的决策分析, 具有较广泛的实用性。

第一节 AHP 决策分析法的原理、步骤与计算方法

一、基本原理

层次分析法的基本原理可以用以下的简单事例分析来说明。假设有 n 个物体 $A_1, A_2, \dots, A_n$ ，它们的重量分别记为 $W_1, W_2, \dots, W_n$ 。现将每个物体的重量两两进行比较如下：

	A_1	A_2	$\vdots$	A_n
A_1	W_1 / W_1	W_1 / W_2	$\vdots$	W_1 / W_n
A_2	W_2 / W_1	W_2 / W_2	$\vdots$	W_2 / W_n
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
A_n	W_n / W_1	W_n / W_2	$\vdots$	W_n / W_n

若以矩阵来表示各物体的这种相互重量关系，即

$$\bullet \quad A = \begin{pmatrix} W_1 / W_1 & W_1 / W_2 & \cdots & W_1 / W_n \\ W_2 / W_1 & W_2 / W_2 & \cdots & W_2 / W_n \\ \vdots & \vdots & \ddots & \vdots \\ W_n / W_1 & W_n / W_2 & \cdots & W_n / W_n \end{pmatrix} \quad (1)$$

(1)式中， A 称为判断矩阵。若取重量向量 $W = [W_1, W_2, \dots, W_n]^T$ ，则有：

$$AW = n \cdot W \quad (2)$$

这就是说， W 是判断矩阵 A 的特征向量， n 是 A 的一个特征值。事实上，根据线性代数知识，我们不难证明， n 是矩阵 A 的唯一非零的，也是最大的特征值，而 W 为其所对应的特征向量。

上述事实提示我们，如果有一组物体，需要知道它们的重量，而又没有衡器，那么我们就可以通过两两比较它们的相互重量，得出每对物体重量比的判断，从而构成判断矩阵；然后通过求解判断矩阵的最大特征值 $\lambda_{\max}$ 和它所对应的特征向量，就可以得出这一组物体的相对重量。根据这一思路，在地理科学研究中，对于一些无法测量的因素，只要引入合理的标度，我们也可以用这种方法来度量各因素之间的相对重要性，从而为有关决策提供依据。上述思路就是层次分析法的基本原理。

二、基本步骤

层次分析方法的基本过程，大体可以分为如下六个基本步骤：

(1)明确问题。即弄清问题的范围，所包含的因素，各因素之间的关系等，以便尽量掌握充分的信息。

(2)建立层次结构。在这一个步骤中，要求将问题所含的因素进行分组，把每一组作为一个层次，按照最高层（目标层）、若干中间层（准则层）以及最低层（措施层）的形式排列起来。这种层次结构常用结构图来表示（见图 6-1），图中要标明上下层元素之间的关系。如果某一个元素与下一层的所有元素均有联系，则称这个元素与下一层次存在有完全层次的关系；如果某一个元素只与下一层的部分元素有联系，则称这个元素与下一层次存在有不完全层次关系。层次之间可以建立子层次，子层次从属于主层次中的某一个元素，它的元素与下一层的元素有联系，但不形成独立层次。

(3)构造判断矩阵。这一个步骤是层次分析法的一个关键步骤。判断矩阵表示针对上一层次中的某元素而言，评定该层次中各有关元素相对重要性的状况，其形式如下：

其中， b_{ij} 表示对于 A_k 而言，元素 B_i 对 B_j 的相对重要性的判断值。 b_{ij} 一般取 1, 3, 5, 7, 9 等 5 个等级标度，其意义为：1 表示 B_i 与 B_j 同等重要；3 表示 B_i 较 B_j 重要一点；5 表示 B_i 较 B_j 重要得多；7 表示 B_i 较 B_j 更重要；9 表示 B_i 较 B_j 极端重要。而 2, 4, 6, 8 表示相邻判断的中值，当 5 个等级不够用时，可以使用这几个数值。

显然，对于任何判断矩阵都应满足

$$\begin{cases} b_{ii} = 1 \\ b_{ij} = \frac{1}{b_{ji}} \quad (i, j = 1, 2, \dots, n) \end{cases} \quad (3)$$

因此，在构造判断矩阵时，只需写出上三角（或下三角）部分即可。

一般而言，判断矩阵的数值是根据数据资料、专家意见和分析者的认识，加以平衡后给出的。衡量判断矩阵质量的标准是矩阵中的判断是否具有 consistency。如果判断矩阵存在关系：

$$b_{ij} = \frac{b_{ik}}{b_{jk}} \quad (i, j, k = 1, 2, 3, \dots, n) \quad (4)$$

则称它具有完全一致性。但是，因客观事物的复杂性和人们认识上的多样性，可能会产生片面性，因此要求每一个判断矩阵都有完全的一致性显然是不可能的，特别是因素多、规模大的问题更是如此。为了考察层次分析法得到的结果是否基本合理，需要对判断矩阵进行一致性检验。

(4)层次单排序。层次单排序的目的是对于上层次中的某元素而言，确定本层次与之有联系的元素重要性次序的权重值。它是本层次所有元素对上一层次而言的重要性排序的基础。

层次单排序的任务可以归结为计算判断矩阵的特征根和特征向量问题，即对于判断矩阵 B ，计算满足：

$$BW = \lambda_{\max} W \quad (5)$$

的特征根和特征向量。(5)式中， $\lambda_{\max}$ 为 B 的最大特征根， W 为对应于 $\lambda_{\max}$ 的正规化特征向量， W 的分量 W_i 就是对应元素单排序的权重值。

通过前面的分析，我们知道，当判断矩阵 B 具有完全一致性时， $\lambda_{\max} = n$ 。但是，在一般情况下是不可能的。为了检验判断矩阵的一致性，需要计算它的一致性指标：

$$CI = \frac{\lambda_{\max} - n}{n - 1} \quad (6)$$

(6)式中，当 $CI=0$ 时，判断矩阵具有完全一致性；反之， CI 愈大，则判断矩阵的一致性就愈差。

为了检验判断矩阵是否具有令人满意的一致性，则需要将 CI 与平均随机一致性指标 RI （见表 6-1）进行比较。一般而言，1 或 2 阶判断矩阵总是具有完全一致性的。对于 2 阶以上的判断矩阵，其一致性指标 CI 与同阶的平均随

机一致性指标 RI 之比，称为判断矩阵的随机一致性比例，记为 CR。一般地，当

$$CR = \frac{CI}{RI} < 0.10 \quad (7)$$

时，我们就认为判断矩阵具有令人满意的一致性；否则，当 $CR \geq 0.1$ 时，就需要调整判断矩阵，直到满意为止。

表 6-1 平均随机一致性指标

阶数	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
RI	0	0	0.58	0.90	1.12	1.24	1.32	1.41	1.45	1.49	1.52	1.54	1.56	1.58	1.59

(5)层次总排序。利用同一层次中所有层次单排序的结果，就可以计算针对上一层次而言的本层次所有元素的重要性权重值，这就称为层次总排序。层次总排序需要从上到下逐层顺序进行。对于最高层，其层次单排序就是其总排序。

若上一层次所有元素 $A_1, A_2, \dots, A_m$ 的层次总排序已经完成，得到的权重值分别为 $a_1, a_2, \dots, a_m$ 与 a_j 对应的本层次元素 $B_1, B_2, \dots, B_n$ 的层次单排序结果为 $[b_1^j, b_2^j, \dots, b_n^j]^T$ (这里，当 B_i 与 A_j 无联系时, $b_i^j = 0$)，那么，得到的层次总排序见表 6-2。

表 6-2 层次总排序表

层次A \ 层次B	层次A				B层次的总排序
	A_1	A_2	$\dots$	A_m	
	a_1	a_2	$\dots$	a_m	
B_1	b_1^1	b_1^2	$\dots$	b_1^m	$\sum_{j=1}^m a_j b_1^j$
B_2	b_2^1	b_2^2	$\dots$	b_2^m	$\sum_{j=1}^m a_j b_2^j$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
B_n	b_n^1	b_n^2	$\dots$	b_n^m	$\sum_{j=1}^m a_j b_n^j$

显然，

$$\sum_{i=1}^n \sum_{j=1}^m a_j b_i^j = 1 \quad (8)$$

即层次总排序为归一化的正规向量。

(6)一致性检验。为了评价层次总排序的计算结果的一致性，类似于层次单排序，也需要进行一致性检验。为此，需要分别计算下列指标：

$$CI = \sum_{j=1}^m a_j CI_j \quad (9)$$

$$RI = \sum_{j=1}^m a_j RI_j \quad (10)$$

$$CR = \frac{CI}{RI} \quad (11)$$

在 (9)式中，CI 为层次总排序的一致性指标， CI_j 为与 a_j 对应的 B 层次

中判断矩阵的一致性指标；在 (10) 式中， RI 为层次总排序的随机一致性指标， RI_j 为与 a_j 对应的 B 层次中判断矩阵的随机一致性指标；在 (11) 式中， CR 为层次总排序的随机一致性比例。

同样，当 $CR < 0.10$ 时，则认为层次总排序的计算结果具有令人满意的一致性；否则，就需要对本层次的各判断矩阵进行调整，从而使层次总排序具有令人满意的一致性。

三、计算方法

通过前面的介绍，我们知道，在层次分析方法中，最根本的计算任务是求解判断矩阵的最大特征根及其所对应的特征向量。这些问题当然可以用线性代数知识去求解，并且能够利用计算机求得任意高精度的结果。但事实上，在层次分析法中，判断矩阵的最大特征根及其对应的特征向量的计算，并不需要追求太多的精度。这是因为判断矩阵本身就是将定性问题量化的结果，允许存在一定的误差范围。因此，我们常常用如下的两种近似算法求解判断矩阵的最大特征根及其所对应的特征向量。

(一) 方根法

这一方法的计算步骤如下：

(1) 计算判断矩阵每一行元素的乘积

$$M_i = \prod_{j=1}^n b_{ij} \quad (i = 1, 2, \dots, n) \quad (12)$$

(2) 计算 M_i 的 n 次方根

$$\bar{W}_i = \sqrt[n]{M_i} \quad (i = 1, 2, \dots, n) \quad (13)$$

(3) 将向量 $\bar{W} = [\bar{W}_1, \bar{W}_2, \dots, \bar{W}_n]^T$ 归一化

$$W_i = \bar{W}_i / \sum_{i=1}^n \bar{W}_i \quad (i = 1, 2, \dots, n) \quad (14)$$

则 $W = [W_1, W_2, \dots, W_n]^T$ 即为所求的特征向量。

(4) 计算最大特征根

$$\lambda_{\max} = \sum_{i=1}^n \frac{(AW)_i}{nW_i} \quad (15)$$

(12) 式中， $(AW)_i$ 表示向量 AW 的第 i 个分量。

(二) 和积法

这一方法的计算步骤如下：

(1) 将判断矩阵每一列归一化：

$$b_{ij} = b_{ij} / \sum_{k=1}^n b_{kj} \quad (i, j = 1, 2, \dots, n) \quad (16)$$

(2) 对按列归一化的判断矩阵，再按行求和：

$$\bar{W}_i = \sum_{j=1}^n b_{ij} \quad (i = 1, 2, \dots, n) \quad (17)$$

(3) 将向量 $\bar{W} = [\bar{W}_1, \bar{W}_2, \dots, \bar{W}_n]^T$ 归一化：

$$W_i = \frac{\bar{W}_i}{\sum_{j=1}^n \bar{W}_j} \quad (i = 1, 2, \dots, n) \quad (18)$$

则： $W=[W_1, W_2, \dots, W_n]^T$ 即为所求的特征向量。

(4)计算最大特征根：

$$\lambda_{\max} = \sum_{i=1}^n \frac{(AW)_i}{nW_i} \quad (19)$$

(19)式中， $(AW)_i$ 同样表示向量 AW 中的第 i 个分量。

第二节 AHP 决策分析实例

一、兰州市主导产业决策分析

地处甘肃省中部、黄河上游的兰州市，是甘肃省的省会，全省政治、经济、文化、医疗卫生、教育和科技中心。兰州经济的发展，无疑在全省、乃至全国占有着十分重要的地位。在改革开放深入发展的今天，如何抓住时机，发挥地区优势，促进兰州经济的全面发展，是摆在省、市各级领导面前的一项急待解决的重大决策问题。为了解决这一问题，必须以市场为导向，结合本市的自然、经济、社会和技术条件，综合各种有利和不利因素，选择一批能发挥地区优势，具有较高效益的主导产业，从而带动全市经济的腾飞。下面，我们用层次分析方法，对兰州市主导产业的选择问题作一些初步分析，以供决策者参考。（一）模型层次结构

1. 目标层 (A)：选择带动兰州市经济全面发展的主导产业。

2. 准则层 (C)包括如下三个方面：

(1) C_1 ：市场需求（包括市场需求现状和远景市场潜力）。

(2) C_2 ：效益准则（这里主要考虑产业的经济效益）。

(3) C_3 ：发挥地区优势，合理利用资源。

3. 对象层 (P)包括如下 14 个产业：

(1) P_1 ：能源工业

(2) P_2 ：交通运输业

(3) P_3 ：冶金工业

(4) P_4 ：化工工业

(5) P_5 ：纺织工业

(6) P_6 ：建材工业

(7) P_7 ：建筑业

(8) P_8 ：机械工业

(9) P_9 ：食品加工业

(10) P_{10} ：邮电通讯业

(11) P_{11} ：电器、电子工业

(12) P_{12} ：农业

(13) P_{13} ：旅游业

(14) P_{14} ：饮食服务

上述目标层、准则层及对象层中各元素所构成的层次结构关系，如图 6-2 所示。

（二）计算过程

1. 构造判断矩阵，进行层次单排序。根据上述模型结构，在专家咨询的基础上，我们构造了 A—C 判断矩阵、C—P 判断矩阵，并进行了层次单排序计算，其结果分别如下：

A—C 判断矩阵：

A	C ₁	C ₂	C ₃	W _A
C ₁	1	$\frac{1}{3}$	3	0.260
C ₂		1	5	0.634
C ₃			1	0.106

$$\lambda_{\max}=3.038 \quad CI=0.019$$

$$RI=0.58 \quad CR=0.0328 < 0.10$$

C₁—P 判断矩阵、C₂—P 判断矩阵、C₃—P 判断矩阵、分别见 122 页和 123 页。

2. 层总排序，一致性检验 根据以上层次单排序的结果，经过计算，可得对象层 (P) 的层次总排序 (见表 6-3)。

(三) 基本结论

综合上述计算过程，可以得出如下两点基本结论：

(1) 从 C 层的排序结果来看，兰州市主导产业选择的准则应该是，首先考虑产业的效益 (主要是经济效益)；其次考虑市场需求和远景市场潜力；第三考虑发挥地区优势和资源合理利用问题。

C ₁	P ₁	P ₂	P ₃	P ₄	P ₅	P ₆	P ₇	P ₈	P ₉	P ₁₀	P ₁₁	P ₁₂	P ₁₃	P ₁₄	W ₁
P ₁	1	3	5	2	4	7	6	8	8	8	8	9	9	9	0.235
P ₂		1	3	$\frac{1}{2}$	2	5	4	6	9	8	7	8	9	9	0.143
P ₃			1	$\frac{1}{4}$	$\frac{1}{2}$	3	2	4	9	6	5	7	8	9	0.084
P ₄				1	3	6	5	7	9	8	8	9	9	9	0.186
P ₅					1	4	3	5	9	7	6	8	8	9	0.110
P ₆						1	$\frac{1}{2}$	2	7	4	3	5	6	8	0.046
P ₇							1	3	8	5	4	6	7	9	0.063
P ₈								1	6	3	2	4	5	7	0.035
P ₉									1	$\frac{1}{4}$	$\frac{1}{5}$	$\frac{1}{3}$	$\frac{1}{2}$	2	0.009
P ₁₀										1	$\frac{1}{2}$	2	3	5	0.025
P ₁₁											1	3	4	6	0.026
P ₁₂												1	2	4	0.016
P ₁₃													1	3	0.012
P ₁₄														1	0.008

$$\lambda_{\max}=15.65 \quad CI=0.127 \quad RI=1.58 \quad CR=0.0804 < 0.10$$

C ₁	P ₁	P ₂	P ₃	P ₄	P ₅	P ₆	P ₇	P ₈	P ₉	P ₁₀	P ₁₁	P ₁₂	P ₁₃	P ₁₄	W ₂
P ₁	1	2	3	4	5	9	6	9	9	9	7	8	8	9	0.234
P ₂		1	2	3	4	8	5	9	9	9	6	7	8	9	0.183
P ₃			1	2	3	8	4	9	9	8	5	6	7	9	0.143
P ₄				1	2	7	3	9	8	8	4	5	6	9	0.109
P ₅					1	6	2	9	8	7	3	4	5	8	0.084
P ₆						1	$\frac{1}{5}$	5	3	2	$\frac{1}{4}$	$\frac{1}{3}$	$\frac{1}{2}$	$\frac{1}{4}$	0.017
P ₇							1	9	7	6	2	3	8	4	0.065
P ₈								1	$\frac{1}{3}$	$\frac{1}{4}$	$\frac{1}{8}$	$\frac{1}{7}$	$\frac{1}{6}$	$\frac{1}{2}$	0.008
P ₉									1	$\frac{1}{2}$	$\frac{1}{6}$	$\frac{1}{5}$	$\frac{1}{4}$	2	0.010
P ₁₀										1	$\frac{1}{5}$	$\frac{1}{4}$	$\frac{1}{3}$	3	0.015
P ₁₁											1	2	3	7	0.047
P ₁₂												1	2	8	0.028
P ₁₃													1	5	0.027
P ₁₄														1	0.015

$\lambda_{\max}=15.94$ CI=0.149RI=1.58CR=0.0943<0.10

C ₃	P ₁	P ₂	P ₃	P ₄	P ₅	P ₆	P ₇	P ₈	P ₉	P ₁₀	P ₁₁	P ₁₂	P ₁₃	P ₁₄	W ₃
P ₁	1	6	8	5	2	7	8	3	9	4	9	9	9	9	0.236
P ₂		1	4	$\frac{1}{2}$	$\frac{1}{5}$	2	3	$\frac{1}{4}$	8	$\frac{1}{3}$	5	7	6	9	0.064
P ₃			1	$\frac{1}{5}$	$\frac{1}{8}$	$\frac{1}{3}$	$\frac{1}{2}$	$\frac{1}{7}$	5	$\frac{1}{6}$	2	4	6	3	0.026
P ₄				1	$\frac{1}{4}$	3	4	$\frac{1}{3}$	9	$\frac{1}{2}$	6	8	9	7	0.084
P ₅					1	6	7	2	3	8	9	9	9	9	0.186
P ₆						1	2	$\frac{1}{5}$	7	$\frac{1}{4}$	4	5	7	5	0.045
P ₇							1	$\frac{1}{6}$	6	$\frac{1}{5}$	3	5	7	4	0.035
P ₈								1	9	2	8	9	9	8	0.143
P ₉									1	$\frac{1}{9}$	$\frac{1}{4}$	$\frac{1}{2}$	$\frac{1}{3}$	2	0.009
P ₁₀										1	7	9	9	8	0.110
P ₁₁											1	3	5	2	0.021
P ₁₂												1	3	$\frac{1}{2}$	0.011
P ₁₃													1	$\frac{1}{4}$	0.008
P ₁₄														1	0.015

$\lambda_{\max}=15.64$ CI=0.126

RI=1.58 CR=0.0797<0.10

表 6-3 对象层 (P) 的层次总排序

A	C ₁	C ₂	C ₃	W _总
	0.260	0.634	0.106	
P ₁	0.235	0.234	0.236	0.2345
P ₂	0.143	0.183	0.064	0.1600
P ₃	0.084	0.143	0.026	0.1153
P ₄	0.186	0.109	0.084	0.1271
P ₅	0.110	0.084	0.186	0.1013
P ₆	0.046	0.017	0.045	0.0281
P ₇	0.063	0.065	0.035	0.0613
P ₈	0.035	0.008	0.143	0.0296
P ₉	0.009	0.010	0.009	0.0105
P ₁₀	0.025	0.015	0.110	0.0271
P ₁₁	0.026	0.047	0.021	0.0398
P ₁₂	0.016	0.028	0.011	0.0284
P ₁₃	0.012	0.027	0.008	0.0215
P ₁₄	0.008	0.015	0.015	0.0135

(2)从 P 层总排序的结果来看,兰州市主导产业选择的优先顺序应该是:P₁ (能源工业)>P₂ (交通运输业)>P₄ (化工工业)>P₃ (冶金工业)>P₅ (纺织工业)>P₇ (建筑业)>P₁₁ (电器、电子工业)>P₈ (机械工业)>P₁₂ (农业)>P₆ (建材工业)>P₁₀ (邮电通讯业)>P₁₃ (旅游业)>P₁₄ (饮食服务业)>P₉ (食品加工业)。

二、晋陕蒙三角地区综合开发治理战略决策分析

晋陕蒙三角地区包括山西省的河曲、保德、偏关和兴县,陕西省的神木、府谷和榆林县,内蒙古自治区的伊金霍洛旗、东胜市、准格尔旗、清水河县和达拉特旗,共 12 个县(市、旗)。

本区自然环境恶劣,水资源缺乏,水土流失及风沙危害严重,农、林、牧业都不发达。但是,本区的煤炭资源十分丰富,拥有我国和世界上罕见的特大煤田,探明储量共计 2576 亿吨。为了给本区综合开发治理决策提供依据,倪建华等曾运用 AHP 决策分析法,按总目标、战略目标、发展战略、制约因素和方针政策等五个层次,分析了它们之间的相互联系与相互制约关系,计算出了各层的相对权重,从而得出了这些因素对实现总目标的影响或重要程度,为制定切实可行的方针政策以克服不利因素提供了必要的依据。以下,我们将这一研究成果作一些简单的介绍。

(一)模型结构

1.总目标和战略目标 总目标是晋陕蒙接壤地区的综合开发治理 战略目标:根据本地区的自然、经济和社会条件,我们归纳出下面三个战略目标:

0₁:煤炭开发;0₂:发展农林牧生产;0₃:改善生态环境,力争达到良性循环。

2.发展战略 根据本区特点,开发治理的战略重点是能源、粮食、副食、水土保持、沙化治理等方面,为此我们提出以下十个发展战略:

C₁:发展统配煤矿;

- C₂：发展地方、乡镇煤矿；
- C₃：发展电力工业；
- C₄：发展重工业、化工工业；
- C₅：发展地方工业、乡镇企业；
- C₆：发展粮食生产；
- C₇：建设肉蛋奶基地；
- C₈：建设果品蔬菜基地；
- C₉：水土保持；
- C₁₀：沙漠化治理。

3. 制约因素 晋陕蒙三角地区虽然有不少有利条件，但也有许多不利因素，这对实现总目标必然会产生很大影响，我们归纳了八个方面的制约因素：

- S₁：运输能力低下；
- S₂：资金严重不足；
- S₃：人才、技术力量（包括技术工人，工程技术人员，科研人员，教员等）缺乏；
- S₄：水资源不足；
- S₅：水土流失严重，风沙危害大；
- S₆：粮食及农副畜产品供应紧张；
- S₇：地方乡镇经济不发达；
- S₈：厂矿建设要占用大部分良田。

4. 方针措施 为了克服不利因素，保证总目标实现，可以有如下十九项方针措施：

- P₁：引入国外资金，引进技术；
- P₂：国家投资；
- P₃：地方集资；
- P₄：当地现有水资源开发节流，合理使用；
- P₅：引黄河水；
- P₆：开发地下水；
- P₇：种草种树，发展畜牧；
- P₈：加强农田基建，提高单产；
- P₉：对可能污染环境的厂矿，提前采取措施；
- P₁₀：各省内自行解决人才、技术问题；
- P₁₁：从全国调入人才，引进技术；
- P₁₂：本地区自行解决人才、技术问题；
- P₁₃：各省内解决农副畜产品供应问题；
- P₁₄：地方解决粮食供应；
- P₁₅：省内解决粮食供应；
- P₁₆：从全国调入粮食；
- P₁₇：改善公路运输条件，新建公路；

P_{18} ：修建铁路；

P_{19} ：对重点工矿，加强水保工作及沙化治理。

根据上述分析，可以得出晋陕蒙三角地区综合开发治理战略决策模型的层次结构，如图 9-3 所示。

(二)模型计算结果

根据图 6-3 所示的层次结构，通过构造 AHP 判断矩阵（共构造了 23 个判断矩阵）、层次单排序、层次总排序及一致性检验等步骤，得到了如下几个方面的计算结果：

(1)计算出 3 个战略目标 O_1, O_2, O_3 的相对权重。

(2)计算出发展战略 $C_1, C_2, \dots, C_{10}$ 对每个战略目标的相对权重，并用 O_1, O_2, O_3 的权重对发展战略的相对权重加权后相加，可得出各发展战略的组合权重，它们表示各发展战略对实现总目标的重要程度。

(3)计算出每个制约因素 $S_1, S_2, \dots, S_8$ 对每个发展战略的相对权重，并用发展战略 $C_1, C_2, \dots, C_{10}$ 的组合权重对制约因素的相对权重加权后相加，可得出各制约因素的组合权重，它们表示各制约因素对实现总目标的制约程度。

(4)计算出各方针政策 $P_1, P_2, \dots, P_{10}$ 对每个制约因素的相对权重，并用各制约因素的组合权重对方针措施的相对权重加权后相加，即可得出各方针措施的组合权重。它们表示各方针政策对实现总目标的重要程度。权重越大越重要。因此在实现总目标的过程中，应该首先考虑实施那些权重较大的方针政策。

上述计算结果分别见表 6-4、表 6-5 和表 6-6。

表 6-4 战略目标和发展战略权重

战略目标	O_1	O_2	O_3	组合权重
发展战略	0.595	0.128	0.276	
C_1	0.230	0.028	0.037	0.151
C_2	0.217	0.032	0.020	0.139
C_3	0.044	0.047	0.025	0.039
C_4	0.024	0.020	0.018	0.022
C_5	0.034	0.040	0.024	0.032
C_6	0.137	0.244	0.071	0.133
C_7	0.101	0.206	0.101	0.114
C_8	0.083	0.138	0.149	0.108
C_9	0.070	0.122	0.277	0.134
C_{10}	0.061	0.122	0.277	0.129

(三)结果分析与结论

1. 结果分析

(1)从战略目标来看，要实现晋陕蒙三角地区综合开发与治理，首先要

发挥本地区煤炭资源的优势，其权重为 0.595。但不容忽视的是，必须采取开发与治理并重的总方针，边开发边治理，以开发促治理，从计算结果看， O_3 的权重为 0.276，其重要程度排在第二位。当然，农林牧生产也应得到相应的重视，其权重为 0.128。

(2)从发展战略上来讲，首先应发展统配煤矿，其权重为 0.151；地方乡镇煤矿的发展

表 6-5 发展战略和制约因素的权重

发展战略 制约因素	C_1 0.151	C_2 0.139	C_3 0.039	C_4 0.022	C_5 0.032	C_6 0.133	C_7 0.114	C_8 0.108	C_9 0.133
S_1	0.293	0.253		0.102	0.071				
S_2	0.086	0.274	0.195	0.110	0.408	0.158	0.359	0.262	0.714
S_3	0.041	0.246	0.093	0.101	0.204				
S_4	0.215	0.062	0.515		0.112	0.366	0.241	0.565	
S_5	0.112	0.050		0.388			0.129	0.117	0.214
S_6	0.183	0.032	0.111	0.195		0.337	0.181		
S_7	0.028	0.047	0.043	0.033	0.204	0.074	0.091		
S_8	0.041	0.036	0.041	0.072		0.065		0.565	

表 6-6 制约因素和方针措施的权重

制约因素 方针措施	S ₁ 0.084	S ₂ 0.361	S ₃ 0.053	S ₄ 0.234	S ₅ 0.093	S ₆ 0.106	S ₇ 0.040	S ₈ 0.029	组i
P ₁		0.105							0
P ₂		0.637					0.196	0.135	0
P ₃		0.258					0.657		0
P ₄				0.250					0
P ₅				0.500					0
P ₆				0.250					0
P ₇					0.375				0
P ₈					0.375			0.584	0
P ₉					0.125			0.281	0
P ₁₀			0.594				0.147		0
P ₁₁			0.249						0
P ₁₂			0.157						0
P ₁₃						0.432			0
P ₁₄						0.105			0
P ₁₅						0.116			0
P ₁₆						0.347			0
P ₁₇	0.250								0
P ₁₈	0.750								0
P ₁₉					0.125				

也占有重要地位，其权重为 0.139；水土保持和粮食生产的权重分别为 0.134 和 0.133，处在第三位和第四位，从它们权重的数值可看出与发展采煤业相差不多，可见它们对实现总目标的重要性。沙漠化治理与建设肉蛋奶基地也应放在重要的位置上，其权重分别为 0.129 和 0.114；建设果品和蔬菜基地的权重为 0.108。另外，发展电力工业与发展地方工业乡镇企业的权重分别为 0.039 和 0.032。从计算结果可以看出，发展重化工业的权重为 0.022，在 10 个发展战略中其权重的大小处在最后一位，即本区重化工业的发展对于总目标的实现作用不大，也即本区由于自然条件的影响和某些因素的限制不宜大量发展重化工业。

(3)从制约因素来看，本区资金短缺这一制约因素的影响最为严重，其权重为 0.36；其次，水资源不足对总目标的制约程度也十分严重，其权重为 0.234。另外，粮食和农副产品供应问题、水土流失、风沙危害以及运输能力的不足，也对总目标的实现有较为严重的制约，其权重分别为 0.106，0.093，0.084。人才技术缺乏、地方经济不发达、厂矿占地过多的权重依次为 0.053，0.040，0.029。

(4)从方针措施来说，通过各种渠道解决资金不足的矛盾是最为重要的，包括国家投资、地方集资，其权重分别为 0.242 和 0.119。由于本区干旱少雨，地表径流少，且中小型水库很快就会被泥沙淤满，因此要想开发本

区的煤炭资源，从长远观点来讲必须从黄河提水、引水，引、提黄河水解决本区能源基地建设问题的权重为 0.117，其次本区交通运输问题也是急待解决的矛盾。目前整个地区除了北同宁（武）苛（岚）支线相接的神（池）河（曲）支线（全长 158 公里）已修至阴塔（长 65 公里）以外，没有一条铁路线，因此修建铁路就成为本区能源基地建设的当务之急，其权重为 0.063，本区大力开采地下水和水资源的开源节流工作可能还有一定的潜力，其权重为 0.059。还有，就是本区的农田基本建设工作，其权重为 0.052。

2. 结论

综合上述分析结果，可以得到如下基本结论：

晋陕蒙三角地区综合开发治理总目标的实现，应该采取煤炭工业的发展与环境治理并重的方针，二者不可偏废。必须首先发展统配煤矿与地方乡镇煤矿，同时要注意水土保持、沙漠化防治及粮食生产。而要做到这些，必然要受到自然条件和资源条件的制约，特别是资金、水资源、粮食及农副产品的供应问题、水土流失及风沙危害、运输条件等因素的制约。要解决这些问题，克服不利的制约因素，要求国家增加对本区投资，同时也要积极利用地方资金，广泛集资，发挥资源优势，新建铁路与公路，尽快提高本区的运输能力，从而保证本区综合开发治理战略目标的顺利实现。

第七章 网络分析方法

区域研究是当代地理学的三大研究方向之一，而地理区域的空间结构，即地理事物的空间格局或空间排布，一向被认为是区域研究的核心问题。在一个现实的区域地理系统中，许多地理事物的空间格局或空间排布表现为“点”和“点的连线”构成的地理网络。譬如，构造体系、水系、交通网、城镇体系等。这些由客观存在的地理事物及其相互联系所构成的地理网络，从数学的角度来看，它们又可以被抽象为图论研究中的“图”与“网络”。本章将从地理网络的图论描述入手，对地理系统的网络分析作一些探讨。

第一节 地理网络的图论描述

用图来表示各类地理事物，可以说是地理学的一个基本特征。在早期的地理学研究中，几乎无一例外地都使用各种符号和图象，来代表所要描述的地理实体。因此，对于地理学来说，图一直是一种传统的方法，它的直观性，形象性和使用上的方便性是人人皆知的。但是，这些所谓的图对地理事物之间的相互联系的本质还未能给予充分的揭示，因此，在可比性与地理学的定量分析方面都表现出某些不足，有值得探讨的必要。

一、地理网络的图论描述

在一个现实的地理系统中，许多地理事物及其之间的相互联系，均可以直接地或者经过适当的简化取舍之后间接地抽象为“网络图”的概念。在这种网络图中，地理系统的地理事物被抽象为点，譬如，山岭制高点，河川汇流点，居民点，城市，工矿分布点，交通站点等，甚至人口密度，工业产值，生物种的数量等抽象的概念也可以被看作为分布点。而地理事物之间的相互联系则被抽象为点的连线，譬如，河流，构造带，交通线，通讯线路等，甚至人口流，物质流，能量流，经济流，价值流，信息流等都可以被抽象为地理事物分布点之间的连线。对于许多复杂的地理实体，也可以经过适当的简化取舍将其抽象为地理网络的形象与直观。譬如，由各种基本地貌单元构成的流域地貌系统，就可以被简化为该流域的水系格局，即树状网络图。著名的数学家列昂纳德·欧拉 (Euler) 于 1736 年首次提出的图论问题“七桥问题”，其实质就是一个地理问题。东普鲁士的哥尼斯堡城（现在的加里宁格勒）是建在两条河流的汇合处以及河中的两个小岛上的（见图 7-1），总共有七座小桥将两个小岛及小岛与城市的其它部分连接起来，欧拉的问题是，哥尼斯堡人从其住所出发，能否恰好只经过每座小桥一次而返回原处？图论的研究结果告诉我们，其答案是否定的。

以上分析和事例都充分说明了图论在地理系统研究，特别是地理系统网络分析中的重要应用价值。那么，究竟图是什么？地理网络的图论描述是怎样的？为了便于以后的分析和讨论，以下我们从数学的角度给出地理网络图的定义：

设 V 是由 n 个点 v_i ($i=1, 2, \dots, n$) 所组成的集合，即 $V=\{v_1, v_2, \dots, v_n\}$ ； E 是由 m

条线 e_i ($i=1, 2, \dots, m$) 所组成的集合，即 $E=\{e_1, e_2, \dots, e_m\}$ ，而且 E 中任意一条线，都是以 V 中的点为端点；任意两条线除了端点外没有其它的公共点，那么把 V 与 E 合在一起就称为一个图 G ，记作 $G=(V, E)$ 。 V 中的每一个点 v_i ($i=1, 2, \dots, n$) 称为 G 的顶点； E 中每一条线称为 G 的边，若一条边 e 连接 u, v 两顶点，则记为 $e=(u, v)$ 。如果 G 的每条边给定了方向，即 $(u, v) \neq (v, u)$ ，则称 G 为有向图；如果 G 的每条边都没有方向，即 $(u, v)=(v, u)$ ，则 G 称为无向图。

在网络图 $G=(V, E)$ 中， V 不允许是空集，但 E 可以是空集合。从以上定义可以看出，网络图包含了两个方面的要素：(1) 点集（或者称顶点集）；

(2) 边集（或称弧集）。例如在图 7-2 的网络图 (a) 中，顶点集为 $V=\{v_1, v_2, v_3, v_4, v_5, v_6\}$ ，边集为 $E=\{e_1, e_2, e_3, e_4, e_5,$

$e_6, e_7, e_8\}$ ，记为图 $G=\{V, E\}$ 。在图 (b) 中，顶点集为 $V=\{v_1, v_2, v_3, v_4, v_5, v_6\}$ ，边集为 $A=\{a_1, a_2, a_3, a_4, a_5, a_6, a_7, a_8\}$ ，记为图 $D=(V, A)$ 。图 (a) 与图 (b) 的差别是： G 中的边 (联线) 没有方向，而 D 中的边 (联线) 是有方向的。因此 G 是无向图，而 D 是有向图。如果顶点 $v_1, v_2, v_3, v_4, v_5, v_6$ 分别代表某区域的城镇，而顶点的联线 (边) $e_1, e_2, e_3, e_4, e_5, e_6, e_7, e_8$ 分别代表连接城镇之间的公路，则 (a) 所示的图 $G=(V, E)$ 就表示了该区域的公路交通网络，(b) 所示的有向图 $D=(V, A)$ 则可以理解为由单行线所形成的交通网络。

同理，如果令顶点集 V 表示全国各地所有的铁路站点所构成的集合，边集 E 表示连接所有铁路站点的铁路线构成的集合，则图 $G=(V, E)$ 就表示了全国铁路交通网络；若令顶点集 V 为某一流域地貌系统中所有的河流汇流点组成的集合，而边集 A 代表该流域的所有河流组成的集合，则图 $D=(V, A)$ 就表示了该流域的水系网络；……。

事实上，如果严格地按照图论关于图的定义，只关注点之间是否流通，而不去注重点间的连通方式，那么任何一个图的画法就不是唯一的一种。由地理事物抽象而来的点 (顶点)，可以画成平面上的点，也可以随意地画成没有特定的物理位置的点。而由地理事物之间的相互联系抽象而来的连线 (边)，可以是直线也可以是曲线；可以是有向的，也可以是无向的。其最为重要的概念是表示顶点的“点”和表示连线的“边”，而其相对位置并不是至关重要的。譬如，下图 (图 7-3) 中的 (a) 与 (b) 只不过是同一个图的不同画法。

为了更进一步地对地理网络进行分析，以下我们对图的一些基本概念作一简单介绍。

1. 关联边。若 $e=(u, v) \in E$ ，则称 u, v 是边 e 的端点，称 e 是 u, v 的关联边。

2. 环。若 e 的两个端点相同，即 $u=v$ ，则称为环。

3. 多重边。若联结两个端点的边多于一条以上，则称为多重边。

4. 简单图。无环、无多重边的图称为简单图。

5. 多重图。含有多重边的图称为多重图。

6. 点与次。

(1) 次。以顶点 v 为端点的边的个数称为点 v 的次，记为 $d(v)$ 。

(2) 悬挂点、悬挂边、孤立点。次等于 1 的点称为悬挂点；与悬挂点关联的边称为悬挂边；次为零的点称为孤立点。

(3) 奇点与偶点。次为奇数的点称为奇点；次为偶数的点称为偶点。

7. 路 (链)。若图 $G=(V, E)$ 中的顶点和边交替出现的序列 $P=\{v_{i1}, e_{i1}, v_{i2}, e_{i2}, \dots, e_{ik-1}, v_{ik}\}$ (对于有向图来说，要求排在每一条边之前和之后的顶点分别是这条边的起点和终点) 满足：

$$e_{it} = (v_{it}, v_{i, t+1}) \quad (t=1, 2, \dots, k-1) \quad (1)$$

则称 P 为一条从 v_{i1} 到 v_{ik} 的路 (链)。简记为 $P=\{v_{i1}, v_{i2}, \dots, v_{ik}\}$ 。

8. 回路。若路的第一点与最后一个点相同，即 $v_{i1}=v_{ik}$ ，则称为回路。

9.连通图。在图 G 中，若任何两点之间至少存在一条路（对于有向图，则不考虑其边的方向），则称 G 为连通图。否则称为不连通图。

10.树。不含回路的连通的无向图称为树。

11.赋权图。如果图 $G = (V, E)$ 中的每一条边 (v_i, v_j) 都相应地赋有一个数值 ω_{ij} ，则称 G 为赋权图。其中， ω_{ij} 称为边 (v_i, v_j) 的权值。

12.基础图。从一个有向图 $D = (V, A)$ 中去掉所有边上的箭头所得到的一个无向图，就称为 D 的基础图，记之为 $G(D)$ 。

13.截。从图中移去边的一个集合将增加亚图的数目时，被移去的边集合就称为截。

14.子图。设 $G = (V, E)$ 是一个无向图， V_1 与 E_1 分别是 V 与 E 的子集，即 $V_1 \subseteq V$ ， $E_1 \subseteq E$ 。如果对于任意 $e_i \in E_1$ ，其两个端点都属于 V_1 ，则称 $G_1 = (V_1, E_1)$ 是图 G 的一个子图。

15.支撑子图。设 $G_1 = (V_1, E_1)$ 是图 $G = (V, E)$ 的子图，如果 $V_1 = V$ ，则称 G_1 是 G 的支撑子图。

16.支撑树。设 $G = (V, E)$ 是一个无向图，如果 $T = (V_1, E_1)$ 是 G 的支撑子图，并且 T 是树，则称 T 是 G 的一个支撑树。

17.树的重量。一个树的所有边的权值之和称为该树的重量。

18.最小支撑树。在一个图的所有支撑树中，重量最小的那个叫做该图的最小支撑树。

二、地理网络的有关测度矩阵与测度指标

当一个现实地理系统中的客观地理事物及其之间的相互联系被抽象为图论描述的地理网络图时，我们就可以借助于有关的测度矩阵或者测度指标对其性质、特征及复杂性等进行定量化描述。

(一)关联矩阵与邻接矩阵

1.关联矩阵 关联矩阵是对网络中的顶点与边的关联关系的一种描述。假设网络图 $G = (V, E)$ 的顶点集为 $V = \{v_1, v_2, \dots, v_n\}$ ，边集为 $E = \{e_1, e_2, \dots, e_m\}$ ，则该网络图的关联矩阵为一个 $n \times m$ 矩阵，可表示为：

$$L(G) = \begin{Bmatrix} g_{11} & g_{12} & \cdots & g_{1m} \\ g_{21} & g_{22} & \cdots & g_{2m} \\ \cdots & \cdots & \cdots & \cdots \\ g_{n1} & g_{n2} & \cdots & g_{nm} \end{Bmatrix} \quad (2)$$

(2)式中， g_{ij} 为顶点 v_i 与边 e_j 相关联的次数，显然， g_{ij} 的取值只能是 0, 1, 或者 2。

例如，图 7-3 中 (a) 或 (b) 的关联矩阵为：

$$L(G) = \begin{pmatrix} 1 & 0 & 0 & 0 & 1 & 0 & 1 \\ 1 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 0 \end{pmatrix}$$

2.邻接矩阵 邻接矩阵是对网络图中各顶点之间的连通性程度进行描述的一种矩阵。假设图 $G = (V, E)$ 的顶点集为 $V = \{v_1, v_2, \dots, v_n\}$ ，则邻接矩

阵是一个 $n \times n$ 的方阵，可表示为：

$$A(G) = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \cdots & \cdots & \cdots & \cdots \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{pmatrix} \quad (3)$$

(3)式中， a_{ij} 表示连接顶点 v_i 与 v_j 的边的数目。

例如，图 7-3 中 (a) 或 (b) 的邻接矩阵为：

$$A(G) = \begin{pmatrix} 0 & 1 & 1 & 0 & 1 \\ 1 & 0 & 1 & 0 & 0 \\ 1 & 1 & 0 & 1 & 1 \\ 0 & 0 & 1 & 0 & 1 \\ 1 & 0 & 1 & 1 & 0 \end{pmatrix}$$

(二) 有关的测度指标

对于任何一个网络图，都存在着三种共同的基础指标：①在网络中次级亚网图的分割数目，即互不连接数目；②在网络中的连线数目；③网络中的接点数目。如果将以上三个基础指标分别记为 p 、 m 和 n ，则我们可以得到关于网络结构的一般性测度指标如下。

1) β 指数。 β 指数也称为线点率，它是网络内每一个接点的平均联线数目，即

$$\beta = \frac{m}{n} \quad (4)$$

β 指数是关于地理网络的复杂性程度的简单度量，其数值的范围在 $[0, 3]$ 区间上。 $\beta = 0$ ，表示无网络存在；网络的复杂性增加，则 β 值增大。一种最低限度的连结，被定义为没有孤立点存在的网络，其连线的个数为 $n - P$ ，此时的 β 指数为

$$\beta_1 = \frac{n - P}{n} \quad (5)$$

如果地理网络不包含次级亚网，即 $P=1$ ，则其最低限度连接的 β 指数值为 $1 - \frac{1}{n}$ 。

2) 回路数。通过前面的介绍，我们知道回路是一种闭合路径，其始点同时为终点。除非联线连接两个亚网络，回路就可能存在。如果网络内存在回路，则连线的数目就必须超过 $n - P$ ，因为这个数目定义着一个最低限度的连接网。那么回路数目为实际连线数目减去最低限度连结的联线数目，即

$$k = m - n + P \quad (6)$$

3) α 指数。 α 指数是测度网络回路的指标，它是网络中实际回路的观察数与网络内可能存在的回路最大数目之间的比率。在任何接点超过 2 的网络中，增添一个额外接点，增加的最多连线数为 3。因此，一个网络的连线的最大可能数目为： $3(n - 2P)$ 。所以，网络内可能存在的回路最大数目为联线的最大可能数目减去最低限度连结的联线数目，即

$$\begin{aligned} & 3(n - 2P) - (n - P) \\ & = 2n - 5P \end{aligned}$$

所以 α 指数为

$$\alpha = \frac{m - n + P}{2n - 5P} \quad (7)$$

α 指数的变化范围介于 $[0, 1]$ 区间， $\alpha = 0$ 意味着网络中不存在回路； $\alpha = 1$ ，说明网络中已达到最大限度的回路数目。 α 指数亦可用百分率来表示，即

$$\alpha = \frac{m - n + P}{2n - 5P} \times 100\% \quad (8)$$

4) γ 指数。 γ 指数是网络内连线的实际观察数与连线可能存在的最大数目之间的比率，即

$$\gamma = \frac{m}{3(n - 2P)} \quad (9)$$

γ 指数是测度网络连通性的一种度量指标，其数值变化范围为 $0 \leq \gamma \leq 1$ 。 $\gamma = 0$ ，表示网络内无连线，只有孤立点存在； $\gamma = 1$ ，则表示网络内每一个接点都存在与其它所有接点相连的连线。 γ 指数亦可以用百分比来表示，即

$$\gamma = \frac{m}{3(n - 2P)} \times 100\% \quad (10)$$

表 7-1 给出了网络图 7-4 中 (a)，(b)，(c)，(d) 四个不同网络的有关测度指标。

通过对以上四个不同的测度指标的计算和分析，不但使我们比较容易地定量地了解我们所研究的地理网络的一些基本特征，如网络的复杂性程度、连通性状况等，而且还可以通过这些指标间接地反映该网络所描述的地理系统的其它内容。譬如，通过对一个国家或地区交通网络的 β 指数的计算和分析，我们不但可以了解该国家或地区交通网络的连通性和复杂性，同时由于交通网络的连通性和复杂性也是一个国家或地区的社会经济发展状况的反映，因而

表 7-1 图 7-4 中不同网络的有关测度指标

测度指标类型	网络图			
	(a)	(b)	(c)	(d)
β	0.90	1.40	1.90	2.40
k	0	5	10	15
α	0	0.33	0.67	1.00
γ	0.375	0.583	0.792	1.00

β 指标也间接地反映了该国家或地区的社会经济发展水平；对于一个流域水系网络，如果 β 值较高，则说明该流域水系错综复杂，河流遍布，同时也表示该流域地貌系统具有复杂的地形构成。由此可以看出，以上介绍的有关测度指标，不仅是地理网络的复杂性、连通性和回路性等性质的度量，而且也反映着现实地理系统的其它内容。因而这些指标在有关的地理系统分析中有广泛的应用前景。

第二节 最短路径与选址问题

当地理系统中的地理事物及其之间的相互联系被抽象为地理网络图时，许多地理系统分析问题就转化为图的计算问题，其中最为常见的是路径与顶点的计算。在路径计算中最为常用的是最短路径问题；而在顶点的计算中最为常见的是中心点与中位点的计算问题。目前，对于这些问题，图论研究已给我们提供了一些比较成熟的算法，下面分别作一简单介绍。

一、最短路径问题

最短路径问题，是地理系统网络分析中的一个重要问题之一，这一问题的研究不但具有重大的理论意义，而且具有重要的应用价值。交通网络结构的分析，交通运输线路（公路、铁路、河流航运线、航空线、管道运输线路等）的选择，通讯线路的建造与维护，运输货流的最小成本分析，城市公共交通网络的规划等，都有直接应用的价值。

（一）最短路径的含义

在地理系统的网络分析中，最短路径可以代表多种不同的含义。一般而言，它主要指以下三个方面的含义：

1. “纯距离”意义上的最短路径。例如，某旅行者想要乘汽车沿着一个国家或者地区的公路交通网络中的某一条线路，从一个城市到达另一个城市，那么这个旅行者就要考虑究竟沿怎样的路线（可以不是直达的线路）行进的距离最短？显然，这里所谓的距离是指实际的里程数，即“纯距离”。

2. “经济距离”意义上的最短路径。例如，某公司在世界上某 10 大港口 C_1 、 C_2 、 $\dots$ 、 C_n 均设有货栈，为了更好地适应市场供需的形势，经常要在各港口之间运送各类货物。从 C_i 到 C_j 之间的直接航运价格，可以从市场动态中得到。如果两个港口之间无直接通航路线时，可以通过第三个港口转运。随着市场动态的变化，该公司的经理就希望求出任意两个港口之间最为廉价的运货线路，即两个港口之间最短的“经济距离”。

3. “时间”意义上的最短路径。例如，某家经营公司有一批货物急需从一个城市运往另一个城市，那么在由公路、铁路、河流航运、航空运输等四种运输方式线路所构成的交通网络中，究竟选择怎样的送货路线（运输方式不限）最节省时间，即“时间”意义上的距离最短。

以上三例问题，我们都可以将其抽象为一类问题，即赋权有向图上的最短路径问题。在这里，不同意义的距离都被抽象为网络图中边的权数。自然，这种权即可以代表“纯距离”，又可以代表“经济距离”，也可以代表“时间距离”。针对不同的问题，其距离（权数）和最短路径的含义可以不同。

（二）最短路径的算法

关于最短路径问题，目前为人们所公认的最好求解方法，是 1959 年由 E.W.Dijkstra 提出的标号法。之所以被称为标号法，是由于在这个方法实施的过程中，对网络图中每一个顶点都对应一个数（或几个数），这个数就称之为该顶点的标号。这个方法的一个突出优点是：它不仅求出了起点到终点的最短路径及其长度，而且求出了起点到图中其它各个顶点的最短路径及其长度。更值得一提的是，这一算法同样也适用于求无向图上的最短路径问题。下面我们具体地介绍这一算法。

设 $G = (V, E)$ 是一个赋权有向图，即对于图中的每一条边 $e = (v_i, v_j)$

都赋有权值 ω_{ij} 。在图 G 中指定两个顶点，确定为起点和终点，不妨设 v_i 为起点， v_p 为终点。标号法的基本思想是：首先从 v_1 开始，给每一个顶点记一个数（称为标号），标号分 T 标号和 P 标号两种， T 标号表示从起点 v_1 到这一点的最短路权的上界，称为临时标号； P 标号表示从 v_1 到该点的最短路权，称为固定标号。已得到 P 标号的顶点不再改变，凡是没有标上 P 标号的顶点，标上 T 标号。算法的第一步就是把某一顶点的 T 标号改变为 P 标号。那么，最多经过 $P-1$ 步，就可以得到从起点 v_1 到每一个顶点的最短路径。其具体计算步骤如下：

开始，先给 v_1 标上 P 标号 $P(v_1)=0$ ，其余各点标上 T 标号 $T(v_j)=+\infty$ ($j \neq 1$)

1. 对于刚刚得到 P 标号的点 v_i ，考虑所有这样的点 v_j ，使 $(v_i, v_j) \in E$ ，而且 v_j 的标号是 T 标号，然后修改 v_j 的 T 标号为 $\min[T(v_j), P(v_i) + \omega_{ij}]$ 。

2. 若 G 中没有 T 标号，则停止。否则把点 v_{j_0} 的 T 标号修改为 P 标号，然后再转入(1)。其中 v_{j_0} 满足：

$$T(v_{j_0}) = \min T(v_j)$$

v_{j_0} 是 T 标号。

例：在网络图7-5中，边旁的数字表示该边的权。求从 v_1 到 v_7 的最短路径。

开始，首先给 v_1 标上 P 标号 $P(v_1)=0$ ，表示从 v_1 到 v_1 的最短路权为零。其它点 $(v_2, \dots, v_7)$ 标上 T 标号 $T(v_j)=+\infty$ ($j=2, 3, \dots, 7$)。

第一步：①因为 (v_1, v_2) ， (v_1, v_3) ， $(v_1, v_4) \in E$ ，而且 v_2, v_3, v_4 是 T 标号，则修改这三个点的 T 标号为：

$$T(v_2) = \min[T(v_2), P(v_1) + \omega_{12}] = \min[+\infty, 0+2] = 2$$

$$T(v_3) = \min[T(v_3), P(v_1) + \omega_{13}] = \min[+\infty, 0+5] = 5$$

$$T(v_4) = \min[T(v_4), P(v_1) + \omega_{14}] = \min[+\infty, 0+3] = 3$$

②在所有 T 标号中， $T(v_2)=2$ 最小，于是令 $P(v_2)=2$ 。

第二步：① v_2 是刚得到 P 标号的点，故考察 v_2 。因为 (v_2, v_3) ， $(v_2, v_6) \in E$ ，而且 v_3, v_6 是 T 标号，故修改 v_3, v_6 的 T 标号为

$$T(v_3) = \min[T(v_3), P(v_2) + \omega_{23}] = \min[5, 2+2] = 4$$

$$T(v_6) = \min[T(v_6), P(v_2) + \omega_{26}] = \min[\infty, 2+7] = 9$$

②在图 G 中的所有 T 标号中， $T(v_4)=3$ 最小，于是令 $P(v_4)=3$ 。

第三步：①考察 v_4 ，因为 $(v_4, v_5) \in E$ ，而且 v_5 是 T 标号，故修改 v_5 的 T 标号：

$$(v_5) = \min[T(v_5), P(v_4) + \omega_{45}] = \min[+\infty, 3+5] = 8$$

②在图 G 中的所有 T 标号中， $T(v_3)=4$ 最小，故令 $P(v_3)=4$ 。

第四步：①考察 v_3 ，因为 (v_3, v_5) ， $(v_3, v_6) \in E$ ，而且 v_5, v_6 为 T 标号，故修改 v_5, v_6 的 T 标号为：

$$T(v_5) = \min[T(v_5), P(v_3) + \omega_{35}] = \min[8, 4+3] = 7$$

$$T(v_6) = \min[T(v_6), P(v_3) + \omega_{36}] = \min[9, 4+5] = 9$$

②在所有的 T 标号中, $T(v_5) = 7$, 令 $P(v_5) = 7$ 。

第五步: ①考察 v_5 , (v_5, v_6) , $(v_5, v_7) \in E$, 而且 v_6 和 v_7 都是 T 标号, 故修改 v_6 、 v_7 的 T 标号:

$$T(v_6) = \min[T(v_6), P(v_5) + \omega_{56}] = \min[9, 7+1] = 8$$

$$T(v_7) = \min[T(v_7), P(v_5) + \omega_{57}] = \min[\infty, 7+7] = 14$$

②在所有 T 标号中, $T(v_6) = 8$ 最小, 令 $P(v_6) = 8$ 。

第六步: ①考察 v_6 , $(v_6, v_7) \in E$, 而且 v_7 为 T 标号, 故修改 v_7 的 T 标号:

$$T(v_7) = \min[T(v_7), P(v_6) + \omega_{67}] = \min[14, 8+5] = 13。$$

②令 $P(v_7) = 13$

所以从 v_1 到 v_7 的最短路径为 $(v_1, v_2, v_3, v_5, v_6, v_7)$, 最短路径长度为 13。

二、选址问题

选址问题, 是当代地理学三大方向之一——区位论研究的主要理论与应用课题之一。选址问题涉及人类生产、生活、文化、娱乐等各个方面。例如, 城市商业区、生活区、工厂、仓库、车站、机场、消防站点、文化娱乐场所等布局问题, 都是选址问题。选址问题的数学模型取决于可供选址的范围, 以及怎样判定选址的质量两个方面的条件。由于存在大量的各种各样的选址问题, 所以有关文献中也有各种各样的选址问题的数学模型及求解方法。但是, 我们这里讨论的仅限于选址的范围是一个网络图, 而且选址位置必须位于网络图的某一个或几个顶点上 (当然亦可以位于一条边的某一个位置上, 关于这方面的讨论, 有兴趣的读者可以参看有关的图论著作)。对这样的选址问题, 根据其选址的质量判据, 我们又可以将其归纳为求网络图的中心点与中位点两类问题。

(一) 中心点选址问题

中心点选址问题的质量判据, 是使最佳选址位置所在的顶点与图中其它顶点之间的最大距离达到最小。这类选址问题适宜于医院、消防站点等一类服务设施的布局问题。例如, 某县要在其所辖的六个乡镇之一修建一个消防站, 为六个乡镇服务, 要求消防站至最远乡镇的距离达到最小。这就是选址问题, 实际上就是求网络图的中心点问题。其数学描述为:

设 $G = (V, E)$ (其中 $V = \{v_1, v_2, \dots, v_n\}$, $E = \{e_1, e_2, \dots, e_n\}$ 是一个无向简单连通赋权图), 联接两个顶点的边的权值代表该两顶点之间的距离, 对于每一个顶点 v_i , 它与各顶点之间的最短路径长度为 $d_{i1}, d_{i2}, \dots, d_{in}$ 。这几个距离中的最大数称为顶点 v_i 的最大服务距离, 记为 $e(v_i)$ 。那么, 中心点选址问题, 就是求图 G 的中心点 v_{i_0} , 使得

$$e(v_{i_0}) = \min\{e(v_i)\} \quad (1)$$

假若上例中的六个乡镇及其之间的距离被抽象为图 7-6 所示的无向简单连通赋权圈, 试问消防站点应设在哪一个乡镇, 即哪一个顶点?

首先，用标号法求出每一个顶点 v_i 至其它各顶点 v_j 的最短路径长度 d_{ij} ($i, j=1, 2, \dots, 6$)，写出距离矩阵：

$$D = \begin{pmatrix} d_{11} & d_{12} & d_{13} & d_{14} & d_{15} & d_{16} \\ d_{21} & d_{22} & d_{23} & d_{24} & d_{25} & d_{26} \\ d_{31} & d_{32} & d_{33} & d_{34} & d_{35} & d_{36} \\ d_{41} & d_{42} & d_{43} & d_{44} & d_{45} & d_{46} \\ d_{51} & d_{52} & d_{53} & d_{54} & d_{55} & d_{56} \\ d_{61} & d_{62} & d_{63} & d_{64} & d_{65} & d_{66} \end{pmatrix} = \begin{pmatrix} 0 & 3 & 6 & 3 & 6 & 4 \\ 3 & 0 & 3 & 4 & 5 & 7 \\ 6 & 3 & 0 & 3 & 2 & 4 \\ 3 & 4 & 3 & 0 & 5 & 7 \\ 6 & 5 & 2 & 0 & 5 & 2 \\ 4 & 7 & 4 & 7 & 2 & 0 \end{pmatrix}$$

其次求距离矩阵中每行的最大值得 $e(v_1)=6, e(v_2)=7, e(v_3)=6, e(v_4)=7, e(v_5)=6, e(v_6)=7$ 。显然： $e(v_1)=e(v_3)=e(v_5)=\min\{e(v_i)\}=6$ 。所以，消防站应建在 v_1, v_3, v_5 点所在的乡镇即可。

(二)中位点选址问题

中位点选址问题的质量判据，是使最佳选址位置所在的顶点到网络图中其它顶点的距离（亦可以是加权距离）总和达到最小。这类选址问题的数学描述为：

设 $G = (V, E)$ （其中 $V = \{v_1, v_2, \dots, v_n\}$ ， $E = \{e_1, e_2, \dots, e_n\}$ ）是一个简单连通赋权无向图，联结两个端点的边的权值为该两点之间的距离；对于每一个顶点 v_i ($i=1, 2, \dots, n$)，有一个正的负荷 $a(v_i)$ ，而且它与其它各顶点之间的最短路径长度为： $d_{i1}, d_{i2}, \dots, d_{in}$ ，那么，中位点选址问题就是求图 G 中的中位点 v_{i0} ，使得：

$$S(v_{i0}) = \min_i S(v_i) = \min_i \left\{ \sum_{j=1}^n a(v_j) d_{ij} \right\} \quad (2)$$

例如，某县下属七个乡镇，各乡镇所拥有的人口 $\alpha(v_i)$ ($i=1, 2, \dots, 7$)，以及各乡镇之间的距离 ω_{ij} ($i, j=1, 2, \dots, 7$)如图 7-7 中所示，现要在七个乡镇之一建一个邮局，为全县所属的七个乡镇服务，试问哪一个乡镇所在地（图中的哪一个顶点）应该作为邮局的选址？显然，邮局的最佳应该是图 7-7 的中位点。

首先，用标号法求出每一个顶点 v_i 至其它各个顶点 v_j 的最短路径长度 d_{ij} ($i, j=1, 2, \dots, 7$)，得到距离矩阵：

$$D = \begin{pmatrix} d_{11} & d_{12} & d_{13} & d_{14} & d_{15} & d_{16} & d_{17} \\ d_{21} & d_{22} & d_{23} & d_{24} & d_{25} & d_{26} & d_{27} \\ d_{31} & d_{32} & d_{33} & d_{34} & d_{35} & d_{36} & d_{37} \\ d_{41} & d_{42} & d_{43} & d_{44} & d_{45} & d_{46} & d_{47} \\ d_{51} & d_{52} & d_{53} & d_{54} & d_{55} & d_{56} & d_{57} \\ d_{61} & d_{62} & d_{63} & d_{64} & d_{65} & d_{66} & d_{67} \\ d_{71} & d_{72} & d_{73} & d_{74} & d_{75} & d_{76} & d_{77} \end{pmatrix} \dots$$

$$= \begin{pmatrix} 0 & 3 & 5 & 6.3 & 9.3 & 4.5 & 6 \\ 3 & 0 & 2 & 3.3 & 6.3 & 1.5 & 3 \\ 5 & 2 & 0 & 2 & 5 & 3.5 & 5 \\ 6.3 & 3.3 & 2 & 0 & 3 & 1.8 & 3.3 \\ 9.3 & 6.3 & 5 & 3 & 0 & 4.8 & 6.3 \\ 4.5 & 1.5 & 3.5 & 1.8 & 4.8 & 0 & 1.5 \\ 6 & 3 & 5 & 3.3 & 6.3 & 1.5 & 0 \end{pmatrix}$$

由上述距离矩阵可以求得：

$$S(v_1) = \sum_{j=1}^7 a(v_j) d_{1j} = 122.3$$

$$S(v_2) = \sum_{j=1}^7 a(v_j) d_{2j} = 71.3$$

$$S(v_3) = \sum_{j=1}^7 a(v_j) d_{3j} = 69.5$$

$$S(v_4) = \sum_{j=1}^7 a(v_j) d_{4j} = 69.5$$

$$S(v_5) = \sum_{j=1}^7 a(v_j) d_{5j} = 108.5$$

$$S(v_6) = \sum_{j=1}^7 a(v_j) d_{6j} = 72.8$$

$$S(v_7) = \sum_{j=1}^7 a(v_j) d_{7j} = 95.3$$

因此 $S(v_3) = S(v_4) = \min\{S(v_i)\}$

$$= \min_i \left\{ \sum_{j=1}^7 a(v_j) d_{ij} \right\} = 69.5$$

即 v_3 和 v_4 都是图 7-7 的中位点。所以邮局的最佳选址应该为 v_3 或 v_4 。

第三节 流的计算

地理系统中的流，泛言之，就是指目标从一个地点向另一个地点转移的运动过程。譬如，将原料从产地运往加工厂，将产品由加工厂运往销售地，将人或货物从一个地点运向另一个地点，河水由一个汇流点流向另一个汇流点，等等，都可以认为是流。虽然，地理系统中的流可以有不同的形式，但我们所关心的是各种流的共性问题，例如，我们希望从一个地点经由运输系统到达另一个地点的运输量达到最大，或者从一个地点经由运输系统将一定数量的目标送至另一地点所需的费用达到最小。

一、有关概念及流的增值算法

地理系统中各种形式的流，均可被描述为网络图中的流，即从一个顶点沿着赋权有向图的边（或称弧），将目标移送至另一个顶点的一条路径。目标开始其行程的顶点称为发点，记为 s ；目标终止其行程的顶点称为收点，记为 t 。从发点流向收点的目标称为单位流量或单位。如上所述，单位流量可以是原料、产品、其它货物或者人，也可以是其它物质、能量或信息。

如果通过边 (u, v) 的单位流量是有限制的，则称边 (u, v) 为有容量的边。我们用 $m(u, v)$ 表示边 (u, v) 的最大容量，用 $c(u, v)$ 表示通过边 $e = (u, v)$ 移运一个单位的费用。例如，一支由 12 辆军用卡车组成的战时车队必须从军需工厂（发点）运输军需用品至部队所在地（收点）。联结军需工厂与部队所在地的公路可以用图表示，图的边对应于未防护的公路，顶点对应于公路汇合点。为了安全上的理由，经过每条边的卡车数量有一个限制（边的容量）。根据以往的经验，调度员可以知道卡车在每一个公路区段上费用（边的费用）。求 12 辆军用卡车中每一辆的最佳路线，就是对应的公路网络中，在不超过边的容量的条件下，从出发点向收点运输 12 个单位流量的费用达到最小。

设给定一个图，其中有一定数量的单位由发点运至收点，而且每个单位所采取的路线也为已知。通过边 (u, v) 的单位的数量称为边 (u, v) 的流，用 $f(u, v)$ 表示。显然， $0 \leq f(u, v) \leq m(u, v)$ 。图中的边可以分为三类：

- 1) N ，流不能有任何增加或减少的边集合；
- 2) I ，流可以增加的边集合；
- 3) R ，流可以减少的边集合。

例如，具有零容量或过高通行费用的边属于集合 N ；具有未使用容量的边属于集合 I ；已有流的边属于集合 R 。集合 I 中的边称为可增加的边；集合 R 中的边称为可减少的边；显然，每一条边必然至少属于这三种集合 N, I, R 之一。一条边还可能同时属于两种集合 I 和 R ，当每一条边已经运送单位流量，但是可以有增加或减少的流时，就可以产生这种情况。这样的边称为中间边。

令 $i(u, v)$ 表示边 (u, v) 中的流可以增加的最大量， $r(u, v)$ 表示边 (u, v) 中的流可以减少的最大量。于是， $i(u, v) = m(u, v) - f(u, v)$ ； $r(u, v) = f(u, v)$ 。

设我们欲从 s 再运送某些单位至 t ，有几种方法可以完成这一任务（当然，前提是从 s 可以运送更多的单位至 t ）。首先，如果能够找出一条从 s 至 t 的路 p ，而且这条路全部由可增加的边组成，则从 s 沿着路 p 就可以再运送追加的流至 t （见图 7-8）。从 s 沿着路 p 可以运送多少追加的单位流量至 t

呢？因为 $i(u, v)$ 表示边 (u, v) 中的流可以增加的最大量，所以至多可以有 $\min_{(u,v) \in P} \{i(u, v)\}$ 追加的单位流量从 s 运送至 t 。对图 7-8 中的路 p ，从 s 至 t 可以运送一个追加的单位流量，因为：

$$\min \{i(s, a), i(a, b), i(b, t)\} = \min \{3, 2, 1\} = 1。$$

其次，如果我们能够找到一条从 t 到 s 的路 p ，而这条路全部由可以减少的边组成，则这条路中每条边中的流可以减少，形成从 t 到 s 的较小的流，因而就形成了从 s 到 t 的较大的净流（见图 7-9）。这条路中可以减少的流的最大量是多少呢？由于每条边 (u, v) 最多可以有减少的流 $r(u, v)$ ，所以路 p 中可以减少的最大流为 $\min_{(u,v) \in P} \{r(u, v)\}$ 。对于图 7-9 中的路，

从 t 到 s 可以取消一个单位流量，因为 $\min \{r(t, b), r(b, a), r(a, s)\} = \{1, 2, 1\} = 1$ 。能否找到另一种方法来增加从 s 到 t 的净流呢？这是可能的。将上述两种方法加以合并，我们可以找到一条从 s 至 t 的链，这条链具有以下性质：

(1) 从 s 到 t 方向的边称为前向边，它们属于集合 I ；

(2) 从 t 到 s 方向的边称为后向边，它们都可以属于集合 R 。

例如，考虑图 7-10 中从 s 至 t 的链 C 。前向边为 (s, a) ， (a, b) 和 (d, t) ；后向边为 (c, b) 和 (d, c) 。如果每条前向边都属于 I ，并且每条后向边都属于 R ，则沿着可增加的前向边前进，沿着可减少的后向边倒退，这样，从 s 沿着链 C 就可以运送追加的流至 t 。从 s 沿着这样一条链至 t 可以运送的追加的流的最大量是下列两个数的最小值：

$$\min \{i(u, v) : (u, v) \text{ 为前向边}\}$$

$$\min \{r(u, v) : (u, v) \text{ 为后向边}\}$$

这两个数的最小值称为链 C 的最大流的增值。在图 7-10 中，前向边的可以增加的流的最大量为 $\min \{i(s, a), i(a, b), i(d, t)\} = \min \{4, 3, 3\} = 3$ ；后向边的流可以减少的最大量为 $\min \{r(d, c), r(c, b)\} = \min \{5, 2\} = 2$ 。所以沿着这条链从 s 至 t 可以再运送 2 个追加的单位流量。

从 s 至 t 有任一条链（如上面所讨论的三种不同形式那样），沿着这条链可以运送追加的单位流量，这样的链称为流的增值链。我们怎样确定从 s 至 t 是否存在一条流的增值链呢？应用下列算法，即所谓流的增值算法很容易解决这一问题，这种算法的基本思路是，从发点 s 生长出一棵由已着色的边组成的树，追加的单位流量可以沿着这些边从 s 运出。将边按上述分类方法分为集合 R ， I ， N ，从 s 至 t 已着色的树中唯一链就是从 s 至 t 的流的增值链时，算法对收点 t 着色；从 s 至 t 不存在流的增值链时，算法不对收点 t 着色。流的增值算法的步骤是：

第一步：确定哪些边属于集合 N 、 I 和 R 。因为在集合 N 中的边、流不能改变，所以这些边可以忽略，不予考虑。对顶点 s 着色。

第二步：按照下列规则，对边和顶点着色，直至顶点 t 已被着色或不可能再进行着色为止。

如果顶点 u 已着色，而顶点 v 尚未着色，则在下列情形之一时，可对顶点 v 和边 (u, v) 着色：

(1) 如果边 $(u, v) \in I$ ，则可对顶点 v 和边 (u, v) 着色。

(2)如果边 $(u, v) \in R$, 则对顶点 v 和边 (u, v) 着色。

如果顶点 t 已被着色, 则从 s 至 t 存在一条由已着色边组成的唯一链。这条链就是流的增值链。如果在算法终止以后, t 仍未被着色, 则从 s 至 t 不存在流的增值链。

例: 应用流的增值算法, 找出图 7-11 中从 s 至 t 的一条流的增值链。每条边旁的字母表示这条边是可增加 (可减少或者不变) 的边。

开始对顶点 s 着色。从顶点 s , 可对顶点 a 和边 (s, a) 着色, 因为 $(s, a) \in I$ 。从顶点 s , 不能对顶点 c 和边 (s, c) 着色, 因为 $(s, c) \notin I$ 。这样就完成了从顶点 s 向外的着色。然后, 考虑从顶点 a 向外的着色。因 $(a, b) \notin I$, 故不能对顶点 b 和边 (a, b) 着色。 $(a, c) \in I$, 故可对顶点 c 和边 (a, c) 着色。又 $(d, a) \notin R$, 故对顶点 d 和边 (d, a) 不能着色。这样就完成了从顶点 a 向外的着色。同理, 考虑顶点 c 向外的着色可对 d 和 (c, d) 及 b 和 (b, c) 着色。

我们再考虑从顶点 b 向外的着色。顶点 a 、 c 、 d 已经着色, 所以不予考虑。 $(b, t) \in I$, 故顶点 t 和边 (b, t) 可以着色。此时, 因为 t 已着色, 所以算停止。已着色的边和顶点如图 7-12 所示。从 s 至 t 的流的增值链为: (s, a) 、 (a, c) 、 (b, c) 、 (b, t) , 最大流的增值为 $\min \{i(s, a), i(a, c), r(b, c), i(b, t)\} = \min \{4, 3, 2, 2\} = 2$ 。前向边 (s, a) , (a, c) , (b, t) 可以增加流 2 个单位; 后向边 (b, c) 可以减少流 2 个单位。这意味着以前经过 (b, c) 的两个单

位流量应改变路线而沿着 (b, t) 行进, 在顶点 c 处, 则由 s 点经过 (s, a) 和 (a, c) 到达的追加的两个单位流量来代替。

二、最大流的计算

最大流问题, 就是在一个有容量的图中从发点到收点找出一条可以运送最大数量单位流量的途径, 而且不得超过边的容量。例如, 有一旅行代办人必须为某天由 s 机场至 t 机场的 10 名游客的飞行作出安排。在那天, s 机场至 t 机场的直达航线只有 7 个座席, 而 s 机场至 a 机场的航线有 5 个座席, a 机场至 t 机场的航线有 4 个座席。那么该旅行代办人应该怎样安排? 这一问题可以被描述为最大流的问题。首先可构造一个网络图, 其中每一条边表示一条航线, 这样共有三条边 (s, t) , (s, a) , (a, t) 。对每条边指定一个容量, 这个容量代表每条边所表示的航线上的坐席数, 即 $m(s, t)=7$, $m(s, a)=5$, $m(a, t)=4$ 。如果这个网络容许 10 个或 10 个以上单位的流从 s 点 (发点) 至 t 点 (收点) 而不超过任一条边的容量, 则旅行代办人在选定日期内可将全部游客送走。

如前所述, 流就是从 s (发点) 至 t (收点) 的任何一种运送。如果令 $f(u, v)$ 表示经由边 (u, v) 运送的单位数量, 则对从 s 至 t 的任何一个流它都具有以下几个性质:

(1) 从每一个顶点 u ($u \neq s, u \neq t$) 出去的单位数量一定等于到达该顶点的单位数量, 即

$$\sum_{v \in V} f(u, v) - \sum_{v \in V} f(v, u) = 0 \quad (1)$$

(1) 式中, V 表示网络图中所有顶点组成的集合。

(2)经由边 (u, v) 运送的流的单位数量一定不能大于该边的容量, 即

$$0 \leq f(u, v) \leq m(u, v) \quad ((u, v) \in E) \quad (2)$$

(2)式中, E 为网络图中所有边组成的集合。

(3)从发点出去的单位流量的净数 Q 一定等于到达收点的单位流量的净数, 即 Σ

$$f(s, v) - \sum_{v \in V} f(v, s) = Q \quad (3)$$

$$\sum_{v \in V} f(v, t) - \sum_{v \in V} f(t, v) = Q \quad (4)$$

从 s 至 t 的每一个流都必须满足上述四个条件 (1) — (4), 而且, 如果可以找到一个满足上述四个条件的数值 $f(u, v) [(u, v) \in E]$ 的集合, 则这些数值就对应于一个从 s 至 t 的流。因此, 对于 $(u, v) \in E$, 数值 $f(u, v)$ 的集合是一个流的充要条件是关系式 (1) — (4) 成立。

最大流问题就是求 Q 的最大值, 它必须受关系式 (1) — (4) 条件的约束。显然, 这是一个线性规划问题, 可以单纯形方法求解。但是, 福特和富尔克逊 (Ford, Fulkerson) 于 1962 年提出了一个更好的、更直观的求解方法, 即最大流的算法。这个算法的思路十分简单: 从 s 至 t 的任何一个流开始, 用流的增值算法寻找流的增值链。如果找出一条从 s 至 t 的流的增值链, 则沿着这条链运送尽可能多的单位流量。然后, 再开始寻找另一条流的增值链, ……。如果再找不出流的增值链, 则算法停止, 因为现有的从 s 至 t 的流已是从 s 至 t 的最大流。根据这一思路, 我们可以将最大流的算法步骤描述如下:

第一步: 令 s 表示发点, t 表示收点。选择任一个从 s 至 t 的起始流, 亦即, 选择满足关系式 (1) — (4) 的 $f(u, v)$ 数值的任一个集合, 如果不知道这样的流, 则对所有 $(u, v) \in E$, 用 $f(u, v) = 0$ 作为起始流。

第二步: 如果 $f(u, v) < m(u, v)$, 令 $i(u, v) = m(u, v) - f(u, v)$, $(u, v) \in I$; 如果 $f(u, v) > 0$, 令 $r(u, v) = f(u, v)$, $(u, v) \in R$ 。

第三步: 对第二步中定义的集合 I 和 R , 执行流的增值算法。如果应用流的增值算法未发现流的增值链, 则算法停止, 现有的流就是最大流。否则, 沿着流的增值算法所发现的流的增值链, 作出最大可能流的增值, 再返回第二步。

例: 在图 7-13 中, 每条边旁的数表示边的容量, 即 $m(u, v)$ 。试求从 s 至 t 的最大流。

第一步: 算法由零流开始, 即对所有的边 $(u, v) \in E$, 令 $f(u, v) = 0$ 。

第二步: 因为对所有的边 $(u, v) \in E$, $f(u, v) < m(u, v)$, 所以每条边都是集合 I 的一个元素, 令 $i(u, v) = m(u, v) - f(u, v) = m(u, v)$ 。因为对所有的边 $(u, v) \in E$, $f(u, v) = 0$, 所以没有边属于 R 。

第三步: 现在, 用流的增值算法寻找一条从 s 至 t 的流的增值链。显然, 可以看出, 从 s 至 t 有几条流的增值链。假设产生一条流的增值链为 (s, a) , (a, b) , (b, t) , 沿着这条链的最大流的增值是 $\min \{i(s, a), i(a, b), i(b, t)\} = \min \{2, 3, 2\} = 2$ 。因而, 在这条链中, 将每一条边增加两个单位, 即 $f(s, a) = 2 + 0 = 2$, $f(a, b) = 2 + 0 = 2$, $f(s, b) = 0$, $f(b, t) = 2 + 0 = 2$, $f(a, c) = 0$, $f(c, t) = 0$ 。以下返回第二步。

第二步：对于上面新产生的流，

$f(s, a) = 2 = m(s, a)$, $(s, a) \notin I$; $(s, a) \in R$, $r(s, a) = 2$ 。
 $f(a, b) = 2 < m(a, b)$, $(a, b) \in I$, $i(a, b) = m(a, b) - f(a, b) = 3 - 2 = 1$;
 $(a, b) \in R$, $r(a, b) = 2$ 。
 $f(b, t) = 2 = m(b, t)$, $(b, t) \in I$; $(b, t) \in R$, $r(b, t) = 2$ 。
 $f(s, b) = 0 < m(s, b)$, $(s, b) \in I$, $i(s, b) = m(s, b) - f(s, b) = 3 - 0 = 3$;
 $(s, b) \in R$ 。
 $f(a, c) = 0 < m(a, c)$, $(a, c) \in I$, $i(a, c) = m(a, c) - f(a, c) = 4 - 0 = 4$;
 $(a, c) \notin R$ 。
 $f(c, t) = 0 < m(c, t)$, $(c, t) \in I$, $i(c, t) = m(c, t) - f(c, t) = 1 - 0 = 1$;
 $(c, t) \notin R$ 。

第三步：再次用流的增值算法寻找另一条流的增值链。算法将产生唯一的流的增值链为 (s, b) , (a, b) , (a, c) , (c, t) 。沿着这条链的最大可能流的增值是 $\min \{i(s, b), r(a, b), i(a, c), i(c, t)\} = \min \{3, 2, 4, 1\} = 1$ 。因此，有一个追加的单位流量沿着这条链从 s 运送至 t 。这条链的每一条前向边 (s, b) , (a, c) , (c, t) 的流增加一个单位，其后向边 (a, b) 的流减少一个单位，最终得到的流是： $f(s, a) = 2$, $f(a, b) = 1$, $f(b, t) = 2$, $f(s, b) = 1$, $f(a, c) = 1$, $f(c, t) = 1$ 。现在，单位流量通过下列路径：

1 个单位从 s 经由边 (s, a) , (a, b) , (b, t) 至 t 。

1 个单位从 s 经由边 (s, b) , (b, t) 至 t 。

1 个单位从 s 经由边 (s, a) , (a, b) , (b, t) 至 t 。以下返回第二步。

第二步：对上面产生的流，

$f(s, a) = 2 = m(s, a)$, $(s, a) \notin I$; $(s, a) \in R$, $r(s, a) = 2$ 。
 $f(a, b) = 1 < m(a, b)$, $(a, b) \in I$, $i(s, a) = m(a, b) - f(a, b) = 3 - 1 = 2$;
 $(a, b) \in R$, $r(a, b) = 1$ 。
 $f(b, t) = 2 = m(b, t)$, $(b, t) \notin I$; $(b, t) \in R$, $r(b, t) = 2$ 。
 $f(s, b) = 1 < m(s, b)$, $(s, b) \in I$, $i(s, b) = m(s, b) - f(s, b) = 3 - 1 = 2$;
 $(s, b) \in R$, $r(s, b) = 1$ 。
 $f(a, c) = 1 < m(a, c)$, $(a, c) \in I$, $i(a, c) = m(a, c) - f(a, c) = 4 - 1 = 3$;
 $(a, c) \in R$, $r(a, c) = 1$ 。
 $f(c, t) = 1 = m(c, t)$, $(c, t) \notin I$; $(c, t) \notin R$, $r(c, t) = 1$ 。

应用流的增值算法，将以上的流作为起始流，我们再也找不到追加的流的增值链 [流的增值算法可以对顶点 s, b, a, c (按此顺序) 着色，但不能对顶点 t 着色]。所以，最大流算法终止。终止流 (如上所述) 就是最大流。在算法终止后，由几条边构成的一个已着色端点和一个未着色端点的截是 (b, t) , (c, t) ，这个截的容量是 $m(b, t) + m(c, t) = 2 + 1 = 3$ ，这就是由最大流算法得到的从 s 运送至 t 的单位数量。

三、最小费用流的算法

在一个有容量的图中，通过每一条边运送一个单位时，都有一个费用，最小费用流问题，就是考虑怎样从发点至收点以最小费用运送一个给定数量 Q 的单位流量。例如，一个制造商欲将成品从工厂运至仓库， he 可以从许多条运输路线中进行选择。在不同的路线上，单位重量成品的运输费用是不同的。每条路线只能担负有限总重量的货物运输。制造商欲将全部成品从工厂运至仓库，用什么方法使运费达到最小？如果用顶点 s 表示工厂，用另一个顶点 t 表示仓库，两条或两条以上路线的每一个交汇点都用一个顶点表示；令每条边的容量等于相应的一段路线所能担负的货物运输的最大重量，并且令每条边的费用等于通过该边所对应的一段路线的单位重量的运费，则制造商的问题就是在这个图上求从 s 至 t 的最小费用流问题。

令 $c(u, v)$ 表示沿着边 (u, v) 运送一个单位流量的费用。首先，我们假设 $c(u, v)$ 都是正整数（后面再对这一假设进行修正）。如前所述，令 $f(u, v)$ 表示通过边 (u, v) 的流的单位数量。当然 $f(u, v) \geq 0$ 。令 Q 表示从发点运送至收点的单位数量。则最小费用流问题可以表示为，

$$\min \left\{ \sum_{(u,v)} c(u, v) f(u, v) \right\} \quad (5)$$

满足约束条件：

$$\sum_v [f(s, v) - f(u, s)] = Q \quad (6)$$

$$\sum_v [f(u, v) - f(v, u)] = 0 \quad (u \neq s, u \neq t) \quad (7)$$

$$\sum_v [f(t, u) - f(v, t)] = -Q \quad (8)$$

$$0 \leq f(u, v) \leq m(u, v) \quad [\text{对所有的 } (u, v) \in E] \quad (9)$$

显然，最小费用流问题也是一个线性规划问题（实际上，最大流问题可视为最小费用流问题。其中所有各条边的费用为零，并且 Q 为最大可能流的数值）。

令 (5) 式由下式代替：

$$\max \left\{ PQ - \sum_{(u,v)} (u, v) f(u, v) \right\} \quad (10)$$

(10) 式中， P 是任一个大数，例如，大于从 s 至 t 运送一个单位的最大总费用的数。如果将 P 解释为从 s 至 t 运送一个单位所得的利润，则 (10) 式可解释为扣除运费以后的最大可能净利润。因而，按照这一解释，使 (10) 式最大的任一个流也将使 (5) 式最小，反之亦然。

最小费用流的算法，首先从 s 至 t 运送尽可能多的全程每单位总费用为零的单位。其次，最小费用流算法从 s 至 t 运送尽可能多的全程每单位总增加费用为 1 的单位，……。当总数为 Q 的单位已从 s 运送至 t 时，或者，当从 s 至 t 没有更多的单位可以运送时，算法终止。换句话说，算法首先对 $P=0$ 解 (6) — (10) 的问题，其次对 $P=1$ 解题，再次对 $P=2$ ，……。

设算法已从 s 至 t 运送尽可能多的总增加费用为 $P-1$ 或稍少的单位。算法怎样确定从 s 至 t 运送每单位总增加费用为 P 的单位流量呢？为做到这一步，算法必须找出一条从 s 至 t 的增值链，这条链具有这样的性质：“沿着这条链”运送一个单位，总的增加费用等于 P 。为了有助于明了这一思路，考虑图 7-14。在这个图中，从 s 至 t 运送一个单位流量的费用最小的路线是

沿着路 (s, a) , (a, b) , (b, t) 进行。单位流量所担负的总费用为 $2+1+2=5$, 如果 $Q=1$, 就构成了一个最小费用流。

假设 $Q=2$, 根据检查结果, 剩下仅有的增值链为 (s, b) , (a, b) , (a, t) , 这条链可以担负运送从 s 至 t 的一个追加单位。如果第二个单位流量从 s 出发, 则第一个单位流量将在顶点 a 处改变方向, 取路 (s, a) , (a, t) 。第二个单位流量将在顶点 b 处代替第一个单位流量, 取路 (s, b) , (b, t) 。这两个单位流量的总费用为 $c(s, a)+c(a, t)+c(s, b)+c(b, t)=2+4+4+2=12$, 比原先的费用增加了 7。所增加的 7 是从流的增值链 (s, b) , (a, b) , (a, t) 中按下列方式形成的:

在前进方向, 通过边 (s, b) , 形成+4 的费用;

在后退方向, 通过边 (a, b) , 形成-1 的费用;

在前进方向, 通过边 (a, t) , 形成+4 的费用。

因此, 算法必然要找出一条流的增值链, 这条链具有这样的性质: 链的前向边费用之和减去后向边费用之和等于 P 。

要完成这项计算, 须对图中每一个顶点 u 指定一个整数 $P(u)$, 这些顶点数值 $P(u)$ 具有这样的性质: $P(s)=0$, $P(t)=P$, 对所有各顶点 $u \neq s, u \neq t$, $0 \leq P(u) \leq P$ 。算法仅沿着边 (u, v) 改变流量, 对于边 (u, v) , 有

$$P(v)-P(u)=c(u, v) \quad (11)$$

如果算法找出一条从 s 至 t , 由所有满足方程 (11) 的边组成的流的增值链, 则从 s 沿着这条链运送至 t 的每个单位的总增加费用等于 P 。以此为依据, 最小费用流的算法步骤如下: 第一步 (开始): 首先令每条边 (u, v) 中的流 $f(u, v)=0$, 对所有各顶点 u , 令 $P(u)=0$ 。第二步 (决定哪些边可以改变流量): 令 I 为满足条件:

$$\begin{aligned} &P(v)-P(u)=c(u, v) \\ \text{及} &f(u, v) < m(u, v) \end{aligned}$$

的所有边组成的集合。

令 R 为满足条件:

$$\begin{aligned} &P(v)-P(u)=c(u, v) \\ \text{及} &0 < f(u, v) \end{aligned}$$

的所有边组成的集合。

令 N 为集合 $E-I \cup R$ (I 与 R 中的那些边仅为可以考虑改变流量的那些边。因而, 只能在满足方程 (11) 的边上改变流量)。

第三步 (改变流量): 按上述第二步定义的 I , R 和 N , 执行最大流算法。当从 s 已经运送总数为 Q 的单位流量至 t 时, 或者, 对 I 、 R 、 N 集合的现有组成部分来说, 当从 s 至 t 已不可能再运送更多的流时, 算法停止。如果先发生前一种情况, 因为最终流就是从 s 运送 Q 单位至 t 的最小费用流, 所以算法停止。如果先发生后一种情况, 检查一下, 现有的流是否为从 s 至 t 的最大流 (检查方法: 验证用流的增值算法最后着色所产生的截是否已经饱和)。如果就是最大流, 因为从 s 至 t 已不能再运送更多的单位流量, 并且最终流就是最小费用流, 所以算法停止, 如果不是最大流, 则转向第四步。

第四步 (改变顶点数值): 考虑流的增值算法所作的最后着色 (流的增值算法是最大流算法中的一个子程序, 而最大流算法又是最小费用流算法的一

个子程序)。将每一个未着色顶点 u 的顶点数值 $P(u)$ 增加+1 (注意, t 尚未着色, 否则就已发现一条流的增值链, 所以对 $P(t)$ 增加+1)。返回第二步。

例：在图 7-15 中, 每条边 (u, v) 旁边的第一个数字表示边的费用 $c(u, v)$, 第二个数字表示边的容量 $m(u, v)$ 。试用最小费用流算法找出最小费用流。

开始时, 所有顶点数值都是零, 并且所有顶点除发点外都未着色, 发点 s 总是已着色的。迭代计算的结果分别列表如下 (见表 7-2)。

从表 7-2 知道, 顶点 t 已经着色。从 s 至 t 沿着路 $(s, a), (a, b), (b, t)$ 可以运送 2 个单位流量。因此, $f(s, a)=2, f(a, b)=2, f(b, t)=2$ 。继续迭代, 计算结果列于表 7-3。

从表 7-3 可以看出, 顶点 t 已经着色。从 s 至 t 沿着链 $(s, b), (a, b), (a, t)$ 可以运送 1 个单位流量。因此, $f(s, a)=2, f(s, b)=1, f(a, b)=1, f(a, t)=1, f(b, t)=2$, 继续迭代, 计算结果见表 7-4。

因为从已经着色顶点至未着色顶点的边 [边 $(s, a), (s, b)$] 已经饱和, 所以从 s 至 t 不能再

表 7-2 最小费用流计算表 (之一)

迭代	$P(s)$	$P(a)$	$P(b)$	$P(t)$	着色边	着色顶点
0	0	0	0	0	无	s
1	0	1	1	1	(s, a)	s, a
2	0	1	2	2	$(s, a)(a, b)$	s, a, b
3	0	1	2	3	$(s, a)(a, b)(b, t)$	$s, a, b, t,$

表 7-3 最小费用流计算表 (之二)

迭代	$P(s)$	$P(a)$	$P(b)$	$P(t)$	着色边	着色顶点
3	0	1	2	3	无	s
4	0	2	3	4	$(s, b)(a, b)$	s, a, b
5	0	2	3	5	$(s, b)(a, b)(a, t)$	s, a, b, t

表 7-4 最小费用流计算表 (之三)

迭代	$P(s)$	$P(a)$	$P(b)$	$P(t)$	着色边	着色顶点
5	0	2	3	5	无	s

运送追加的单位流量。因此, 现有的三个单位的流就是具有最小费用的最大流。

到现在为止, 以上我们的讨论都是假定边的费用 $c(u, v)$ 为正整数。以下我们来修正最小费用流的算法, 以适应具有非整数边的费用。

因为边的费用均假设为正整数, 所以从 s 至 t 所运送的所有单位流量担

负的总增加费用是正整数。因此，算法只须检验 P 的整数值，并且顶点数值总是增加整数+1。

当边的费用无需为整数时，则从 s 至 t 的流增值路的总费用不一定是整数，因而，算法必须检验不是整数的 P 值。算法应该检验 P 的哪些值呢？算法应该作出顶点数值的什么样的增量呢？

设算法对 $P=P(t)$ 的某些值已经执行过。某些边将有一个已着色端点和一个未着色端点。我们必须考虑以下两种情形：

第一种情形：如果顶点 u 已经着色，而顶点 v 未着色，则仅当 $(u, v) \in I$ 且 $P(v)$ 可以改变，使 $P(v)-P(u)=c(u, v)$ 时，边 (u, v) 是可供选择为着色的边。如果 $(u, v) \in I$ ，则 $P(v)-P(u) \leq c(u, v)$ 。因而在边 (u, v) 可以着色之前， $P(v)$ 必须增加 $c(u, v)-P(v)+P(u)$ 个单位。令 $\delta(u, v)=c(u, v)-P(v)+P(u)$ 。

第二种情形：如果顶点 v 已着色，而顶点 u 未着色，则仅当 $(u, v) \in R$ 且 $P(u)$ 可以改变使 $P(v)-P(u)=c(u, v)$ 时，边 (u, v) 是可供选择为着色的边。如果 $(u, v) \in R$ ，则 $P(v)-P(u) \geq c(u, v)$ 。因而，在边 (u, v) 可以着色之前， $P(u)$ 必须增加 $P(v)-P(u)-c(u, v)$ 个单位。令 $\delta(u, v)=P(v)-P(u)-c(u, v)$ 。

如果边 (u, v) 不适合于第一种情形或第二种情形，则令 $\delta(u, v) = \infty$ 。令 $\delta = \min_{(u,v)} \{ \delta(u, v) \} \quad (u, v) > 0$ 。

因此，如果每个未着色顶点 u 的对偶的数值 $P(u)$ 增加 δ ，则至少将有另外一条边可以着色，并且也许将发现一条新的流的增值链，所以， $P(t)$ 的值将增加 δ 。同样地，为保证至少有另外一条边可以着色，需要最小的增量，确定这个增量，就可以算出以后各个顶点数值的增量。如果 $\delta = \infty$ ，则没有另外的边可以着色，在这种情形时，现有的流就是最大流。于是，由修正顶点数值的增量过程，最小费用流算法就可以适用于非整数的边的费用。

修正的算法是否可以在有限个步骤以后结束呢？答案是肯定的。因为 $P(t)$ 必然总是等于从 s 至 t 沿着某一条链的各条边的费用的总和，而且，因为从 s 至 t 仅存在有限条不同的链，从而在修正算法中， $P(t)$ 只能取有限个不同的数值。因此，修正的算法必然仅在有限个步骤以后结束。

第八章 控制论及其应用

控制论，是自本世纪 40 年代以来形成的一门新兴学科。它是在数学、计算机、通信技术及物理学、电子学、生物学、生理学、社会学、行为科学等学科相互渗透、高度综合的基础上形成的。如果说物理学、电子学、生物学、生理学、社会学等学科是从各个不同的方面研究那些同类或彼此相近的现象，那么控制论则是从一个既定的方面研究不同客体所共同的控制规律，而这些客体可能分属于物理学、电子学、生物学、生理学、社会学等各种不同的学科。控制论的特征是利用不同客体之间的类比以及广泛地运用数学方法，其许多方法具有跨学科和普遍适用的性质。本章，我们拟结合有关实例，介绍和探讨控制论方法在地理学研究中的应用问题。

第一节 控制论概述

一、控制论的发展

1948 年,维纳 (N.Wiener)所著的《动物和机器中的控制与通信》一书的出版,标志着控制论的正式诞生,迄今已有 40 多年的历史。在这 40 多年中,控制论的发展大致经历了三个时期,即 40 年代末期到 50 年代末期的经典控制论时期,60 年代初期到 70 年代初期的现代控制论时期及 70 年代中期到现在的大系统理论时期。

(一)经典控制论时期

40 年代末期,在工业生产和武器装备方面,已经开始采用简单的自动控制系统。例如,锅炉水位的自动控制,蒸气温度、水轮机转速、发电机电压和频率、电动机转速等的自动控制,高射炮的自动跟踪装备,军用飞机和舰船的自动驾驶仪等。

40 年代末期到 50 年代末期是控制论发展的初期阶段。这一时期,控制论所研究的控制系统都比较简单,且多属于工程方面,单输入单输出的线性系统是其主要研究对象。对于系统动态过程及动态特性的描述,主要是采用微分方程(或差分方程)与传递函数方法。该时期的主要成就是建立了系统、信息、控制、反馈、稳定性、“黑箱”等基本概念与分析方法,为控制论的进一步发展打下了基础。但是,在现实世界中,能用微分方程(或差分方程)与传递函数方法描述的系统很有限。因此,经典控制论的应用范围受到了一定的限制。

(二)现代控制论时期

到了 60 年代初期,随着导弹、人造卫星、航天工程、高能物理、电子计算机等科学技术的迅猛发展,控制论已从经典控制论发展成为现代控制论。

60 年代初到 70 年代初,现代控制论具有以下几个特点:

1)系统描述的状态空间化。在经典控制论时期,对于系统动态过程及动态特性的描述,主要采用微分方程(或差分方程)与传递函数方法。而到了现代控制论时期,则引进了状态和状态空间的概念,将描述系统动态过程及动态特性的数学模型表示成了状态空间模型。系统描述的状态空间化是现代控制论的最大特点。将系统的数学模型表示成状态空间模型,使问题的处理大为简化,也便于运用计算机求解。

2)系统的广义化。经典控制论主要研究单输入单输出的线性系统。而现代控制论研究的系统不仅局限于单输入单输出的线性系统,而且对于多输入多输出的非线性系统等更为复杂的系统,从理论上来说都能处理,这就使它的应用范围扩大了很多。

3)理论分析的计算机化。由于系统的复杂性使相应的理论分析日趋复杂。因此,现代控制论的许多分析方法往往着眼于如何利用计算机,使理论分析更有效、更可靠。

4)性能(目标函数)最优化。经典控制论的设计方法常采用试凑法,而现代控制论则引进了“性能指标”(或目标函数)的概念,在满足一定约束条件的前提下,寻求一个最优控制,使性能指标取最优值。在这一最优控制的作用下,系统能“动态最优地”达到预期的目标。

(三)大系统理论时期

到了 70 年代中期,控制论的研究对象发展到更加复杂的大系统。

对于大系统来说，其变量众多，结构更加庞大而复杂，所以现代控制论的状态空间分析法已经不能完全适应了。处理大系统的方法主要有：分解-协调原理；分散最优控制；多级递阶控制；大系统模型降阶理论；向量李雅普诺夫稳定性理论等。但是，目前，对于大系统的分析和研究还没有一个统一的办法，正处在发展阶段。

随着大系统理论的发展，控制论也已不仅仅是研究动物和机器的控制与通信问题，而进入到政治、经济、军事、环境、医学、社会学、心理学，甚至哲学、美学和艺术等各个领域。目前，控制论已经成为一门具有十分广泛的适用性的方法论学科。

二、控制论的基本概念

系统、控制、信息、反馈、输入、输出、状态等是控制论的基本概念。控制论就是阐述这些概念的定义、关系、数学描述方法及其在各种具体问题中的应用的科学。

(一)系统

系统，是控制论的一个基本概念。但是，在早期的控制论著作中，如在维纳的《控制论》中，作者只是对系统的主要特征作了描述，并没有给出系统明确的定义。随着学科的发展，控制论学者认识到必须给出系统确切的定义。基于这种认识，许多控制论专家都曾从不同的角度提出了系统的定义。如，我国著名的控制论专家钱学森先生曾在《工程控制论》一书修订版序言中明确提出控制论的研究对象就是系统，他指出：“所谓系统，是由相互制约的各个部分组织成的具有一定功能的整体。”这一定义揭示了系统的一些重要特性：系统是一个整体，整体是由部分组成的，各个组成部分之间相互联系、相互制约，整体表现出一定的功能。

在控制论中，特别强调系统的功能，它认为对系统控制的目的在于获取特定的功能。为了实现系统的功能，需要对系统的各个构成部分进行组织，把控制与组织联系起来。

控制论所研究的系统都处在一定的环境中。从其外部规定性来看，它是按研究者所关心的问题而从错综复杂、相互联系的事物中相对孤立出来作为研究对象的一部分事物。一般而言，控制论所研究的这种具有控制作用的相对孤立系统，是由两个功能不同的子系统，即控制系统与受控系统组成的。

从控制论的角度来看，地理学的研究对象——地理环境与人类活动相互作用的地理系统，实际上是一个地理调控系统，其调控系统、受控系统与环境之间的相互关系可以用图 8-1 来表示。

在图 8-1 中，调控系统是指对人类活动具有支配和控制作用的社会经济系统或其某些组成部分，它是地理调控系统的主体，支配人类活动的决策行为从这里发出，并通过有关信息传递和执行环节产生行动，使人类活动作用于受控系统——地理环境或其某些组成部分。这里，受控系统（地理环境或其某些组成部分）是系统的客体，是被调控对象。人类活动的目的就是企图通过对其所生存的地理环境的调控，以便为人类的生存和发展提供更多的产品（包括物质产品、精神产品、文化产品等）。而图 8-1 中的环境则是一个不同于地理环境的概念，它是相对于研究对象——地理调控系统而言的外部环境。如果我们研究某一地理系统的调控问题，则其它地理系统就可能成为这一系统的环境或者环境的组成部分。以农田生态经济调控系统为例，其调控

系统就是农田生产系统，受控系统就是农田土壤系统，环境则不但包括了农作物生长发育的非生物环境（包括光、水、热等在内的气候环境及地貌条件），而且还包括其它社会经济环境。这一调控系统的运行过程就是在农田生产系统的调控下，通过人类的生产活动将一组来自于环境的输入要素（包括光、水、热等自然要素的输入，以及籽种、肥料、农药、电力、农用机械等输入要素）转化为农产品输出给社会的过程

（二）控制

控制，是控制论的中心概念。控制论就是研究各种不同系统所共同具有的控制规律的科学。凡控制总要涉及到施控者（控制系统）和受控者（受控系统）两种实体，它是施控者作用于受控者的一种主动行为。领导、指挥、管理、教育、设计，这些都是主动的控制行为，即施控者对受控者施加作用的行为。控制是有目的的，它能导致受控对象发生合乎目的的变化。例如，企业管理的目的是为了提高经济效益，增强企业活力；教育的目的是为了使受教育者增长知识，掌握技能。控制的目的是根据受控对象的功能输出来衡量。某控制系统的目的如果是一个，则称为单目标控制系统；如果是多个，则称为多目标控制系统。

在控制论研究中，受控对象一般都有多种可能的行为（只有一种可能行为对象无控制的必要），其中，有些合乎目标，有些则与目标相悖。控制就是在多种可能的行为中选择一个最合乎目标的行为，避免不合乎目标的行为发生。对于施控者而言，应有多种手段（只有一种手段无控制的可能），不同手段作用于受控对象有不同的效果。控制就是在多种手段中选择最有效的手段，以达到最佳的控制效果。

综上所述，所谓控制就是施控者选择适当的手段作用于受控者，以引起受控者的行为发生预期变化的一种策略性的主动行为。

（三）输入和输出

控制论既强调系统与环境之间的明确界限，也强调系统与环境之间的相互联系和作用。输入与输出就是描述系统与环境之间相互联系和作用的概念。

在控制论中，环境对系统的作用称为输入，而系统对环境的作用称为输出。其中，输入又可以进一步分成两类，一类是体现控制目标的控制作用，记作 u ；一类是有碍于控制目标实现的干扰作用，记作 m 。如果将系统的输出记作 y ，则可以将一个控制系统抽象地表示为图 8-2 所示的形式。

用输入和输出描述系统，在控制论中具有基本的方法论意义。在这种描述下，我们可以把控制理解为选择适当的输入以获得预期输出的操作。这样，控制问题的研究可以归结为系统的输入-输出关系的研究。从这种观点出发，一些控制论学者曾把系统定义为“从输入到输出的变换器”或“由一些分别表示输入和输出的时间函数的有序对来描述的一种抽象的对象”。

输入是环境对系统的激励，输出是系统对输入激励的响应。输出对输入的响应特性，表现了系统的基本特性。如果将输入空间记为 U ，输出空间记为 Y ，则系统的输出对输入的响应特性可定义为从 U 到 Y 的映射

$$f: U \rightarrow Y$$

映射 f 表示输入与输出之间的因果关系。这种关系通常都具有某种不确定性、不确知性。存在不确定性才使控制成为必要和可能。控制的意义正在

于使系统在不确定的条件下达到比较确定的目标。控制论所研究的就是如何描述这种不确定性，寻找处理这种不确定性的控制手段。因此，维纳在《控制论》一书中强调控制问题不属于牛顿决定论的范畴，而属于统计力学范畴。

(四)状态

输入-输出关系描述的是系统的外部特性。但是，为了更有效地控制系统，还需要了解其内部特性。描述系统内部特性的是状态概念。这一概念较为精确的定义是：系统的现有状态可以看成是，当所有现在和未来时刻的输入为已知时，为完全描述系统将来的行为所需要的关于“过去”的最少量信息。借助于状态概念，现代控制论获得了对系统概念的深刻认识。

(五)反馈

把系统受上一步控制作用而产生的效果(输出)作为决定对系统下一步如何控制(输入)的依据，这种行为或策略称为反馈。反馈量被用来加强控制量对系统的作用，称为正反馈；反馈量被用来抵消控制量对系统的作用，称为负反馈。反馈代表一种控制原理，它能为各类控制系统的运行机制作出科学的解释，特别有助于人们对活性机体和社会组织中的许多现象的理解。反馈是一种技术方法，它是控制论方法宝库中极具特色的一种，在工程、经济、政治、生物、环境等问题的控制研究中有着广泛的应用。

(六)信息

一个具体的控制过程可能是物理的、生理的或社会的，但是贯穿于一切控制过程的共同本质是信息的获取、加工、传输、存储和利用。借助于信息的运动过程，可以把空间上相互分立或时间上前后相继的不同环节联接成为一个功能整体，对控制目标和手段进行选择，对系统各个构成部分进行组织。与没有控制机制的物质系统相比，控制系统的显著特点是，巨大物质的运动与巨大能量的传输变换，可以通过携带信息的能量不大的信号来指挥和控制。信息是控制论最基本的概念之一。

第二节 地理系统动态的控制论描述

地理系统是一类动态系统，它的状态是在不断地变化的。对于这样一类动态系统，我们所关心的问题之一是其动态过程是怎样进行的？其未来状态是否可以预测？这一过程是否可以调控？这些问题的回答都依赖于人们对地理系统的认识程度，而对地理系统动态过程的描述就是人们对其认识程度的一种反映。本节，我们将从控制论角度出发，探讨地理系统动态过程的描述问题。

一、传递函数法

(一)地理系统动态的框图模型

根据经典控制论，为了研究上的方便，对于一个复杂的地理系统，我们可以撇开它们与系统调控过程无关的具体内容，把整个系统划分为许多的控制环节，每一个环节都代表一个实际过程，然后再将每一个环节按其相应过程的输入-输出关系联接起来，这样就得到一个实际系统简化了的框图模型。譬如，对于农田生态系统，基于目前人们认识世界和改造自然能力的限制，人类对它的调控还只能局限于播种、施肥、灌溉、使用农药等措施上。因此，对于农田生态系统这一简单的地理调控系统，通过对其每一实际过程与输入-输出环节的分析，我们可以将其抽象为如图 8-3 的框图模型。

在图 8-3 中， y 为农田生态系统的经济收益； x_1 为籽种播种量； x_2 为肥料施用量； x_3 为灌溉用水量； x_4 的农药使用量； x_5 为籽种投资； x_6 为肥料投资； x_7 为灌溉投资； x_8 为农药投资； u 为初始投资； $W_i(s)$ ($i=1, 2, \dots, 8$) 为各环节的传递函数； $E(s)$ 为反馈环节的传递函数。

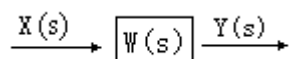
图 8-3 表明，农田生态系统的经济收益（系统输出）决定于籽种、肥料、灌溉用水、农药使用量的输入；而籽种、肥料、灌溉用水、农药使用量等作为输出，又决定于各自的投资输入。在一定的输入（初始投资）下，经过系统的各运行环节就会得到一定的输出（经济收益），然后又可以把输出引向输入端，将部分收益再用于投资。这样就构成了农田生态系统的反馈控制过程。在这个框图模型中，传递函数的功能是描述各输入-输出环节的动态关系。

(二)地理系统动态的传递函数

根据经典控制理论，对线性定常系统的动态特性的研究，可以用传递函数来描述系统的输入和输出之间的关系。所谓某环节的传递函数，就是在零初始条件下，该环节输出量（或称响应函数）的拉普拉斯（Laplace）变换 $Y(s)$ 与输入量（或称驱动函数）的拉普拉斯变换 $X(s)$ 之比，即

$$W(s) = \frac{Y(s)}{X(s)} \quad (1)$$

其框图模型描述如下：



传递函数的概念一般只适用于线性定常系统，然而它也可以被扩充到一些非线性系统的研究之中。

这里需要说明的是，系统动态的微分方程描述与传递函数描述是一致的。在进行系统动态特征分析时，为了避免复杂的微分方程的求解工作，常

常需要将系统的微分方程描述形式转化为传递函数描述形式。假设某系统的某输入-输出环节的微分方程描述为

$$\begin{aligned} & \frac{d^n Y}{dt^n} + a_{n-1} \frac{d^{n-1} Y}{dt^{n-1}} + \cdots + a_1 \frac{dY}{dt} + a_0 Y \\ &= b_0 X + b_1 \frac{dX}{dt} + \cdots + b_{m-1} \frac{d^{m-1} X}{dt^{m-1}} + b_m \frac{d^m X}{dt^m} \end{aligned} \quad (2)$$

在 (2) 式中, Y 代表输出, X 代表输入。

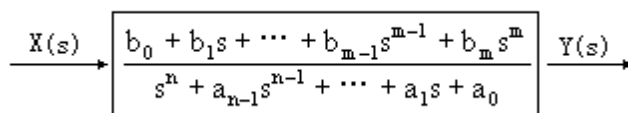
在零初始条件下, 对方程 (2) 式两端作拉普拉斯变换得:

$$\begin{aligned} & (s^n + a_{n-1}s^{n-1} + \cdots + a_1s + a_0)Y(s) \\ &= (b_0 + b_1s + \cdots + b_{m-1}s^{m-1} + b_ms^m)X(s) \end{aligned}$$

由上式容易得到这一环节的传递函数:

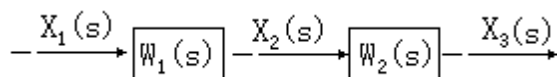
$$W(s) = \frac{Y(s)}{X(s)} = \frac{b_0 + b_1s + \cdots + b_{m-1}s^{m-1} + b_ms^m}{s^n + a_{n-1}s^{n-1} + \cdots + a_1s + a_0} \quad (3)$$

其框图模型描述如下:

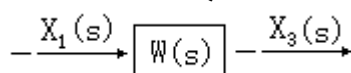


复杂的地理系统, 往往包含着若干个输入-输出环节, 各个环节以不同的方式联接就构成了复杂的系统动态过程。为了分析整个系统的动态过程, 下面我们来讨论传递函数的合成运算。

1. 串联系统 如果一个系统的动态过程由两个输入-输出环节构成, 第一个环节的输出又是第二个环节的输入, 则这种联接方式构成了如下的串联系统:



如果记整个系统的传递函数为 W(s), 则上述框图可被简化为:



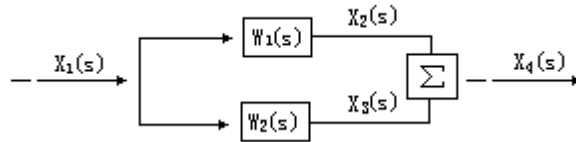
那么

$$W(s) = \frac{W_3(s)}{X_1(s)} = \frac{X_3(s)}{X_2(s)} \cdot \frac{X_2(s)}{X_1(s)} = W_1(s) \cdot W_2(s) \quad (4)$$

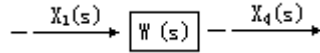
即系统传递函数等于两个环节传递函数的乘积。同样的道理可以证明, 由 n 个输入-输出环节构成的串联系统, 其传递函数等于各个环节传递函数的乘积, 即

$$W(s) = W_1(s) \cdot W_2(s) \cdots W_n(s) = \prod_{i=1}^n W_i(s) \quad (5)$$

2. 并联系统 所谓并联系统, 是指流入系统的输入分别进入了不同的输入-输出环节, 由各个环节的分别作用, 共同决定系统的输出, 即如下框图模型所示:



如果记系统的传递函数为 $W(s)$ ，则上述框图就被简化为：



$$\begin{aligned} \text{由于 } X_4(s) &= X_2(s) + X_3(s) \\ &= W_1(s)X_1(s) + W_2(s)X_1(s) \\ &= [W_1(s) + W_2(s)] X_1(s) \end{aligned}$$

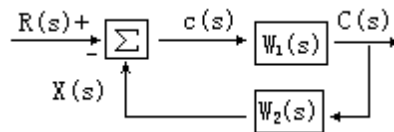
所以系统的传递函数为：

$$W(s) = \frac{X_4(s)}{X_1(s)} = W_1(s) + W_2(s) \quad (6)$$

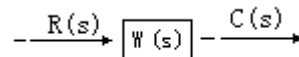
同理可以证明，由 n 个输入-输出环节构成的并联系统，其传递函数等于各个环节传递函数之和，即

$$W(s) = W_1(s) + W_2(s) + \cdots + W_n(s) = \sum_{i=1}^n W_i(s) \quad (7)$$

3. 反馈系统 所谓反馈系统，就是把系统的输出再引向输入端，从而使系统的输入发生改变的一类控制系统，即如下图所示：



其中， $W_1(s)$ 为系统输入-输出环节的传递函数， $W_2(s)$ 为反馈环节的传递函数。如果记整个反馈控制系统的传递函数为 $W(s)$ ，则以上框图可被简化为



对于反馈控制系统，又可以根据不同的反馈机制，将它分为正反馈和负反馈两种情形。

(1) 正反馈

由于

$$\begin{cases} C(s) = R(s) + X(s) \\ X(s) = C(s) \cdot W_2(s) \\ C(s) = e(s) \cdot W_1(s) \end{cases}$$

因而

$$\begin{aligned} R(s) &= e(s) - X(s) = e(s) - C(s) \cdot W_2(s) \\ &= e(s) - e(s) \cdot W_1(s) \cdot W_2(s) \\ &= [1 - W_1(s)W_2(s)]e(s) \end{aligned}$$

所以，系统的传递函数为

$$\begin{aligned}
 W(s) &= \frac{C(s)}{R(s)} = \frac{e(s)W_1(s)}{[1 - W_1(s) \cdot W_2(s)] \cdot e(s)} \\
 &= \frac{W_1(s)}{1 - W_1(s) \cdot W_2(s)}
 \end{aligned} \tag{8}$$

(2) 负反馈

由于

$$\begin{cases} e(s) = R(s) - X(s) \\ X(s) = C(s)W_2(s) \\ C(s) = e(s)W_1(s) \end{cases}$$

因而，

$$\begin{aligned}
 R(s) &= e(s) + X(s) = e(s) + C(s)W_2(s) \\
 &= e(s) + e(s)W_1(s)W_2(s) \\
 &= [1 + W_1(s)W_2(s)]e(s)
 \end{aligned}$$

所以，系统的传递函数为

$$\begin{aligned}
 W(s) &= \frac{C(s)}{R(s)} = \frac{W_1(s)e(s)}{[1 + W_1(s)W_2(s)]e(s)} \\
 &= \frac{W_1(s)}{1 + W_1(s)W_2(s)}
 \end{aligned} \tag{9}$$

对于图 8-3 描述的农田生态系统，通过运用以上传递函数的三种运算法则，经过合成运算可以得到系统的传递函数：

$$W(s) = \frac{[W_1(s)W_5(s) + W_2(s)W_6(s) + W_3(s)W_7(s) + W_4(s)W_8(s)]}{1 - E(s)[W_1(s)W_5(s) + W_2(s)W_6(s) + W_3(s)W_7(s) + W_4(s)W_8(s)]}$$

任何一个地理系统，其复杂的动态关系都是以上三种关系（串联、并联、反馈）的组合与合成。因此有了以上传递函数的运算法则，就可以由各个环节的传递函数经过合成运算得到整个系统的传递函数。通过对系统传递函数的分析，又可以把握地理系统的动态特征，由此就可以对系统进行未来发展趋势的预测，从而为系统的调控提供决策依据。

二、状态空间法

上面我们介绍了传递函数在定常的地理系统动态分析中的应用。但许多复杂的地理系统可能是多输入、多输出，也可能是时变的系统。对这类性能指标要求严格的复杂系统，经典控制论的传递函数就显得无能为力了。本世纪 60 年代发展起来的现代控制论则可用于多输入、多输出的，线性的或者非线性的，定常的或者时变的系统。在现代控制论中，描述系统动态的基本方法就是状态空间法，描述系统的数学模型就是状态方程。

所谓状态空间是指状态变量构成的欧氏空间，它包含了描述系统的状态变量的一个最小集合。状态空间有时也被称为相空间，系统的状态随时间变化在状态空间中形成的运动轨迹被称为相轨线。为了描述系统在状态空间中的动态行为，需要状态方程（即输入对状态变量的作用关系式）和输出方程（或称为观察方程，即输出和状态变量的关系式）。一般地，一个系统的状态空间描述可以写成：

$$\begin{cases} \frac{dx}{dt} = F(x, u) \\ Y(t) = G(x, u) \end{cases} \quad (10)$$

在 (10) 式中, x 为状态变量向量, u 为输入 (或控制) 变量向量, Y 为输出变量向量。对于线性系统来说, 其状态空间描述为

$$\begin{cases} \dot{x} = A(t)x + B(t)u \\ Y = C(t)x + D(t)u \end{cases} \quad (11)$$

(11) 式中, $A(t)$ 为状态矩阵, $B(t)$ 为输入矩阵, $C(t)$ 为输出矩阵, $D(t)$ 为输入矩阵。因为矩阵 $A(t)$ 、 $B(t)$ 、 $C(t)$ 、 $D(t)$ 就完全决定了系统的动态特征, 因而方程 (11) 式表示的系统也可被记为形如 $[A(t), B(t), C(t), D(t)]$ 的四联矩阵。当 $A(t)=A$, $B(t)=B$, $C(t)=C$, $D(t)=D$ 都为常系数矩阵时, $[A, B, C, D]$ 就表示了一个线性定常系统。事实上, 由微分方程或传递函数表示的线性定常系统都可以转化为状态空间的表示形式。设有某一连续的线性定常系统, 其微分方程描述为

$$y^{(n)} + a_1 y^{(n-1)} + a_2 y^{(n-2)} + \dots + a_{n-1} y + a_n y = u \quad (12)$$

(12) 式中, u 为输入变量, y 为输出变量, $y^{(i)}$ ($i=1, 2, \dots, n$) 为输出变量 y 的 i 阶导数, a_i ($i=1, 2, \dots, n$) 为常系数。下面我们将它转化为状态空间描述形式。

在上述微分方程描述的系统中, 若已知初始变量 $y(0)$ 及其各阶导数 $y^{(1)}(0)$, $y^{(2)}(0)$, $\dots$, $y^{(n-1)}(0)$, 并且已知 $t \geq 0$ 的输入变量 $u(t)$, 则可以完全确定系统的未来行为。因此, 我们可以取

$$X = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ M \\ x_n \end{pmatrix} = \begin{pmatrix} y \\ y^{(1)} \\ \vdots \\ M \\ y^{(n-1)} \end{pmatrix}$$

作为状态变量向量。将状态向量的每一个分量代入系统的微分方程中就得

$$\begin{cases} \dot{x}_1 = x_2 \\ \vdots \\ \dot{x}_2 = x_3 \\ \vdots \\ \dot{x}_3 = x_4 \\ \vdots \\ \dot{x}_{n-1} = y^{(n-1)} = x_n \\ \dot{x}_n = y^{(n)} = u - a_1 y^{(n-1)} - a_2 y^{(n-2)} - \dots - a_{n-1} y - a_n y \\ \quad = u - a_1 x_n - a_2 x_{n-1} - \dots - a_{n-1} x_2 - a_n x_1 \end{cases}$$

所以, 只要取

$$A = \begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 \\ -a_n & -a_{n-1} & -a_{n-2} & -a_1 \end{pmatrix} \quad B = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ M \\ 0 \\ 1 \end{pmatrix}$$

就得到系统的状态方程：

$$\dot{x} = Ax + Bu \quad (13)$$

显然，系统的输出方程为

$$y = x_1$$

因此，只要取 $C = [1, 0, 0, \dots, 0]$ ，就可以得到其矩阵形式：

$$y = Cx \quad (14)$$

这样，联立 (13) 与 (14) 式就得到系统 (12) 式的状态空间描述：

$$\begin{cases} \dot{x} = Ax + Bu \\ y = Cx \end{cases} \quad (15)$$

以上我们所讨论的是连续系统情形，但是在地理学研究中，有关信息的获取总是在离散的时间间隔上进行的，所以离散系统是经常可见的。对于离散系统，其状态空间的描述可以表述为

$$\begin{cases} X(t+1) = \Phi(t+1, t)X(t) + F(t)U(t) \\ Y(t) = H(t)X(t) + G(t)U(t) \end{cases} \quad (16)$$

$$\text{或者} \quad \begin{cases} X(t+1) = \Phi(t+1, t)X(t) + F(t)U(t+1) \\ Y(t) = H(t)X(t) + G(t)U(t) \end{cases} \quad (17)$$

上式各矩阵与连续情形的叫法一样。系统 (16) 式与系统 (17) 式的差别在于 (17) 式的输入是同时的，而 (16) 式的输入是有时间延迟的。在地理调控系统中，延迟是不可避免的，譬如，生产投资过程、环境治理过程等都是有时间延迟的过程。当 $\Phi(t+1, t)$ ， $F(t)$ ， $H(t)$ 及 $G(t)$ 与时间无关时，上述系统就称作离散的定常系统。

由以上分析可以看出，要建立某一系统的状态空间模型，首先需要确定系统的状态变量、输入变量(或控制变量)及输出变量，然后再建立输入变量对状态变量的作用关系及输出变量和状态变量的关系。若有可能的话，将状态变量取作输出变量，则状态方程与输出方程就同为一个方程，这样对系统建模将带来极大的方便。譬如，考虑某地区的人口动态系统，假设该区域由几个小区域组成，如果记 $x_i(t)$ ($i=1, 2, \dots, n$) 为第 i 个小区域在第 t 时刻(第 t 年)的人口数量，那么向量

$$X_{(t)} = \begin{pmatrix} x_1(t) \\ x_2(t) \\ \vdots \\ M \\ x_n(t) \end{pmatrix} \quad (18)$$

既可以作为该区域人口系统的状态变量，又可以作为其输出变量。显然，该区域人口系统的状态不仅与每一个小区域的人口增长率有关，而且还与各小区域之间的相互作用，以及整个大区与外界环境之间的人口扩散有关。我们取来自于系统之外的人口迁入量 $u(t)$ ，以及从系统流向外界的人口迁出量 $u'(t)$ 作为系统的输入变量，即 $u(t) = \begin{bmatrix} u(t) \\ u'(t) \end{bmatrix}$ ，则可建立该区域人口系统的状态方程。

为了建立模型的方便，我们引入如下有关参量：对于每一个小区域 i ，记 r_i 为人口生灭增长率(显然， r_i 与该小区域的人口生育政策、出生率、死

亡率、卫生医疗条件等有关。),记 e_i 为就业率, $e_i' = 1 - e_i$ 为未就业率. $e_i' = 1 - e_i$ 为未就业率。设 r_{ij} 为 i 小区的经济中心至 j 小区的经济中心的距离, 显然 $r_{ij} = r_{ji}$ 。

如果记 $x_{ij}(t)$ 为第 t 年由 j 小区流向 i 小区的人口数量; 而 $x_{ji}(t)$ 为第 t 年由 i 小区流向 j 小区的人口数量 $x_{i0}(t)$ 为第 t 年由系统外界迁入 i 小区的人口数量, 而 $x_{0i}(t)$ 为第 t 年由 i 小区迁出系统外界的人口数量。那么我们有

$$x_i(t+1) = (1+r_i)x_i(t) + \sum_{j=1}^n x_{ij}(t) - \sum_{j=1}^n x_{ji}(t) + x_{i0}(t) - x_{0i}(t) \quad (i = 1, 2, \dots, n) \quad (19)$$

由外界迁入系统的总人口 $u(t)$ 及由系统迁出外界的总人口 $u'(t)$, 是由系统的总体调控行为决定的, 但是它们在系统内部各小区域之间的数量分配却是由各小区域的就业率与未就业率决定的, 即

$$x_{i0}(t) = \frac{e_i}{\sum_{j=1}^n e_j} u(t) \quad (20)$$

$$x_{0i}(t) = \frac{e_j'}{\sum_{j=1}^n e_j'} u'(t) \quad (21)$$

另外, 一般地认为, 由第 j 小区流向 i 小区的人口数量与 i 小区的就业率 e_i 及 j 小区的人口数量成正比, 而与该两小区经济中心之间的距离成反比, 即

$$x_{ij}(t) = k \frac{e_i}{r_{ij}} x_j(t) \quad (22)$$

(22)式中, k 为经验常数。

$$\text{同理,} \quad x_{ji}(t) = k \frac{e_j}{r_{ji}} x_i(t) \quad (23)$$

将(20)式和(23)式代入(19)式得

$$x_i(t+1) = \left[(1+r_i) - \sum_{j=1}^n k \frac{e_j}{r_{ji}} \right] x_i(t) + \sum_{j=1}^n k \frac{e_i}{r_{ij}} x_j(t) + \frac{e_i}{\sum_{j=1}^n e_j} u(t) - \frac{e_j'}{\sum_{j=1}^n e_j'} u'(t) \quad (i = 1, 2, \dots, n) \quad (24)$$

如果记

$$\Phi = \begin{pmatrix} (1+r_1) - \sum_{j=1}^n k \frac{e_j}{r_{j1}} & k \frac{e_1}{r_{12}} & \dots & k \frac{e_1}{r_{1n}} \\ k \frac{e_2}{r_{21}} & \left[(1+r_2) - \sum_{j=1}^n k \frac{e_j}{r_{j2}} \right] & \dots & k \frac{e_2}{r_{2n}} \\ k \frac{e_n}{r_{n1}} & k \frac{e_n}{r_{n2}} & \dots & \left[(1+r_n) - \sum_{j=1}^{n-1} k \frac{e_j}{r_{jn}} \right] \end{pmatrix}$$

$$F = \begin{pmatrix} \frac{e_1}{\sum_{j=1}^n e_j} & -\frac{e_1'}{\sum_{j=1}^n e_j'} \\ \frac{e_2}{\sum_{j=1}^n e_j} & -\frac{e_2'}{\sum_{j=1}^n e_j'} \\ \vdots & \vdots \\ \frac{e_n}{\sum_{j=1}^n e_j} & -\frac{e_n'}{\sum_{j=1}^n e_j'} \end{pmatrix}$$

则该区域人口系统的状态方程为：

$$X(t+1) = \Phi X(t) + FU(t) \quad (25)$$

第三节 最大值原理及其应用

地理系统调控，其根本目的就是希望通过采取各种调控方法和手段，使人类活动与地理环境之间相互适应、相互协调，从而使人与地理环境相互作用的地理系统朝着良性有序的方向发展。最佳调控策略，是地理调控系统所追求的理想目标。本节，我们将运用最佳控制理论的有关思想和方法，提出地理系统最佳调控的数学模型，并介绍最大值原理在地理系统调控中的运用。

一、地理系统最佳调控的数学模型

由上节的介绍，我们知道，对于一个具有 n 个状态变量和 m 个控制或输入变量的地理系统，其运动规律可以由以下形式的常微分方程组，即状态方程来描述：

$$\begin{aligned}\frac{dx_i}{dt} &= f_i(x_1, x_2, \dots, x_n, u_1, u_2, \dots, u_m) \\ i &= 1, 2, \dots, n, t_0 \leq t \leq T\end{aligned}\quad (1)$$

(1)式中， x_i ($i=1, 2, \dots, n$)为状态变量， u_j ($j=1, 2, \dots, m$)为控制或输入变量。

(1)式也可以改写成向量形式：

$$\frac{dX}{dt} = F(X, U) \quad t \leq 0 \leq t \leq T \quad (2)$$

(2)式中， $X=[x_1, x_2, \dots, x_n]^T$ 是 R_n 空间中的向量，而 $U=[u_1, u_2, \dots, u_m]^T$ 是 R_m 空间中的向量， $F(X, U)=[f_1(x, u), f_2(x, u), \dots, f_n(x, u)]^T$ 为 n 维向量函数。

在有关实际问题的分析中，我们认为，对于状态空间的每一个点 X ，以及 R^m 空间中的每一个点 U ，诸函数 $f_i(x, u)$ ($i=1, 2, \dots, n$)是有定义的，而且对 x_i, u_j ($i=1, 2, \dots, n; j=1, 2, \dots, m$)连续，对于 x_i ($i=1, 2, \dots, n$)连续可微。即函数 $f_i(x_1, x_2, \dots, x_n; u_1, u_2, \dots, u_m)$ 和

$$\frac{\partial f_i(x_1, x_2, \dots, x_n; u_1, u_2, \dots, u_m)}{\partial x_i} \quad (i, j=1, 2, \dots, n) \text{ 对变量 } x_i, u_j \text{ 是连续的。}$$

一般来说，在系统控制理论中，对控制 $U(t)=[u_1(t), u_2(t), \dots, u_m(t)]^T$ 都要加一定的限制，譬如，要求诸控制 $u_j(t)$ ($j=1, 2, \dots, m$)在时间区间 $[t_0, T]$ 内分段连续，并且要求控制函数 $U(t)$ 满足如下的约束条件：

$$\varphi_i(x, u) \leq 0 \quad i=1, 2, \dots, l(l \leq m) \quad (3)$$

从而保证控制域在 R^m 空间中的某一个有界闭集 U_t 内取值，即

$$U(t) \in U_t \quad (4)$$

这样的控制 $U(t)$ 称为允许控制。

其次，初值条件

$$X(t_0)=X^{(0)} \quad (5)$$

也是系统调控所必不可少的条件。另外，在一些系统的调控问题中还可以加入终端条件：

$$\Psi_i(X(t), T)=0 \quad i=1, 2, \dots, k(k \leq n) \quad (6)$$

这里 T 可以是固定的，也可以是不固定的。

在 R^n 空间中，所有满足(6)式的点 $X(T)=[x_1(T), x_2(T), \dots, x_n(T)]^T$ 组成的集合：

$$\Omega=\{X(T) \mid \Psi_i(X(T), T)=0, X(T) \in R^n\} \quad (7)$$

称为目标集。具有终端约束(6)式的问题，叫做可变右端问题。如果目标集 Ω 蜕化成一个点，即 $X(T)=X_T$ ，(X_T 是给定的)，则称为固定右端问题。而称无终端约束问题为自由端问题。给定允许控制 $U(t) \in U_t$ ，只要它满足终端条件(6)式，那么对于任何可能的系统初始状态

$X(t_0) = X^{(0)} \notin \Omega$ ，都能唯一地确定受控对象之运动规律，使系统(2)式自 $X^{(0)}$

出发，在 $t=T$ 时刻到达目标集 Ω 。这样的允许控制称为可行控制。

地理系统的最佳调控，就是要在给定的系统初始状态 $X^{(0)} \notin \Omega$ 下，确定一个可行控制 $U^*(t)$ ，使系统(2)式自 $X^{(0)}$ 出发，在 $t=T$ 时刻到达目标集 Ω ，而且使：

$$\begin{aligned} \max(\min)J[U(t)] = \max(\min)\{G[X(T), T] \\ + \int_{t_0}^T f_0[X(t), U(t)]dt\} = J[U^*(t)] \end{aligned} \quad (8)$$

这样的控制 $U^*(t)$ 称为最优控制，与这个控制相应的控制轨迹 $X^*(t)$ 称为最佳控制轨迹。而泛函 $J[U(T)]$ 叫做评价系统质量优劣的目标泛函或质量指标或性能指标，简称指标，其中， $G[X(t), T]$ 这一项称为终端指标。

二、最大值原理

通过以上的讨论，我们知道，要确定地理系统的最佳调控策略，就需要在状态方程(1)(或(2))式，允许控制约束(3)(或(4))式，初始条件(5)式，以及终端约束条件(6)式下，优化系统的性能指标 $J[U(t)]$ ，使之达到最大(或最小)值。那么，怎样(满足什么条件)才能使性能指标 $J[U(t)]$ 达到最大值呢？为了回答这一问题，以下我们来介绍最大值原理。

为了讨论问题的方便，我们不妨将终端约束条件(6)式改成显式形式：

$$x_i(T) = x_T^0 \quad i = 1, 2, \dots, k \quad (0 \leq k \leq n) \quad (6')$$

这就意味着状态变量 $x_1, x_2, \dots, x_k$ 存在终端约束，而 $x_{k+1}, x_{k+2}, \dots, x_n$ 无终端约束。

为了给出最大值原理(性能指标 $J[U(t)]$ 取最大值的必要条件)，我们首先引入哈密顿函数：

$$\begin{aligned} H[X(t), U(t), \lambda(t)] = \lambda_0 f_0[X(t), U(t)] \\ + \sum_{i=1}^n \lambda_i(t) f_i[X(t), U(t)] \end{aligned} \quad (9)$$

(9)式中， λ_0 是常数， $\lambda_i(t)$ ($i=1, 2, \dots, n$) 是附加的未知函数，称为伴随变量。首先，伴随变量必须满足伴随方程：

$$\frac{d\lambda_i}{dt} = \frac{\partial H}{\partial x_i} = - \sum_{i=1}^n \lambda_i \frac{\partial f_i}{\partial x_i} \quad (i = 1, 2, \dots, n) \quad (10)$$

而且必须有 $(\lambda_0)^2 + \sum_{i=1}^n [\lambda_i(t)]^2 \neq 0$ (即所有乘子不能同时变为零)及

$$\lambda_0=0 \text{ 或 } \lambda_0=1 \quad (11)$$

其次

$$\frac{dx_i}{dt} = \frac{\partial H}{\partial \lambda_i} = f_i(X(t), U(t)) \quad (i = 1, 2, \dots, n) \quad (12)$$

再次, 对于所有 $t_0 \leq t \leq T$, 必须有

$$H[X(t), U(t), \lambda(t)] = \max_{U \in U_t} H[X(t), U(t), \lambda(t)] \quad (13)$$

最后, 对于(6')式没有规定终端值的每一个状态变量, 应有相应的横截条件:

$$\lambda(T) = (i=k+1, k+2, \dots, n) \quad (14)$$

如果终端时间 T 没有规定, 则有附加条件:

$$H(T) = 0 \quad (15)$$

这就是说, 如果 $U(t)$ 是使性能指标 $J[U(t)]$ 达到最大值的最佳控制, $X(t)$ 是相应的最佳轨迹, 则必须存在一个伴随向量 $\lambda(t)$, 使得(10)–(15)式中的所有条件均满足。

需要强调的是, 在上述条件(11)式中, $\lambda_0=0$ 时的情况是反常的; $\lambda_0=1$ 时的情况是正常的。但是, 这一附加条件却使上述最大值原理的描述更具一般性。读者不难验证, 一维控制问题的最大值原理只不过是上述原理的一个特例。

三、可更新资源的最佳利用策略

可更新资源主要是指生物资源和某些动态非生物资源。农业生产所利用的资源大部分属于可更新资源, 例如, 农作物、森林、牧草、野生动物、牲畜、淡水和海洋水产品(鱼类、虾等)以及水资源、土壤、人力资源等。可更新资源可借助于自然生长、繁殖不断地进行自我更新而维持一定的数值。如果对这些资源进行科学管理和合理地开发利用, 将会取之不尽, 用之不竭; 反之, 管理和使用不当, 将会使这些资源受到损害, 甚至完全枯竭, 从而给人类的经济活动带来巨大的损失。一般地讲, 可更新资源枯竭的危险要比非可更新资源更大。这是因为所有可更新资源都要受到自然更新能力的限制, 如果人们的生产经营活动不够得当, 就可能超出这种限制从而导制资源枯竭, 甚至会导致不可逆转的永久性枯竭和灭绝。如果研究单一种类可更新资源的管理利用问题, 则这一问题就是一维系统的最佳调控问题。这里, 我们取资源的储备量 $x=x(t)$ 为系统的状态变量, 资源收获量 $u=u(t)$ 为控制变量。如果记该种可更新资源的自然增长率函数为 $F(x)$, 则该调控系统的动态描述为:

$$\begin{cases} \frac{dx}{dt} = F(x) - u(t) \\ x(0) = x_0 \end{cases} \quad (16)$$

在收获过程中, 一个最基本的要求就是

$$x(t) > 0 \quad (17)$$

为了维护资源的再生能力, 我们对收获量提出以下约束

$$u_{\min} \leq u(t) \leq u_{\max} \quad (18)$$

如果资源的单位收获量价格为 P , 成本为 $C(x)$, 那么持续的经济收益为

$$J_0 = \int_0^{+\infty} [P - C(x)]u(t)dt \quad (19)$$

若货币贴现率为 δ ($\delta > 0$)，那么上述收益折成现值后的总经济收益为

$$J(u) = \int_0^{+\infty} E^{-\delta t} [P - C(x)u(t)] dt \quad (20)$$

我们的目的是在维持生态平衡的前提下获得持续的最佳经济效益。因此，问题的实质就是在约束条件(16)，(17)，(18)式下求目标泛函(20)式的极大值。这就意味着确定一个允许的收获量 $u^*(t)$ ，使资源储量稳定在某个水平上而获得最佳收益。利用最大值原理可知最优状态水平 x^* 满足如下方程：

$$\frac{d\rho(x)}{dx} = \sigma [\rho - C(x)] \quad (21)$$

(21)式中， $\rho(x) = F(x)[P - C(x)]$ 表示资源储量为 x 时的持续经济利润。

若方程(21)式有唯一的解 x^* ，那么资源的最佳利用策略为：

$$u^*(t) = \begin{cases} u_{\max} & \text{当 } x > x^* \text{ 时} \\ F(x^*) & \text{当 } x = x^* \text{ 时} \\ u_{\min} & \text{当 } x < x^* \text{ 时} \end{cases}$$

如果在种群水平上探讨可更新资源(如森林资源、鱼类资源、草场资源等生物资源)的管理利用问题，则状态变量 $x=x(t)$ 就代表在 t 时刻资源(林木、鱼类、牧草等)的种群水平，而控制变量 $u=u(t)$ 则代表森林的采伐速率、鱼的捕获速率、牧草的采食速率等。对于这类生物资源，其种群自然增长过程可以用逻辑斯谛(Logistic)方程描述，即自然增长率函数为

$$F(x) = rx \left(1 - \frac{x}{K}\right) \quad (22)$$

(22)式中 r 代表资源种群的内容增长率， K 代表环境容纳量。

如果认为收获成本与种群水平成反比，则可以假设：

$$C(x) = \frac{C}{x} \quad (23)$$

(23)式中， C 为常数。将(22)式与(23)式代入(21)式可以得一正数解：

$$x^* = \frac{K}{4} \left\{ \left(\frac{C}{PK} + 1 - \frac{\delta}{r} \right) + \left[\left(\frac{C}{PK} + 1 - \frac{\delta}{r} \right)^2 + \frac{8\delta C}{\delta} \right]^{\frac{1}{2}} \right\} \quad (24)$$

此时，资源的最佳利用策略为

$$u^*(t) = \begin{cases} u_{\min} & \text{当 } x < x^* \text{ 时} \\ rx^* \left(1 - \frac{x^*}{K}\right) & \text{当 } x = x^* \text{ 时} \\ u_{\max} & \text{当 } x > x^* \text{ 时} \end{cases}$$

第四节 大系统理论及其应用

本世纪 70 年代以后,控制论的研究对象已经从简单系统发展到更加复杂的大系统。对于大系统来说,由于其变量众多、结构复杂,因而经典控制论的传递函数法和现代控制论的状态空间法已经不能完全适用了。在这种情形下,大系统理论也就应运而生了。目前,处理大系统的主要方法有:分解-协调原理,分散最优控制,多级递阶控制,大系统模型降阶理论,向量李雅普诺夫稳定性理论等。本节,我们拟结合国内有关学者的研究成果,对大系统理论及其在地理系统调控中的应用作一简单的介绍。

一、大系统理论简述

(一)大系统的概念

什么是“大系统”?这是一个至今尚未给出确切回答的问题。一般说来,人们总是认为,规模庞大、结构复杂的系统是大系统。譬如,现代化大型企业管理系统,大的区域性或全国性的电力网络管理调度系统,国民经济计划管理系统,城市生态系统,区域农业生态系统,等等。

(二)大系统理论的主要研究内容

目前,大系统理论还正处于发展阶段,其主要的研究内容可以概括为如下两个方面:

1.大系统分析与建模 大系统分析,就是对系统的技术性能、经济指标、社会效果、生态环境影响等进行分析,作出评价;对系统现有的运行状态进行观测与估算;对系统未来的发展动态与趋势进行预测与模拟;对系统的环境条件及其变化进行观测与分析等等。大系统分析的目的是寻求改进现有系统或筹建新系统的方法,为系统方案的选择与制定决策提供依据,从而实现大系统管理、控制、运行的最优化。

为了进行系统分析,需要建立模型。大系统分析采用的模型是多种多样的,譬如,用来描述系统的动态及静态特性、性能指标、运行状态的数学模型;用来表示信息流与物质流、时间顺序、逻辑关系等相互联系的网络图模型;用来模仿实际系统的物理过程、运行状态、生理或心理活动的技术模型等等。

2.大系统综合 大系统综合,就是在大系统分析的基础上,按系统构成要素之间的相互关系,各级子系统之间的联结规律,彼此逐级联结起来,形成从简单到复杂,从低级到高级的大系统的过程。其任务就是要对尚待筹建的大系统进行决策、规划与设计,对大系统的筹建过程与实际运行过程进行科学的计划协调与组织管理,根据大系统的总任务与总目标,选择设计方案,确定控制策略,制定管理方法,解决大系统的优化设计、优化控制和优化管理问题。

二、大系统的结构方案

(一)多级递阶结构方案

多级递阶控制方法是一种受到广泛注意的大系统结构方案。以三级递阶结构为例,其结构方案如图 8-4 所示。

在图 8-4 中,第一级为局部控制级(最低决策层),它直接控制大系统的各个局部对象或过程;第二级为递阶控制级(中间决策层),它对第一级的各控制机构进行协调控制;第三级为协调控制级(最高决策层),它对第二级进

行协调控制，根据大系统的总目标，通过递阶结构，完成大系统的控制任务。

(二) 多层控制结构方案

多层控制方案的特点是按任务或功能分层。较高层的任务功能较复杂，自扰动因素变动较慢；较低层的任务功能简单，自扰动因素变化较快。各层之间的关系具有“分工”性质，上下层之间也有隐含的“支配”与“被支配”的关系。这种结构具有纵向的信息交换，除了由上层到下层的信息交换之外，还有从受控对象或过程到各层的反馈信息(如图 8-5 所示)。

在图 8-5 中，第一层为直接控制层，其功能是根据最优控制层的指令，直接控制受控对象或过程的运行状态，克服快变化的影响；第二层为最优化层，其功能是根据给定的目标函数、约束条件及系统的数学模型，进行最优控制；第三层为自适应层，其功能是根据对象特性或运行情况变化，对最优化层的目标函数、约束条件及数学模型进行适当的修正，以保证系统最优运行状态；第四层为自组织层，其功能是根据大系统的总目标与总任务，考虑环境变化的条件，制订决策、计划协调与组织管理。

(三) 多段控制结构方案

多段控制结构方案的特点，是按受控过程的时间分段，根据受控过程的时间顺序，将全过程划分为若干段，每一段形成一个小系统。不过，这里的协调段按衔接条件进行协调，它把前段终点边界条件与后段初始边界条件衔接起来，完成全过程的控制任务。相邻的各段控制小系统之间通过协调段与受控对象具有横向信息交换，上下级之间也有纵向的信息交换(如图 8-6 所示)。

三、大系统的分解与协调

在复杂的大系统控制过程中，为了保证系统优化控制策略的实现，常常需要依据大系统理论的分解协调原理，将系统的总目标和总模型分解为子目标和子模型，并根据于目标和子模型的相互作用以及外部环境及约束条件的变化，来反复协调各个子目标和子模型的关系，以至调整大系统的总目标和总模型。

(一) 大系统分解

大系统分解的基本思想是在优化过程中，先将高阶大系统划分为若干个低阶子系统，然后再用通常的优化方法使各个子系统优化。大系统分解时，需要解决两个问题，一是目标函数(性能指标)的分解；二是模型关联的分解。而实际上，常见的情形是大系统最优化的总目标是可以分离的，例如大系统的目标函数是各个子系统的目标函数之和；而系统的模型是相互关联的，例如子系统的状态变量之间存在着相互影响。因此，大系统分解主要是解决系统模型关联的分解问题。下面我们以线性大系统为例说明之。

考虑具有二次性能指标的线性大系统，其状态方程与输出方程如下：

$$\begin{cases} \mathbf{x}(k+1) = \mathbf{A}_{\mathbf{x}(k)} + \mathbf{B}_{\mathbf{u}(k)} \\ \mathbf{y}(k) = \mathbf{C}_{\mathbf{x}(k)} \end{cases} \quad (1)$$

式中 $\mathbf{x}(k)$ 为 n 维的状态变量；

$\mathbf{u}(k)$ 为 m 维的控制或输入变量；

$\mathbf{y}(k)$ 为 r 维的输出变量；

A, B, C 分别为 $n \times n$, $n \times m$, $r \times n$ 的常数矩阵。

此大系统的总目标函数是：

$$J = \sum_{k=0}^{k_f} [y(k)^T Q y(k) + u(k)^T R u(k)] \quad (2)$$

式中, Q、R 皆为常数矩阵； k_f 为终止时间。

设 B、C、Q、R 皆为对角形分块矩阵, 即它们分别具有以下形式：

$$B = \begin{pmatrix} B_1 & & 0 \\ & B_2 & \\ & & O \\ 0 & & & B_N \end{pmatrix} \quad C = \begin{pmatrix} C_1 & & 0 \\ & C_2 & \\ & & O \\ 0 & & & C_N \end{pmatrix}$$

$$Q = \begin{pmatrix} Q_1 & & 0 \\ & Q_2 & \\ & & O \\ O & & & Q_N \end{pmatrix} \quad R = \begin{pmatrix} R_1 & & 0 \\ & R_2 & \\ & & O \\ O & & & R_N \end{pmatrix}$$

相应地, 把 A 也写成如下的分块形式：

$$A = \begin{bmatrix} A_{11} & A_{12} & \cdots & A_{1N} \\ A_{21} & A_{22} & \cdots & A_{2N} \\ M & M & O & M \\ A_{N1} & A_{N2} & \cdots & A_{NN} \end{bmatrix}$$

上述这一系列假定, 实质上是说, 大系数(1)式中的各子系统只通过状态变量发生联系, 而大系统的总目标函数是可分解的, 它可以表示为

$$J = \sum_{i=1}^N J_i \quad (3)$$

其中 N 为子系统的数目, 各子系统的目标函数分别为

$$J_i = \frac{1}{2} \sum_{k=0}^{k_f} [x_i(k)^T C_i^T Q_i C_i x_i(k) + u_i^T(k) R_i u_i(k)]$$

$$i=1, 2, \dots, N \quad (4)$$

大系统的状态方程(1)式可分解为：

$$\begin{cases} x_i(k+1) = A_{ii} x_i(k) + B_i u_i(k) + \sum_{j=1}^N A_{ij} x_j(k) \\ y_i(k) = C_i x_i(k) \quad i = 1, 2, \dots, N \end{cases} \quad (5)$$

由上述分析可以看出, 就目标函数而言, 我们已经把总目标函数分解成了 N 个子系统的目标函数之和, 从而能够把 N 个子系统的目标函数分离出来。

现在的问题是怎样处理状态变量之间的关联项 $\sum_{j=1}^N A_{ij} x_j(k)$, 我们才能获得 N

个子系统的相对独立的状态方程。解决这一问题, 主要有下面两种方法：

1. 非现实法 这种方法的基本思想是引入一系列的伪变量来描述状态变量关联的影响从而达到企图“切断”各子系统之间的模型关联的目的, 在 (5) 式中令

$$m_i = \sum_{j=1}^N A_{ij} x_j$$

则有

$$\begin{cases} x_i(k+1) = A_{ii}x_i(k) + B_i u_i(k) + m_i(k) \\ y_i(k) = C_i x_i(k) \end{cases} \quad (6)$$

于是原大系统的最优化问题就变成了根据无关联的状态方程(6)式在附加约束条件

$$m_i(k) = \sum_{j=1}^N A_{ij} x_j(k) \quad i = 1, 2, \dots, N \quad (7)$$

之下, 求控制策略 $u(0), \dots, u(k_f-1)$, 使

$$J = \sum_{i=1}^N J_i$$

取最小值。

为了求解上述问题, 我们用拉格朗日乘子将约束条件(7)式引入到目标函数中, 于是得出修正后的大系统目标函数:

$$\begin{aligned} J^* = \sum_{i=1}^N \sum_{k=0}^{k_i} \left\{ \frac{1}{2} [x_i^T(k) C_i^T Q_i C_i x_i(k) + u_i^T(k) R_i u_i(k)] \right. \\ \left. + \lambda_i^T(k) \left[\sum_{j=1}^N A_{ij} x_j(k) - m_i(k) \right] \right\} = \sum_{i=1}^N J_i^* \end{aligned} \quad (8)$$

此处, J_i^* 是子系统修正后的目标函数, 具体地说:

$$\begin{aligned} J_i^* = \frac{1}{2} \sum_{k=0}^{k_i} [x_i^T(k) C_i^T(K) C_i^T Q_i C_i x_i(k) + u_i^T(k) R_i u_i(k)] \\ + \sum_{k=0}^{k_i} \left[\lambda_i^T(k) \sum_{j=1}^N A_{ij} x_j(k) - \lambda_i^T(k) m_i(k) \right] \quad i = 1, 2, \dots, N \end{aligned} \quad (9)$$

这样就可以把大系统分解成为 N 个子系统, 于是即可分别对每个子系统进行最优化运算。

2. 现实法这种分解方法, 不是假想“切断”模型关联, 而是将状态变量的关联设置为某些预估值, 即置

$$\begin{aligned} A_{ij} x_j(k) = m_{ij}(k) \quad i \neq j \quad i = 1, 2, \dots, N \\ j = 1, 2, \dots, N \end{aligned} \quad (10)$$

且令:

$$\phi_i(k) = \sum_{j=1}^N m_{ij}(k) \quad i = 1, 2, \dots, N \quad (11)$$

把这样的 ϕ_i ; 代入(5)式中, 可在形式上把大系统状态方程(1)式分解成 N 个独立的子系统状态方程的联合形式:

$$\begin{cases} x_i(k+1) = A_{ii}x_i(k) + B_i u_i(k) + \phi_i(k) \\ y_i(k) = C_i x_i(k) \end{cases} \quad (12)$$

$$i = 1, 2, \dots, N$$

考虑到预估关联的关系式是作为约束而出现的, 自然可利用拉格朗日乘子, 把约束条件引到目标函数之中, 于是得到修正的大系统目标函数为:

$$J^* = \sum_{i=1}^N J_i^* \quad (13)$$

(13)式中, J_i^* 是子系统修正的目标函数, 具体地有:

$$\begin{aligned} J_i^* = & \frac{1}{2} \sum_{k=0}^{k_i} [x_i^T(k) C_i^T Q_i C_i x_i(k) + u_i^T(k) R_i u_i(k)] \\ & + \sum_{k=0}^{k_i} \sum_{j=1}^N [\lambda_{ij}^T(k) A_{ij} x_j(k) - \lambda_{ij}^T(k) m_{ij}(k)] \\ & i = 1, 2, \dots, N \end{aligned} \quad (14)$$

这样, 根据低阶的状态方程(12)式和分离的目标函数(14)式, 把 $\lambda_{ij}(k)$, $m_{ij}(k)$ 看成是待定的参量, 可用通常方法求得子系统的局部最优解。

(二)大系统协调

作为大系统优化的第二步, 就是协调, 即在分解后各子系统局部优化的基础上, 使总目标函数极大(小)化, 以实现大系统的全局优化。在协调过程中, 首先需要解决的是协调原则问题, 即根据什么原则, 对各个子系统进行协调控制。在大系统的分解协调理论中。最基本的协调原则有关联平衡和关联预估原则。这两种协调原则, 都是按协调偏差所进行的反馈闭环控制, 只是所选择的协调变量不同而已。

综上所述, 大系统的多级优化控制。一般地分两个步骤进行。第一步为系统分解, 即用非现实法和现实法把大系统分解为若干低阶子系统, 进行局部优化。第二步为系统协调, 即用关联平衡原则或关联预估原则, 按协调偏差进行反馈控制, 以实现大系统的全局优化。当然, 这两步是不能截然分开的, 而是反复进行的统一过程。

四、大系统理论的应用举例

作为大系统理论在地理系统优化调控中的应用实例, 下面我们来介绍汤兵勇先生等(1990)所研制的吉林市水污染控制规划模型。

(一)数学模型

根据松花江段环境标准, 按 90%频率, 吉林市 BOD_5 总削减量为 85.36 吨, 集中于五个主要处理厂。在该污水处理系统进行削减量分配, 以使经济费用最小, 所提出的优化问题数学模型为:

$$\begin{cases} \min F = \sum_{i=1}^5 C_i x_i^2 \\ \sum_{i=1}^5 w_i x_i = 85.36 \\ 0 \leq x_i \leq b_i \quad (i = 1, 2, 3, 4) \end{cases} \quad (15)$$

吉林市现有及正在设计的污水处理系统技术经济情况见表 8-1。

由于吉林化工综合污水处理系统已建成, 故其削减量和投资运行支出费按设计数值选择, 则上述数学模型可改写为

$$\begin{cases} \min Z = \sum_{i=0}^4 C_i x_i^2 \\ \sum_{i=1}^4 w_i x_i = 50.5 \\ 0 \leq x_i \leq b_i \quad (i = 1, 2, 3, 4) \end{cases} \quad (16)$$

(二)二级递阶分解协调的计算格式

利用拉格朗日乘子法，按 x_i 分解模型：

$$L(x, \lambda) = \sum_{i=1}^4 C_i x_i^2 + \lambda \left(\sum_{i=1}^4 w_i x_i - 50.5 \right) = \sum_{i=1}^4 L_i \quad (17)$$

(17)式中， $L_i = C_i x_i^2 + \lambda w_i x_i - 12.625 \lambda$ ($i = 1, 2, 3, 4$) 于是，我们得到四个带有参数 λ 的子系统，其选取拉格朗日乘子 λ 作为协调量，问题的优化可通过一个二级递阶分解协调的方式来完成。

第一级：分别对四个子系统(即对吉林造纸厂、吉林九站造纸厂、吉林化纤厂以及城市

表 8-1 吉林市污水处理系统设计技术经济情况表

序号	工厂名称	污 水 流 量 (m^3/d)	BOD ₅ 需处理量 (t/d)	去除量 (t/d)	去除率 (%)	建设 投资 (万元)	9 年运 支 费 (万元)	总费用 (万元)	费
1	吉林造纸厂	40000	7	6	85.7	400	1188	1588.0	2
2	吉林九站造纸厂	1920	1.34	0.69	51.5	28.9	57	85.9	3
3	吉林化纤厂	14400	4.32	4.04	93.3	260	641.5	901.5	10
4	城市污水系统	175000	52.82	48.89	91.8	4934.8	3372.3	8307.1	98
5	化工综合污水处理系统	120000	38.7	34.86	90.2	5975.9	7676.7	13652.6	165
合计	/	351320	104.18	94.48	/	11599.6	12935.5	24535.1	

污水处理系统)进行去除率的优化：

$$\begin{cases} \min L_i = c_i x_i^2 \div \lambda w_i x_i \\ 0 \leq x_i \leq b_i \quad (i = 1, 2, 3, 4) \end{cases} \quad (18)$$

这可采用任何一种单变量寻优算法。

第二级：根据对偶理论，以 $\frac{\Delta L}{\lambda}$ 作为 λ 的寻优方向，构成 λ 的改善关系式：

$$\lambda^{k+1} = \lambda^k + a \cdot \frac{\Delta L}{\lambda^k} \quad (19)$$

当满足平衡关系

$$\sum_{i=1}^4 w_i x_i = 50.5 \quad (20)$$

时，即达到问题的整体最优。

根据上述推导结果，这一问题的具体计算步骤是：

(1) 给定初始值 λ_0 、 α_0 ，分别进行一级优化，采用 0.618 法，求得 x_i^k ($i = 1, 2, 3, 4$; $k = 0, 1, 2, \dots$)。

(2) 根据第一级计算结果，判别 $\sum_{i=1}^4 w_i x_i$ 是否为零？若满足约束条件，转 (4)；否则转 (3)。

(3) 以梯度作为寻优方向，按已叙述的 λ 递推关系式计算新的协调变量 λ^{k+1} ，即

$$\lambda^{k+1} = \lambda^k + \alpha \cdot \frac{\Delta L}{\lambda^k} \quad (21)$$

此处步长 α 的选择一般可利用一维搜索确定最优步长。计算完毕，返回 (1)。

(4) 输出 x_i ($i = 1, 2, 3, 4$) 据此可计算各处理厂的削减量并计算总费用。

(三) 计算结果

上述问题的计算结果如表 8-2 所示。

表 8-2 优化计算结果

项目	工厂名称序号				
	1	2	3	4	合计
去除率(x_i , 单位 :%)	49.79	51.24	64.28	82.89	/
去除量(t/d)	3.4854	0.6866	2.777	437856	50.7346
费用(万元)	536.0059	85.0754	427.7028	6776.6779	7825.46

整个吉林市污水处理系统(五个处理厂)的 BOD_5 削减量为：

$$50.7346 + 34.86 = 85.5946 \text{ (t/d)}$$

整个系统的经济费用为

$$7825.46 + 13652.6 = 21478.06 \text{ (万元)}。$$

第九章 模糊数学方法

模糊数学方法，是一种研究和处理模糊现象的新型数学方法。这一方法，是由美国自动控制专家查德(L.A.Zadeh)于 1965 年首次提出来的。20 多年来，模糊数学方法在自然科学和社会科学研究的各个领域得到了广泛应用。

第一节 模糊数学基本知识简介

一、模糊子集及其运算

在经典集合论中，一个元素对于一个集合，要么属于，要么不属于，二者必居其一，绝不允许模棱两可。这一要求就从根本上限定了以经典集合论为基础的常规数学方法的应用范围，它只能用来研究那些具有绝对明确的界限的事物和现象。但是，在现实世界中，并非所有事物和现象都具有明确的界限。譬如，“高与矮”，“好与坏”，“美与丑”，……，这样一些概念之间就没有绝对分明的界限。严格说来，这些概念就是没有绝对的外延，这些概念被称之为模糊概念，它们不能用一般集合论来描述，而需要用模糊集合论去描述。

(一)模糊子集及其表示方法

1. 模糊子集

(1)隶属函数：在经典集合论中，一个元素 x 和一个集合 A 之间的关系只能有 $x \in A$ 或者 $x \notin A$ 这两种情况。集合可以通过其特征函数来刻画，每一个集合 A 都有一个特征函数 $C_A(x)$ ，其定义如下：

$$C_A(x) = \begin{cases} 1 & \text{当 } x \in A \\ 0 & \text{当 } x \notin A \end{cases} \quad (1)$$

(1)式所表示的特征函数的图形，如图 9-1 所示。由于经典集合论的特征函数只允许取 0 与 1 两个值，故与二值逻辑 $\{0, 1\}$ 相对应。

模糊数学是将二值逻辑 $\{0, 1\}$ 拓广到可取 $[0, 1]$ 闭区间上任意的无穷多个值的连续值逻辑。因此，也必须把特征函数作适当的拓广，这就是隶属函数 $\mu(x)$ ，它满足：

$$0 \leq \mu(x) \leq 1 \quad (2)$$

(1)式也可以记作 $\mu(x) \in [0, 1]$ ，一般情形下，其图形如图 9-2 所示。

(2)模糊子集的定义：1965 年，查德首次给出了模糊子集的定义：设 U 是一个给定的论域(即讨论对象的全体范围)， $\mu_A: x \rightarrow [0, 1]$ 是 U 到 $[0, 1]$ 闭区间上的一个映射，如果对于任何 $x \in U$ ，都有唯一的 $\mu_A(x) \in [0, 1]$ 与之对应，则该映射便给定了论域 U 上的一个模糊子集 $\tilde{A}$ ， μ_A 称做 $\tilde{A}$ 的隶属函数， $\mu_A(x)$ 称做 x 对 $\tilde{A}$ 的隶属度。

2. 模糊子集的表示方法 通过上述关于模糊子集的定义可以看出，一个模糊子集完全由其隶属函数所刻画。因此，模糊子集通常有以下几种表示方法：

(1)如果论域 U 是有限集时，可以用向量来表示模糊子集 $\tilde{A}$ 。一般地，若论域为 $U = \{x_1, x_2, \dots, x_n\}$ ，则模糊子集 $\tilde{A}$ 可表示为：

$$\tilde{A} = [\mu_1, \mu_2, \dots, \mu_n] \quad (3)$$

在(3)式中， $\mu_i \in [0, 1]$ ($i=1, 2, \dots, n$) 为第 i 个元素 x_i 对 $\tilde{A}$ 的隶属度。

(2)查德表示方法：①如果论域 U 是有限集时，采用查德记号可以将模糊

子集 $\underline{A}$ 表示为：

$$\underline{A} = \frac{\mu_1}{x_1} + \frac{\mu_2}{x_2} + \cdots + \frac{\mu_n}{x_n} = \sum_{i=1}^n \frac{\mu_i}{x_i} \quad (4)$$

应该注意，(4)式的记号决不是分式求和，而只是一个记号而已，其“分母”表示论域 U 中的元素，“分子”是相应元素的隶属度，当隶属度为 0 时，那一项可以不写入。②如果论域 U 是无限集时，采用查德记号可以将模糊子集 $\underline{A}$ 表示为：

$$\underline{A} = \int_{x \in U} \mu_A(x) / x \quad (5)$$

在(5)式中，“积分号”不是普通的积分，也不代表求和，而是表示各个元素与其隶属度对应关系的一个总括。

(3)如果给出了论域 U 上的模糊子集 $\underline{A}$ 的隶属函数的解析表达式，则也就表示出了模糊子集 $\underline{A}$ 。

(二)模糊子集的运算及其性质

1.模糊子集的运算 论域 U 上两个模糊子集 $\underline{A}$ 和 $\underline{B}$ 之间的相等、包含关系及并、交、补运算，分别规定如下：

$$(1) \underline{A} = \underline{B} \Leftrightarrow \mu_A(x) = \mu_B(x) \quad \forall x \in U$$

$$(2) \underline{A} = \Phi \Leftrightarrow \mu_A(x) = 0 \quad \forall x \in U$$

$$(3) \underline{A} \subseteq \underline{B} \Leftrightarrow \mu_A(x) \leq \mu_B(x) \quad \forall x \in U$$

$$(4) \underline{A} \Leftrightarrow \mu_{\bar{A}}(x) = 1 - \mu_A(x) \quad \forall x \in U$$

$$(5) \underline{A} \cup \underline{B} \Leftrightarrow \mu_{A \cup B}(x) = \max\{\mu_A(x), \mu_B(x)\}$$

$$= \mu_A(x) \vee \mu_B(x) \quad \forall x \in U$$

$$(6) \underline{A} \cap \underline{B} \Leftrightarrow \mu_{A \cap B}(x) = \min\{\mu_A(x), \mu_B(x)\}$$

$$= \mu_A(x) \wedge \mu_B(x) \quad \forall x \in U$$

上述记号“ $\vee$ ”和“ $\wedge$ ”是运算符号，简称为算子，“ $\vee$ ”表示取最大值，“ $\wedge$ ”表示取最小值

2.模糊子集运算的基本性质 对于模糊子集的运算，它具有如下几个基本性质。

$$(1) \text{幂等律：} \underline{A} \cup \underline{A} = \underline{A}, \underline{A} \cap \underline{A} = \underline{A}$$

$$(2) \text{交换律：} \underline{A} \cup \underline{B} = \underline{B} \cup \underline{A}, \underline{A} \cap \underline{B} = \underline{B} \cap \underline{A}$$

$$(3) \text{结合律：} (\underline{A} \cup \underline{B}) \cup \underline{C} = \underline{A} \cup (\underline{B} \cup \underline{C})$$

$$(\underline{A} \cap \underline{B}) \cap \underline{C} = \underline{A} \cap (\underline{B} \cap \underline{C})$$

$$(4) \text{吸收律: } \underline{\underline{A}} \cup (\underline{\underline{A}} \cap \underline{\underline{B}}) = \underline{\underline{A}}$$

$$\underline{\underline{A}} \cap (\underline{\underline{A}} \cup \underline{\underline{B}}) = \underline{\underline{A}}$$

$$(5) \text{分配律: } \underline{\underline{A}} \cup (\underline{\underline{B}} \cap \underline{\underline{C}}) = (\underline{\underline{A}} \cup \underline{\underline{B}}) \cap (\underline{\underline{A}} \cup \underline{\underline{C}})$$

$$\underline{\underline{A}} \cap (\underline{\underline{B}} \cup \underline{\underline{C}}) = (\underline{\underline{A}} \cap \underline{\underline{B}}) \cup (\underline{\underline{A}} \cap \underline{\underline{C}})$$

$$(6) \text{对合律: } \overline{\underline{\underline{A}}} = \underline{\underline{A}}$$

$$(7) \text{De Morgan法则: } \overline{\underline{\underline{A}} \cup \underline{\underline{B}}} = \overline{\underline{\underline{A}}} \cap \overline{\underline{\underline{B}}}$$

$$\overline{\underline{\underline{A}} \cap \underline{\underline{B}}} = \overline{\underline{\underline{A}}} \cup \overline{\underline{\underline{B}}}$$

$$(8) \text{常数运算法则: } \underline{\underline{A}} \cup U = U \quad \underline{\underline{A}} \cap U = \underline{\underline{A}}$$

二、模糊子集的 α -截集及其性质

(一) 模糊子集的 α -截集

定义：设 $\underline{\underline{A}}$ 是论域 U 上的一个模糊子集，其隶属函数为 μ_A ， x 对 $\underline{\underline{A}}$ 的隶属度为 $\mu_A(x)$ 。对于任- $\alpha \in [0, 1]$ ，称集合

$$\underline{\underline{A}}_{\alpha} = \{x \mid \mu_A(x) > \alpha, x \in U\} \quad (6)$$

为 $\underline{\underline{A}}$ 的强 α -截集；称集合

$$\underline{\underline{A}}_{\alpha}^{-} = \{x \mid \mu_A(x) \geq \alpha, x \in U\} \quad (7)$$

为 $\underline{\underline{A}}$ 的弱 α -截集。

有时也将强 α -截集与弱 α -截集统称为 α -截集。

(二) 模糊子集的 α -截集的性质

模糊子集的 α -截集，具有下述几个基本性质：

$$(1) \text{若 } \alpha \leq \beta, \text{ 则 } \underline{\underline{A}}_{\beta} \subseteq \underline{\underline{A}}_{\alpha}, \quad \underline{\underline{A}}_{\beta}^{-} \subseteq \underline{\underline{A}}_{\alpha}^{-}$$

$$(2) \text{对于任意 } \alpha \in [0, 1], \text{ 都有:}$$

$$\underline{\underline{A}}_{\alpha} \subseteq \underline{\underline{A}}_{\alpha}^{-}$$

$$(3) \text{对于任意 } \alpha \in [0, 1], \text{ 都有:}$$

$$(\underline{\underline{A}} \cap \underline{\underline{B}})_{\alpha} = \underline{\underline{A}}_{\alpha} \cap \underline{\underline{B}}_{\alpha}, \quad (\underline{\underline{A}} \cup \underline{\underline{B}})_{\alpha} = \underline{\underline{A}}_{\alpha} \cup \underline{\underline{B}}_{\alpha}$$

$$(4) \text{对于任意 } \alpha \in [0, 1], \text{ 都有:}$$

$$(\underline{\underline{A}} \cap \underline{\underline{B}})_{\alpha}^{-} = \underline{\underline{A}}_{\alpha}^{-} \cap \underline{\underline{B}}_{\alpha}^{-}, \quad (\underline{\underline{A}} \cup \underline{\underline{B}})_{\alpha}^{-} = \underline{\underline{A}}_{\alpha}^{-} \cup \underline{\underline{B}}_{\alpha}^{-}$$

$$(5) \bigcup_{\alpha < \beta} \underline{\underline{A}}_{\beta} = \underline{\underline{A}}_{\alpha}, \quad \bigcap_{\alpha > \beta} \underline{\underline{A}}_{\alpha}^{-} = \underline{\underline{A}}_{\alpha}^{-}$$

$$(6) \bigcap_{\alpha > \beta} \underline{\underline{A}}_{\beta} = \underline{\underline{A}}_{\alpha}, \quad \bigcup_{\alpha < \beta} \underline{\underline{A}}_{\beta}^{-} = \underline{\underline{A}}_{\alpha}^{-}$$

$$(7) \underline{\underline{A}} = \bigcup_{\alpha \in [0,1]} \alpha \cdot \underline{\underline{A}}_{\alpha} \quad (8)$$

在(8)式中, $\alpha \cdot A_{\alpha}$ 为模糊子集, 其隶属函数为

$$\mu_{\alpha \cdot A_{\alpha}}(x) = \begin{cases} \alpha & \text{当 } x \in A_{\alpha} \\ 0 & \text{当 } x \notin A_{\alpha} \end{cases}$$

三、模糊关系与模糊变换

(一)模糊关系

1.模糊关系的概念 模糊关系,是一般关系的推广,其定义如下:设 U 和 V 是两个普通集合, U 与 V 的直积

$$U \times V = \{(x, y) \mid x \in U, y \in V\}$$

上的一个模糊子集 R 便称为 U 到 V 上的一个模糊关系。若 $(x, y) \in U \times V$,

则称 $\mu_R(x, y)$ 为 x 与 y 具有关系 R 的程度。一般地, $\mu_R(x, y)$ 也可以记为 $R(x, y)$ 。特别地, 当 $U = V$ 时, 则称 R 为 U 中的模糊关系。

当 U 和 V 为有限集合时, 模糊关系 R 可以用矩阵表示为:

$$R = (r_{ij})_{m \times n} = \begin{pmatrix} r_{11} & r_{12} & \cdots & r_{1n} \\ r_{21} & r_{22} & \cdots & r_{2n} \\ \cdots & \cdots & \cdots & \cdots \\ r_{m1} & r_{m2} & \cdots & r_{mn} \end{pmatrix}$$

在(9)式中, $r_{ij} = \mu_R(x_i, y_j)$, $r_{ij} \in [0, 1]$, $i=1, 2, \cdots, m$; $j=1, 2, \cdots, n$; m 为 U 中所含元素的个数, n 为 V 中所含元素的个数。(8)式所示的矩阵称为模糊关系矩阵, 简称模糊矩阵。

因为模糊关系就是集合 U 与 V 的直积 $U \times V$ 上的模糊子集, 所以它的相等、包含、并、交、补等运算与模糊子集的概念和运算性质完全相同, 这里不再作重复。下面介绍 U 中几个重要的特殊关系。

(1)恒等关系 I :

$$I \Leftrightarrow \mu_I(x, y) = \begin{cases} 1 & \text{当 } x = y \\ 0 & \text{当 } x \neq y \end{cases} \quad \forall x, y \in U$$

(2)零关系 O :

$$O \Leftrightarrow \mu_O(x, y) = 0 \quad \forall x, y \in U$$

(3)全称关系 E :

$$E \Leftrightarrow \mu_E(x, y) = 1 \quad \forall x, y \in U$$

(4)转置关系或逆关系 R^T :

$$R^T \Leftrightarrow \mu_{R^T}(x, y) = \mu_R(y, x) \quad \forall x, y \in U$$

2.模糊关系的合成及其性质

(1)模糊关系的合成

定义：设 U 、 V 、 W 是三个集合， $\widetilde{R}_1$ 是 U 到 V 上的模糊关系， $\widetilde{R}_2$ 是 V 到 W 上的模糊关系，则称 $\widetilde{R}_1 \circ \widetilde{R}_2$ 为关系 $\widetilde{R}_1$ 与 $\widetilde{R}_2$ 的合成，且规定它为 U 到 W 上的模糊关系，其隶属函数为：

$$\mu_{\widetilde{R}_1 \circ \widetilde{R}_2}(x, z) = \bigvee [\mu_{\widetilde{R}_1}(x, y) \wedge \mu_{\widetilde{R}_2}(y, z)]$$

(2) 模糊关系合成的基本性质

① 结合律： $(\widetilde{R}_1 \circ \widetilde{R}_2) \circ \widetilde{R}_3 = \widetilde{R}_1 \circ (\widetilde{R}_2 \circ \widetilde{R}_3)$

② $\widetilde{I} \circ \widetilde{R} = \widetilde{R} \circ \widetilde{I} = \widetilde{R}$

③ $\widetilde{O} \circ \widetilde{R} = \widetilde{R} \circ \widetilde{O} = \widetilde{O}$

④ 若 $\widetilde{R}_1 \subseteq \widetilde{R}_2$ ，则有：

$$\widetilde{R} \circ \widetilde{R}_1 \subseteq \widetilde{R} \circ \widetilde{R}_2, \quad \widetilde{R}_2 \circ \widetilde{R}$$

⑤ $(\widetilde{R}_1 \circ \widetilde{R}_2)^T = \widetilde{R}_2^T \circ \widetilde{R}_1^T$

⑥ $\widetilde{R} \circ (\widetilde{R}_1 \widetilde{Y} \widetilde{R}_2) = (\widetilde{R} \circ \widetilde{R}_1) \widetilde{Y} (\widetilde{R} \circ \widetilde{R}_2), (\widetilde{R}_1 \widetilde{Y} \widetilde{R}_2) \circ \widetilde{R} = (\widetilde{R}_1 \circ \widetilde{R}) \widetilde{Y} (\widetilde{R}_2 \circ \widetilde{R})$

⑦ $\widetilde{R} \circ (\widetilde{R}_1 \widetilde{I} \widetilde{R}_2) = (\widetilde{R} \circ \widetilde{R}_1) \widetilde{I} (\widetilde{R} \circ \widetilde{R}_2), (\widetilde{R}_1 \widetilde{I} \widetilde{R}_2) \circ \widetilde{R} = (\widetilde{R}_1 \circ \widetilde{R}) \widetilde{I} (\widetilde{R}_2 \circ \widetilde{R})$

⑧ $(\widetilde{R}_1 \widetilde{Y} \widetilde{R}_2)^T = \widetilde{R}_1^T \widetilde{Y} \widetilde{R}_2^T$

⑨ $(\widetilde{R}_1 \widetilde{I} \widetilde{R}_2)^T = \widetilde{R}_1^T \widetilde{I} \widetilde{R}_2^T$

⑩ $(\widetilde{R}^T)^T = \widetilde{R}$

3. 模糊相似关系与模糊等价关系

(1) 模糊相似关系

定义：设 $\widetilde{R}$ 是 U 中的模糊关系，若它满足如下性质：

① 自反性： $\forall x \in U, \mu_{\widetilde{R}}(x, x) = 1$

② 对称性： $\forall x, y \in U, \mu_{\widetilde{R}}(x, y) = \mu_{\widetilde{R}}(y, x)$ 则称 $\widetilde{R}$ 为 U 中的模

糊相似关系。

(2) 模糊等价关系

定义：设 $\widetilde{R}$ 是 U 中的模糊相似关系，若它满足传递性，即 $\forall x, y, z \in U, \bigvee [\mu_{\widetilde{R}}(x, z) \wedge \mu_{\widetilde{R}}(z, y)] \leq \mu_{\widetilde{R}}(x, y)$ ，则称 $\widetilde{R}$ 为 U 中的模糊等价关系。

(二) 模糊变换

定义：设 $\widetilde{R}$ 是一个给定的模糊矩阵

$$R = (r_{ij})_{m \times n} = \begin{pmatrix} r_{11} & r_{12} & \cdots & r_{1n} \\ r_{21} & r_{22} & \cdots & r_{2n} \\ \cdots & \cdots & \cdots & \cdots \\ r_{m1} & r_{m2} & \cdots & r_{mn} \end{pmatrix} \quad (10)$$

在(10)式中, $0 \leq r_{ij} \leq 1 (i=1, 2, \dots, m; j=1, 2, \dots, n)$; $\underline{\underline{A}}$ 是一个给定的模糊向量

$$\underline{\underline{A}} = [a_1, a_2, \dots, a_m] \quad (11)$$

在(11)式中, $0 \leq a_i \leq 1 (i=1, 2, \dots, m)$ 。则称 $\underline{\underline{A}}$ 与 $\underline{\underline{B}}$ 的合成运算

$$\begin{aligned} \underline{\underline{B}} &= [b_1, b_2, \dots, b_n] = \underline{\underline{A}} \circ \underline{\underline{R}} \\ &= \left[\bigvee_{k=1}^m (a_k \wedge r_{k1}), \bigvee_{k=1}^m (a_k \wedge r_{k2}), \dots, \bigvee_{k=1}^m (a_k \wedge r_{kn}) \right] \end{aligned} \quad (12)$$

为模糊变换。显然, 在(12)式中, 有 $0 \leq b_j \leq 1 (j=1, 2, \dots, n)$ 。

第二节 模糊聚类分析方法

模糊聚类分析，是从模糊集的观点来探讨事物的数量分类的一类方法。近年来，模糊聚类分析方法在地理分区与地理事物分类研究中得到了广泛地应用。本节，我们将主要介绍基于模糊等价关系与基于最大模糊支撑树的模糊聚类分析方法在地理分区和地理事物分类中的应用。

一、基于模糊等价关系的模糊聚类分析方法

基于模糊等价关系的模糊聚类分析方法的基本思想是：由于模糊等价关系 $\tilde{R}$ 是论域集 U 与自己的直积 $U \times U$ 上的一个模糊子集，因此可以对 $\tilde{R}$ 进行

分解，当用 λ - 水平对 $\tilde{R}$ 作截集时，截得的 $U \times U$ 的普通子集 R_λ 就是 U

上的一个普通等价关系，也就得到了关于 U 中被分类对象元素的一种分类。当 λ 由 1 下降到 0 时，所得的分类由细变粗，逐渐归并，从而形成一个动态聚类谱系图。由此可见，分类对象集 U 上的模糊等价关系 $\tilde{R}$ 的建立是这种

聚类分析方法中的一个关键性的环节。

(一)建立模糊等价关系

为了建立分类对象集合 U 上的模糊等价关系 $\tilde{R}^*$ ，通常需要首先计算各个分类对象之间的相似性统计量，建立分类对象集合 U 上的模糊相似关系 $\tilde{R}$ 。

1. 模糊相似关系的建立 关于各分类对象之间相似性统计量 r_{ij} 的计算，除了采用夹角余弦公式和相似系数计算公式(分别见第二章第三节中(10)和(11)式)以外，还可以采用如下几个计算公式。

(1)数量积法：

$$r_{ij} = \begin{cases} 1 & \text{当 } i = j \\ \sum_{k=1}^n x_{ik} x_{jk} / M & \text{当 } i \neq j \end{cases} \quad (i, j = 1, 2, \dots, m) \quad (1)$$

在(1)式中， M 是一个适当选择之正数，一般而言，它应满足：

$$M > \max_{i \neq j} \left\{ \sum_{k=1}^n x_{ik} x_{jk} \right\}$$

(2)绝对值差数法：

$$r_{ij} = \begin{cases} 1 & \text{当 } i = j \\ 1 - c \sum_{k=1}^n |x_{ik} - x_{jk}| & \text{当 } i \neq j \end{cases} \quad (i, j = 1, 2, \dots, m) \quad (2)$$

在(2)式中， c 为适当选择之正数，使 $0 \leq r_{ij} < 1 (i \neq j)$ 。

(3)最大最小值法：

$$r_{ij} = \frac{\sum_{k=1}^n \min(x_{ik}, x_{jk})}{\sum_{k=1}^n \max(x_{ik}, x_{jk})} \quad (i, j = 1, 2, \dots, m) \quad (3)$$

(4)算术平均最小法：

$$r_{ij} = \frac{\sum_{k=1}^n \min(x_{ik}, x_{jk})}{\frac{1}{2} \sum_{k=1}^n (x_{ik} + x_{jk})} \quad (i, j = 1, 2, \dots, m) \quad (4)$$

(5)绝对值指数法：

$$r_{ij} = e^{-\sum_{k=1}^n |x_{ik} - x_{jk}|} \quad (i, j = 1, 2, \dots, m) \quad (5)$$

(6)指数相似系数法：

$$r_{ij} = \frac{1}{n} \sum_{k=1}^n e^{-\frac{3}{4} \frac{(x_{ik} - x_{jk})^2}{s_k^2}} \quad (i, j = 1, 2, \dots, m) \quad (6)$$

在(6)式中， s_k 是第 k 个指标的方差，即

$$s_k = \sqrt{\frac{1}{m} \sum_{i=1}^m (x_{ik} - \bar{x}_k)^2}$$

2.将模糊相似关系 $\tilde{R}$ 改造为模糊等价关系 $\tilde{R}^*$ 。通过上节的介绍，我们知道，模糊相似关系 $\tilde{R}$ 满足自反性和对称性，但一般而言，它并不满足传递性，也就是说它并不是模糊等价关系。因此，为了聚类，我们必须采用传递闭合的性质将这种模糊相似关系 $\tilde{R}$ 改造为模糊等价关系 $\tilde{R}^*$ 。改造的办法是将 $\tilde{R}$ 自乘，即

$$\tilde{R}_2 = \tilde{R} \circ \tilde{R}$$

$$\tilde{R}^4 = \tilde{R}^2 \circ \tilde{R}^2$$

这样下去，就必然会存在一个自然数 K ，使得：

$$\tilde{R}^{2k} = \tilde{R}^k \circ \tilde{R}^k = \tilde{R}^k$$

这时， $\tilde{R}^* = \tilde{R}^k$ 便是一个模糊等价关系了。

显然，对于第二章中表 2-12 所描述的九个农业区域，用夹角余弦公式计算所得的相似系数矩阵

$$\widetilde{R} = \begin{bmatrix} 1 & 0.88 & 0.49 & 0.88 & 0.30 & 0.24 & 0.20 & 0.93 & 0.77 \\ 0.88 & 1 & 0.38 & 0.94 & 0.06 & 0.05 & 0.01 & 0.95 & 0.93 \\ 0.49 & 0.38 & 1 & 0.67 & 0.76 & 0.80 & 0.71 & 0.45 & 0.55 \\ 0.88 & 0.94 & 0.67 & 1 & 0.30 & 0.30 & 0.24 & 0.92 & 0.95 \\ 0.30 & 0.06 & 0.76 & 0.30 & 1 & 0.99 & 0.98 & 0.21 & 0.21 \\ 0.24 & 0.05 & 0.80 & 0.30 & 0.99 & 1 & 0.99 & 0.18 & 0.23 \\ 0.20 & 0.01 & 0.71 & 0.24 & 0.98 & 0.99 & 1 & 0.14 & 0.19 \\ 0.93 & 0.95 & 0.45 & 0.92 & 0.21 & 0.18 & 0.14 & 1 & 0.90 \\ 0.77 & 0.93 & 0.55 & 0.95 & 0.21 & 0.23 & 0.19 & 0.90 & 1 \end{bmatrix}$$

就是这九个农业区域所构成的分类对象集合上的一个模糊相似关系，经过自乘计算后可以验证：

$$\blacksquare \widetilde{R} = \widetilde{R} \circ \widetilde{R} = \widetilde{R}$$

$$\widetilde{R}^* = \widetilde{R} \circ \widetilde{R} \circ \widetilde{R} \circ \widetilde{R} = \widetilde{R}^4 = \begin{bmatrix} 1 & 0.93 & 0.67 & 0.93 & 0.67 & 0.67 & 0.67 & 0.93 & 0.93 \\ 0.93 & 1 & 0.67 & 0.94 & 0.67 & 0.67 & 0.67 & 0.95 & 0.94 \\ 0.67 & 0.67 & 1 & 0.67 & 0.80 & 0.80 & 0.80 & 0.67 & 0.67 \\ 0.93 & 0.94 & 0.67 & 1 & 0.67 & 0.67 & 0.67 & 0.94 & 0.95 \\ 0.67 & 0.67 & 0.80 & 0.67 & 1 & 0.99 & 0.99 & 0.67 & 0.67 \\ 0.67 & 0.67 & 0.80 & 0.67 & 0.99 & 1 & 0.99 & 0.67 & 0.67 \\ 0.67 & 0.67 & 0.80 & 0.67 & 0.99 & 0.99 & 1 & 0.67 & 0.67 \\ 0.93 & 0.95 & 0.67 & 0.94 & 0.67 & 0.67 & 0.67 & 1 & 0.94 \\ 0.93 & 0.94 & 0.67 & 0.95 & 0.67 & 0.67 & 0.67 & 0.94 & 1 \end{bmatrix}$$

即： $\widetilde{R}^*$ 是一个模糊等价关系。

(二)在不同的截集水平下进行聚类

用上述模糊等价关系 $\widetilde{R}^*$ ，在不同的截集水平下聚类，得到如下聚类

结果：

(1)取 $\lambda=1$ ，得：

$$\widetilde{R}_1^* = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

在 $\widetilde{R}_1^*$ 中，由于各行均不相同，故 $G_1, G_2, G_3, G_4, G_5, G_6, G_7, G_8, G_9$ 各自成为一类。

(2)取 $\lambda=0.99$ ，得：

$$R_{0.99}^* \sim \begin{Bmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{Bmatrix}$$

在 $R_{0.99}^*$ 中，由于第 5、6、7 行相同，而其它各行均不相同，故将 G_5, G_6, G_7 归并为一类，而 $G_1, G_2, G_3, G_4, G_8, G_9$ 各自成为一类。

(3) 取 $\lambda = 0.95$ ，得：

$$R_{0.95}^* \sim = \begin{Bmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \end{Bmatrix}$$

在 $R_{0.95}^*$ 中，由于第 2、8 行相同，第 4、9 行相同，第 5、6、7 行相同，而第 1

行与第 3 行和其它各行均不相同，故 G_2 与 G_8 聚为一类， G_4 与 G_9 聚为一类， G_5, G_6, G_7 聚为一类，而 G_1 和 G_3 各自成为一类。

(4) 取 $\lambda = 0.94$ ，得：

$$R_{0.94}^* \sim = \begin{Bmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \end{Bmatrix}$$

在 $R_{0.94}^*$ 中，由于第 2、4、8、9 行相同，第 5、6、7 行相同，第 1 行与第 3 行和其它各行均不相同，故 G_2, G_4, G_8, G_9 聚为一类， G_5, G_6, G_7 聚为一类， G_1 和 G_3 各自聚为一类。

(5) 取 $\lambda = 0.93$ ，得：

$$\underset{\sim}{R}_{0.93}^{\ast} = \begin{cases} 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \end{cases}$$

在 $\underset{\sim}{R}_{0.93}^{\ast}$ 中, 由于第1、2、4、8、9行相同, 第5、6、7行相同, 第3行与其它各行均不相同, 故 G_1 、 G_2 、 G_4 、 G_8 、 G_9 聚为一类, G_5 、 G_6 、 G_7 聚为一类, G_3 各自成为一类。

(6)取 $\lambda=0.80$, 得:

$$\underset{\sim}{R}_{0.8}^{\ast} = \begin{cases} 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 \end{cases}$$

在 $\underset{\sim}{R}_{0.8}^{\ast}$ 中, 由于第1、2、4、8、9行相同, 第3、5、6、7行相同, 故 G_1 、 G_2 、

G_4 、 G_8 、 G_9 聚为一类, G_3 、 G_5 、 G_6 、 G_7 聚为一类。

(7)取 $\lambda=0.67$, 得:

$$\underset{\sim}{R}_{0.67}^{\ast} = \begin{cases} 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \end{cases}$$

在 $\underset{\sim}{R}_{0.67}^{\ast}$ 中, 由于各行都相同, 故 G_1 、 G_2 、 G_3 、 G_4 、 G_5 、 G_6 、 G_7 、 G_8 、

G_9 均聚为一类。

综合上述聚类结果, 可以作出如下聚类谱系图, 如图 9-3 所示。

二、基于最大模糊支撑树的模糊聚类分析方法

除了依据模糊等价关系进行聚类分析外，还可以应用最大模糊支撑树进行聚类分析。基于最大模糊支撑树的聚类分析过程，可按如下步骤进行。

第一步：建立分类对象集上的模糊相似关系，构造模糊图。这一步骤的工作可按如下作法进行：

(1) 计算各个分类对象之间的相似性统计量 $r_{ij} (i, j=1, 2, \dots, m)$ ，建立分类对象集 U 上的模糊相似关系 $\tilde{R} = (r_{ij})_{m \times m}$ 。

(2) 将 $\tilde{R}$ 表示成一个由 m 个结点所构成的模糊图 $G = (V, E)$ ，使 G 中的任意两个结点 V_i 与 V_j 之间都有一条边相联结，且赋该边的权值为 r_{ij} 。

假若，对于某五个地理区域所构成的分类对象集合 $V = \{v_1, v_2, v_3, v_4, v_5\}$ ，经过选择聚类要素并对其原始数据进行标准化处理后，计算各分类对象之间的相似性统计量，得到如下的模糊相似关系

$$\tilde{R} = \begin{bmatrix} 1 & 0.7 & 0.6 & 0.1 & 0.3 \\ 0.7 & 1 & 0.7 & 0.3 & 0.8 \\ 0.6 & 0.7 & 1 & 0.4 & 0.9 \\ 0.1 & 0.3 & 0.4 & 1 & 0.1 \\ 0.3 & 0.8 & 0.9 & 0.1 & 1 \end{bmatrix}$$

则按照上述作法，可以将其表示成一个模糊图，如图 9-4 所示。

第二步：构造最大模糊支撑树。构造模糊图 G 上的最大支撑树的算法，可按下述作法进行：

(1) 找出 G 中最大权值的边 r_{ij} ；

(2) 将 r_{ij} 存放在集合 C 中，将 r_{ij} 边上的新结点放入集合 T 中，若 T 中已含有所有 m 个结点时，转(4)；

(3) 检查 T 中每一个结点与 T 外的结点组成的边的权值，找出其中最大者 r_{ij} ，转至(2)；

(4) 结束，此时 G 中的边就构成了 G 的最大模糊支撑树 $T_{\max}$ 。

对于图 9-4 所示的模糊图 G ，按照上述算法，可以求出其最大模糊支撑树 $T_{\max}$ ，如图 9-5 所示。

可以证明， $T_{\max}$ 具有下述三个特点：①它不存在回路，所以是树；②它对原图 G 中所有结点都是连通的，所以它是图 G 的支撑树；③对于 G 的其它任何支撑树 T ，都有： $T_{\max}$ 中各边的权值之和大于或等于 T 中各边的权值之和。所以， $T_{\max}$ 的确是 G 的最大模糊支撑树。

第三步：由最大模糊支撑树进行聚类分析。其具体作法是：选择某一个 λ 值作截集，将 $T_{\max}$ 中小于 λ 的边断开，使相连的各结点构成一类，当 λ 由 1 下降到 0 时，所得的分类由细变粗，各结点所代表的分类对象逐渐归并，从而形成一个动态聚类谱系图。

譬如，对于图 9-5 所示的 G 的最大模糊支撑树 $T_{\max}$ ，当分别选取 $\lambda = 1$ ， $\lambda = 0.9$ ， $\lambda = 0.8$ ， $\lambda = 0.7$ ， $\lambda = 0.4$ 时，就可以得出不同的分类结果，这一过程所形成的聚类谱系图如图 9-6 所示。

第三节 模糊综合评判方法

模糊综合评判方法，是一种运用模糊数学原理分析和评价具有“模糊性”的事物的系统分析方法。它是一种以模糊推理为主的定性与定量相结合、精确与非精确相统一的分析评价方法。由于这种方法在处理各种难以用精确数学方法描述的复杂系统问题方面所表现出的独特的优越性，近年来已在许多学科领域中得到了十分广泛的应用。

一、模糊综合评判模型

(一)单层次模糊综合评判模型

给定两个有限论域

$$U = \{u_1, u_2, \dots, u_m\} \quad (1)$$

$$V = \{v_1, v_2, \dots, v_n\} \quad (2)$$

(1)式中， U 代表所有的评判因素所组成的集合；(2)式中， V 代表所有的评语等级所组成的集合。

如果着眼于第 i ($i=1, 2, \dots, m$) 个评判因素 u_i ，其单因素评判结果为 $R_i = [r_{i1}, r_{i2}, \dots, r_{in}]$ ，则 m 个评判因素的评判决策矩阵为

$$R = \begin{pmatrix} R_1 \\ R_2 \\ \vdots \\ R_m \end{pmatrix} = \begin{pmatrix} r_{11} & r_{12} & \cdots & r_{1n} \\ r_{21} & r_{22} & \cdots & r_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ r_{m1} & r_{m2} & \cdots & r_{mn} \end{pmatrix} \quad (3)$$

就是 U 到 V 上的一个模糊关系。

如果对各评判因素的权数分配为： $\tilde{A} = [a_1, a_2, \dots, a_m]$ (显然， $\tilde{A}$

是论域 U 上的一个模糊子集，且 $0 \leq a_i \leq 1, \sum_{i=1}^m a_i = 1$)，则应用模糊变换的合成运算，可以得到论域 V 上的一个模糊子集，即综合评判结果：

$$\tilde{B} = \tilde{A} \circ \tilde{R} = [b_1, b_2, \dots, b_n] \quad (4)$$

(二)多层次模糊综合评判模型

在复杂大系统中，需要考虑的因素往往是很多的，而且因素之间还存在着不同的层次。这时，应用单层次模糊综合评判模型就很难得出正确的评判结果。所以，在这种情况下，就需要将评判因素集合按照某种属性分成几类，先对每一类进行综合评判，然后再对各类评判结果进行类之间的高层次综合评判。这样，就产生了多层次模糊综合评判问题。

多层次模糊综合评判模型的建立，可按以下步骤进行：

(1)对评判因素集合 U ，按某个属性 c ，将其划分成 m 个子集，使它们满足：

$$\begin{cases} \sum_{i=1}^m U_i = U \\ U_i \cap U_j = \Phi \quad (i \neq j) \end{cases} \quad (5)$$

这样，就得到了第二级评判因素集合：

$$U/c = \{U_1, U_2, \dots, U_m\} \quad (6)$$

在(6)式中, $U_i = \{u_{ik}\}$ ($i=1, 2, \dots, m; k=1, 2, \dots, n_k$)表示子集 U_i 中含有 n_k 个评判因素。

(2)对于每一个子集 U_i 中的 n_k 个评判因素,按单层次模糊综合评判模型进行评判。如果 U_i 中诸因素的权数分配为 $\tilde{A}_i$, 其评判决策矩阵为 $\tilde{R}_i$, 则得

到第 i 个子集 U_i 的综合评判结果:

$$\tilde{B}_i = \tilde{A}_i \circ \tilde{R}_i = [b_{i1}, b_{i2}, \dots, b_{in}] \quad (7)$$

(3)对 U/c 中的 m 个评判因素子集 $U_i (i=1, 2, \dots, m)$, 进行综合评判, 其评判决策矩阵为:

$$\tilde{R} = \begin{pmatrix} \tilde{B}_1 \\ \tilde{B}_2 \\ \vdots \\ \tilde{B}_m \end{pmatrix} = \begin{pmatrix} b_{11} & b_{12} & \cdots & b_{1n} \\ b_{21} & b_{22} & \cdots & b_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ b_{m1} & b_{m2} & \cdots & b_{mn} \end{pmatrix} \quad (8)$$

如果 U/c 中的各因素子集的权数分配为 $\tilde{A}$, 则可得综合评判结果:

$$\tilde{B}^* = \tilde{A} \circ \tilde{R} \quad (9)$$

在(9)式中, $\tilde{B}^*$ 既是 U/c 的综合评判结果,也是 U 中的所有评判因素的综

合评判结果。这里需要强调的是,在(7)或(9)式中,矩阵合成运算的方法通常有两种:一是主因素决定模型法,即利用逻辑算子 $M(\wedge, \vee)$ 进行取大或取小合成,该方法一般仅适合于单项最优的选择;二是普通矩阵模型法,即利用普通矩阵算法进行运算,这种方法兼顾了各方面的因素,因此适宜于多因素的排序。

若 U/c 中仍含有很多因素,则可以对它再进行划分,得到三级以至更多层次的模糊综合评判模型。多层次的模糊综合评判模型,不仅可以反映评判因素的不同层次,而且避免了由于因素过多而难于分配权重的弊病。

二、模糊综合评判实例

作为模糊综合评判方法的应用实例,以下我们将探讨农业生态经济系统功能综合评价问题。

农业生态经济系统,是一类多要素的复杂系统,其内部诸要素之间的相互作用关系及各要素对系统功能的影响程度在量上是难以精确衡量的,即系统具有“模糊性”特征;其次,农业生态经济系统也还是一个包含着若干不同生产层次(或若干子系统)的复合系统,其系统功能从整体上来说是一种综合功能,具有“多属性”特点。因此,农业生态经济系统功能评价是一种多属性或多准则评价问题。这就要求评价者必须根据评价问题的性质、目标、要求等选择适宜的评价模型和方法。在这方面,模糊综合评判模型为我们提供了一种有效的方法。

(一)评价要素指标体系的设置

评价要素指标体系的设置,是对农业生态经济系统功能进行综合评价的前提和基础。指标体系设置得是否合理和准确,直接影响着评价结果的科学

性、可靠性和准确性。因此，农业生态经济系统功能综合评价的首要任务就是根据评价对象的性质、评价目标以及评价决策要求等，建立能够全面、准确地反映评价问题全貌的综合评价要素指标体系。

农业生态经济系统功能，是一种综合性功能，它主要由经济效益、生态效益和社会效益三个方面来反映。所以，对农业生态经济系统功能的考察及评价，必须立足于这三个基本方面。而这三个方面的效益又是由不同的要素来体现的，每一种要素都有表征其属性特征的指标。这些要素指标的组合就构成了农业生态经济系统功能综合评价的指标体系(见图 9-7)。

由图 9-7 可知，评价要素集合为：

$$U = \{u_1, u_2, u_3\}$$

其中，各单要素子集 $u_i (i=1, 2, 3)$ 分别为：

$$U_1 = \{u_{11}, u_{12}, u_{13}, u_{14}, u_{15}\}$$

$$U_2 = \{u_{21}, u_{22}, u_{23}, u_{24}, u_{25}\}$$

$$U_3 = \{u_{31}, u_{32}, u_{33}, u_{34}, u_{35}\}$$

(二)评语集合的确定

根据评价决策的实际需要，将评判等级标准划分为“好”、“较好”、“一般”、“较差”和“差”五个等级。即评语集合为：

$$V = \{v_1, v_2, v_3, v_4, v_5\}$$

$$= \{\text{好, 较好, 一般, 较差, 差}\}$$

(三)评价要素权重子集的确定

在上述农业生态经济系统功能综合评价指标体系中，由于下层各指标对上层某一指标的相对重要程度并非一样，即一些指标的影响程度要大于或超过另一些指标。因此，为了衡量下层各指标对上层指标的相对重要性，需要确定评价指标的权重系数。常见的确定权重系数的方法有：(1)主观经验判断法；(2)专家调查法或专家征询法；(3)评判专家小组集体讨论投票表决法；(4)层次分析法(即 AHP 方法)。为了保证确定的权重系数的客观性、公正性和科学性，常常可将上述几种方法结合起来使用。以下采用主观经验判断法和专家征询法相结合来确定各级评价要素指标的权重系数子集。

各子集权重(一级权重)为：

$$\tilde{A} = [a_1, a_2, a_3]$$

各子集 $U_i (i=1, 2, 3)$ 中诸要素的权重(二级权重)分别为：

$$A_1 = [a_{11}, a_{12}, a_{13}, a_{14}, a_{15}]$$

~

$$A_2 = [a_{21}, a_{22}, a_{23}, a_{24}, a_{25}]$$

~

$$A_3 = [a_{31}, a_{32}, a_{33}, a_{34}, a_{35}]$$

~

(四)评判的实施

所谓评判的实施，就是根据评判对象——农业生态经济系统的各种实际调查访问材料、各种试验与研究数据，采用模糊数学和精确数学方法对各个评价指标进行定量估算，然后由评判专家小组的每一个成员根据已确定的评价等级标准依次对各个指标进行评价。假定评判专家小组有 20 名成员，其中

有 7 名对系统生态效益功能的评价指标之一“水土流失状况(u_{22})”同意“较好(v_2)”的评价等级，即持同意意见的专家占专家小组总人数的 7/20，因此该指标的评价值就是 0.35。依次类推，可分别得出各子集 $u_i (i=1, 2, 3)$ 中单要素的评价决策矩阵 $R_i (i=1, 2, 3)$ 为：

$$\begin{aligned} \tilde{R}_1 &= \begin{pmatrix} r_{111} & r_{112} & r_{113} & r_{114} & r_{115} \\ r_{121} & r_{122} & r_{123} & r_{124} & r_{125} \\ r_{131} & r_{132} & r_{133} & r_{134} & r_{135} \\ r_{141} & r_{142} & r_{143} & r_{144} & r_{145} \\ r_{151} & r_{152} & r_{153} & r_{154} & r_{155} \end{pmatrix} = (r_{1ij})_{5 \times 5} \\ \tilde{R}_2 &= \begin{pmatrix} r_{211} & r_{212} & r_{213} & r_{214} & r_{215} \\ r_{221} & r_{222} & r_{223} & r_{224} & r_{225} \\ r_{231} & r_{232} & r_{233} & r_{234} & r_{235} \\ r_{241} & r_{242} & r_{243} & r_{244} & r_{245} \\ r_{251} & r_{252} & r_{253} & r_{254} & r_{255} \end{pmatrix} = (r_{2ij})_{5 \times 5} \\ \tilde{R}_3 &= \begin{pmatrix} r_{311} & r_{312} & r_{313} & r_{314} & r_{315} \\ r_{321} & r_{322} & r_{323} & r_{324} & r_{325} \\ r_{331} & r_{332} & r_{333} & r_{334} & r_{335} \\ r_{341} & r_{342} & r_{343} & r_{344} & r_{345} \\ r_{351} & r_{352} & r_{353} & r_{354} & r_{355} \end{pmatrix} = (r_{3ij})_{5 \times 5} \end{aligned}$$

然后由各单要素的权重系数向量 $\tilde{A}_i$ 和评价决策矩阵 $\tilde{R}_i$ ，利用合成

运算法则经过合成运算即可得到：

$$\tilde{B}_i = \tilde{A}_i \circ \tilde{R}_i = [b_{i1}, b_{i2}, b_{i3}, b_{i4}, b_{i5}] \quad (i = 1, 2, 3)$$

基于单要素模糊综合评判结果 $\tilde{B}_i$ ，可以得到 U 中各子集的综合评价决策矩阵：

$$\tilde{R} = \begin{pmatrix} \tilde{B}_1 \\ \tilde{B}_2 \\ \tilde{B}_3 \end{pmatrix} = \begin{pmatrix} b_{11} & b_{12} & b_{13} & b_{14} & b_{15} \\ b_{21} & b_{22} & b_{23} & b_{24} & b_{25} \\ b_{31} & b_{32} & b_{33} & b_{34} & b_{35} \end{pmatrix}$$

最后再由 U 的各子集的权重系数向量 $\tilde{A}$ 和综合评价决策矩阵 $\tilde{R}$ ，经过合成运算，即得出对农业生态经济系统功能的模糊综合评价结果：

$$\begin{aligned} \tilde{B} &= \tilde{A} \circ \tilde{R} = [a_1, a_2, a_3] \begin{pmatrix} \tilde{B}_1 \\ \tilde{B}_2 \\ \tilde{B}_3 \end{pmatrix} \\ &= [b_1, b_2, b_3, b_4, b_5] \end{aligned}$$

(五)评价实例计算

对于某农业生态经济系统，经评判专家小组测评结果，分别得各子集 $u_i (i=1, 2, 3)$ 中诸要素的评价决策矩阵：

$$\underset{\sim}{R_1} = \begin{pmatrix} 0.10 & 0.40 & 0.35 & 0.10 & 0.05 \\ 0.25 & 0.35 & 0.25 & 0.15 & 0.00 \\ 0.20 & 0.25 & 0.35 & 0.10 & 0.10 \\ 0.15 & 0.35 & 0.30 & 0.10 & 0.10 \\ 0.12 & 0.36 & 0.32 & 0.14 & 0.06 \end{pmatrix}$$

$$\underset{\sim}{R_2} = \begin{pmatrix} 0.10 & 0.35 & 0.35 & 0.10 & 0.10 \\ 0.15 & 0.45 & 0.25 & 0.10 & 0.05 \\ 0.10 & 0.40 & 0.35 & 0.10 & 0.05 \\ 0.20 & 0.25 & 0.35 & 0.15 & 0.05 \\ 0.12 & 0.38 & 0.26 & 0.09 & 0.12 \end{pmatrix}$$

$$\underset{\sim}{R_3} = \begin{pmatrix} 0.10 & 0.40 & 0.30 & 0.10 & 0.10 \\ 0.08 & 0.20 & 0.50 & 0.12 & 0.10 \\ 0.14 & 0.26 & 0.38 & 0.14 & 0.08 \\ 0.10 & 0.50 & 0.35 & 0.05 & 0.00 \\ 0.11 & 0.37 & 0.37 & 0.10 & 0.05 \end{pmatrix}$$

采用主观经验判断法和专家征询法相结合的方法，可得：

$$\underset{\sim}{A} = [0.38, 0.34, 0.28]$$

$$\underset{\sim}{A_1} = [0.20, 0.18, 0.19, 0.15, 0.28]$$

$$\underset{\sim}{A_2} = [0.25, 0.28, 0.20, 0.14, 0.13]$$

$$\underset{\sim}{A_3} = [0.21, 0.12, 0.30, 0.25, 0.12]$$

采用普通矩阵乘法，经过合成运算，得各子集 $u_i (i=1, 2, 3)$ 的综合评判结果分别为：

$$\underset{\sim}{B_1} = \underset{\sim}{A_1} \cdot \underset{\sim}{R_1} = [0.16, 0.34, 0.32, 0.12, 0.06]$$

$$\underset{\sim}{B_2} = \underset{\sim}{A_2} \cdot \underset{\sim}{R_2} = [0.13, 0.38, 0.31, 0.11, 0.07]$$

$$\underset{\sim}{B_3} = \underset{\sim}{A_3} \cdot \underset{\sim}{R_3} = [0.11, 0.36, 0.37, 0.10, 0.06]$$

因此， U 中各子集的综合评价决策矩阵为：

$$\underset{\sim}{R} = \begin{pmatrix} \underset{\sim}{R_1} \\ \underset{\sim}{R_2} \\ \underset{\sim}{R_3} \end{pmatrix} = \begin{pmatrix} 0.16 & 0.34 & 0.32 & 0.12 & 0.06 \\ 0.13 & 0.38 & 0.31 & 0.11 & 0.07 \\ 0.11 & 0.36 & 0.37 & 0.10 & 0.06 \end{pmatrix}$$

所以，该农业生态经济系统功能模糊综合评价结果为：

$$\underset{\sim}{B} = \underset{\sim}{A} \circ \underset{\sim}{R} = [0.16, 0.34, 0.32, 0.12, 0.07]$$

将其归一化得：

$$\underset{\sim}{B} = [0.158, 0.337, 0.317, 0.119, 0.069]$$

上述评价结果表明，该农业生态经济系统功能还是较好的。

第四节 应用实例

——甘肃黄土高原区农业生态气候条件分析

一、农业生态气候模型

一个地区的气候特征，主要反映在热量和水分的组合关系上。这是大多数气候分类方法的理论基础。农业生态气候是一个复杂的动态系统，其在时间和空间上的变化，不仅使农业生产具有相应的季节性和地域性，而且还直接影响着农作物生长的快慢、质量的高低和产量的多少，并在一定程度上制约着农业生产结构及其空间布局。其次，作为重要生态因子的气候对农业的影响，就农区、牧区和林区而言，主要是分别指对农作物、牧草和森林生长的影响。这既涉及繁多的自然因素，同时也受到人类活动的干扰。所以，农业生态气候处于农业和气候两大系统的界面上，兼有自然因素和人文因素的共同影响，是一个具有复杂生物物理机制和生物化学变化的动态系统，难以用传统的气候学方法进行分类和评价。鉴于上述原因，我国学者顾恒岳、艾南山等先生(1984)曾提出和建立了农业生态气候适宜度理论。

该理论认为，各气候要素对作物生长的适宜度是定义在此气候要素定义域(取值范围)上，在 $[0, 1]$ 区间上取值的模糊子集。譬如，气温适宜度可以表示为：

$$S_T \subseteq [T_0, T_1] \quad (1)$$

$$S_T \triangleq \mu_{S_T}(T) \rightarrow [0, 1] \quad (2)$$

在(1)式中， S_T 表示气温要素的适宜度模糊子集；在(2)式中， $\mu_{S_T}(T)$ 表示气温适宜度——模糊子集 S_T 的隶属函数， S_T 表示气温要素 T 对 S_T 的隶属度。

由各气候要素的适宜度子集，可以诱导出各气候要素随时间 t 变化的适宜度过程，即气候适宜态。如果用 $S_T(t)$ 、 $S_R(t)$ 和 $S_I(t)$ 分别代表气温 T 、相对湿度 R 和日照 I 的适宜态，则农业生态气候的动态过程可用模糊向量

$$\underset{\sim}{S_c}(t) = \begin{pmatrix} S_T(t) \\ S_R(t) \\ S_I(t) \end{pmatrix} \quad (3)$$

表示。 $S_c(t)$ 可以是连续过程，也可以是离散过程。

关于农业生态气候适宜度，可以用下述一组模型描述。

1. 资源模型

$$\underset{\sim}{S_{c1}}(t) \triangleq \frac{1}{3} [\underset{\sim}{S_T}(t) + \underset{\sim}{S_R}(t) + \underset{\sim}{S_I}(t)] \quad (4)$$

连续过程：

$$\underset{\sim}{S_{c1}}(t) = \int_{t \in [0, t_0]} \frac{1}{3} [\underset{\sim}{S_T}(t) + \underset{\sim}{S_R}(t) + \underset{\sim}{S_I}(t)] / t \quad (5)$$

离散过程：

$$\underset{\sim}{S_{C1}}(t) = \sum_{j=1}^n \frac{1}{3} [S_T(t_j) + S_I(t_j) + S_R(t_j)] \quad (6)$$

资源模型表示光、水、热组合过程对作物生长可提供的气候

2. 效能模型

$$\underset{\sim}{S_{C2}}(t) \triangleq \underset{\sim}{S_T}(t) \underset{\sim}{\wedge} \underset{\sim}{S_R}(t) \underset{\sim}{\wedge} \underset{\sim}{S_I}(t) \quad (7)$$

连续过程：

$$\underset{\sim}{S_{C2}}(t) = \int_{t \in [0, t_0]} [\underset{\sim}{S_T}(t) \underset{\sim}{\wedge} \underset{\sim}{S_R}(t) \underset{\sim}{\wedge} \underset{\sim}{S_I}(t)] / t \quad (8)$$

离散过程：

$$\underset{\sim}{S_{C2}}(t) = \sum_{j=1}^n [\underset{\sim}{S_T}(t_j) \underset{\sim}{\wedge} \underset{\sim}{S_R}(t_j) \underset{\sim}{\wedge} \underset{\sim}{S_I}(t_j)] / t_j \quad (9)$$

效能模型反映了光、水热匹配程度及其对作物生长的适宜度过程。3. 结构模型

$$\underset{\sim}{S_{C3}}(t) \triangleq \underset{\sim}{\alpha_1 S_T}(t) + \underset{\sim}{\alpha_2 S_R}(t) + \underset{\sim}{\alpha_3 S_I}(t) \quad (10)$$

连续过程：

$$\underset{\sim}{S_{C3}}(t) = \int_{t \in [0, t_0]} [\underset{\sim}{\alpha_1 S_T}(t) + \underset{\sim}{\alpha_2 S_R}(t) + \underset{\sim}{\alpha_3 S_I}(t)] / t \quad (11)$$

离散过程：

$$\underset{\sim}{S_{C3}}(t) = \sum_{j=1}^n [\underset{\sim}{\alpha_1 S_T}(t_j) + \underset{\sim}{\alpha_2 S_R}(t_j) + \underset{\sim}{\alpha_3 S_I}(t_j)] / t_j \quad (12)$$

结构模型反映了农业生态气候过程的区域差异。权重 α_i ($i=1, 2, 3$) 的选择，在旱生和长日照作物区，可适当增大 α_2 ，减小 α_3 ；而在喜温喜湿作物区，则可酌情增大 α_1 和 α_2 ；在喜凉作物为主的地区，可适当减小 α_1 。

4. 农业生态气候指数

根据资源模型和效能模型，可用实测的多年平均气候资料推求农业生态气候资源指数 C_r ，效能指数 C_e 及气候利用系数 K_c 。

(1) 资源指数：

$$C_T = \int_0^{t_0} \underset{\sim}{S_{C1}}(t) dt = \frac{1}{3} \int_0^{t_0} [\underset{\sim}{S_T}(t) + \underset{\sim}{S_R}(t) + \underset{\sim}{S_I}(t)] dt \quad (13)$$

$$\text{或者} \quad C_T = \frac{1}{3} \sum_{j=1}^n [\underset{\sim}{S_T}(t_j) + \underset{\sim}{S_R}(t_j) + \underset{\sim}{S_I}(t_j)] \quad (14)$$

资源指数 C_r 表示气候资源的潜力， C_r 愈大，气候潜力愈大，资源愈丰富。

(2) 效能指数：

$$C_e = \int_0^{t_0} \underset{\sim}{S_{C2}}(t) dt = \int_0^{t_0} [\underset{\sim}{S_T}(t) \underset{\sim}{\wedge} \underset{\sim}{S_R}(t) \underset{\sim}{\wedge} \underset{\sim}{S_I}(t)] dt \quad (15)$$

$$\text{或者} \quad C_e = \sum_{j=1}^n [\underset{\sim}{S_T}(t_j) \underset{\sim}{\wedge} \underset{\sim}{S_R}(t_j) \underset{\sim}{\wedge} \underset{\sim}{S_I}(t_j)] \quad (16)$$

效能指数从作物生长的角度，反映了光、水、热的匹配程度， C_e 愈大，光、水、热匹配程度愈佳，愈有利于作物生长。

(3) 气候利用系数：

$$K=C_e/C_r \quad (17)$$

气候利用系数反映了农业气候资源在天然条件下被作物生长过程所利用的实际效率，K 愈大，利用率愈高。

二、甘肃黄土高原区农业生态气候条件分析

本节所研究的甘肃黄土高原区以六盘山和渭河源山地为界，可分为三个部分，六盘山以东地区，即通常所谓的陇东黄土高原，海拔较低，水分、热量均较充足；六盘山以西地区，通称陇中黄土高原，海拔平均较前者高 500 米以上，降水偏少，气温偏低；而渭河源山地以西部分，主要指临夏一带黄土高原，因临近青藏高原，水热条件的内部地域分界十分显著，总的特征是降水量较多而气温偏低。按行政区划分，包括庆阳地区、平凉地区、定西地区、兰州市、白银市所管辖的全部地区，及天水地区的北部和临夏回族自治区的东部地区，共四十个县(市)。我们选用本区内各主要气象台站多年气温、相对湿度和日照时数，取月平均值的离散过程。运用以上农业生态气候模型，计算了各台站所代表地区的气候资源指数，效能指数和利用系数，它们与多年平均的日照总时数，年降水量和 $\geq 10^{\circ}\text{C}$ 活动积温的比较结果列于表 9-1。

表 9-1 甘肃黄土高原区多年平均气候适宜度指数

序号	台站名	C_r	C_e	K	年日照总时数 F(小时)	年降水量 P(mm)	$\geq 10^{\circ}\text{C}$ 活动积温 Q($^{\circ}\text{C}$)
1	张家川	5.27	2.65	0.50	2888.0	606.5	2225.1
2	秦安	5.67	2.49	0.44	2068.8	507.7	3395.0
3	漳县	5.46	3.01	0.55	2925.7	456.0	2443.4
4	甘谷	5.85	3.00	0.51	2741.9	472.3	3262.6
5	清水	5.61	2.75	0.49	2815.8	574.9	2898.6
6	武山	5.36	2.67	0.50	2730.8	480.6	3084.5
7	天水台	5.81	2.94	9.51	2760.7	556.2	3359.5
8	天水县	5.80	3.20	0.52	2644.0	508.4	3536.9
9	平凉	5.23	2.60	0.50	2731.3	511.1	2862.8
10	静宁	5.34	2.71	0.51	2721.7	479.4	2539.2

(续表)

序号	台站名	Cr	Ce	K	年日照总时数 F(小时)	年降水量 P(mm)	≥ 10℃活动积温 Q(℃)
11	泾川	5.81	3.10	0.53	2619.8	549.9	3335.6
12	崇信	5.39	2.78	0.52	2723.0	546.6	3262.0
13	庄浪	5.34	2.72	0.51	2572.9	547.6	2677.4
14	华亭	5.56	2.91	0.52	2415.5	606.6	2694.7
15	灵台	5.43	2.91	0.54	2566.6	643.7	2832.6
16	靖远	4.34	1.20	0.28	2413.1	240.0	3224.4
17	会宁	4.19	1.63	0.39	2324.7	426.7	2088.5
18	定西	4.91	2.38	0.48	2449.0	429.2	2239.1
19	临洮	5.22	2.73	0.52	2553.9	578.0	2415.8
20	通渭	5.39	2.63	0.49	2483.9	440.3	2371.4
21	渭源	5.08	2.58	0.51	2450.8	525.6	1938.6
22	陇西	5.27	2.70	0.51	2491.7	444.7	2585.4
23	永靖	4.06	0.71	0.17	2552.2	304.2	2871.5
24	东乡	4.25	1.87	0.44	2392.6	537.6	1584.2
25	临夏台	5.10	2.74	0.54	2528.4	493.9	2328.5
26	广河	5.21	2.78	0.53	2426.2	502.3	2236.9
27	和政	5.36	2.85	0.53	2356.4	627.8	1810.9
28	康乐	5.20	2.85	0.55	2360.5	562.9	2104.3
29	环县	4.25	1.55	0.37	2033.9	417.7	3058.6
30	华池	4.85	2.31	0.48	2120.3	501.9	2896.7
31	庆阳	5.05	2.30	0.46	2190.4	533.8	3209.2
32	合水	4.69	2.25	0.48	2101.1	568.5	2998.0
33	镇原	5.03	2.49	0.49	2136.7	504.8	3191.1
34	正宁	4.79	2.47	0.52	2136.5	624.4	2739.2
35	宁县	5.23	2.76	0.53	1941.2	572.0	2929.9
36	景泰	3.66	0.25	0.07	1928.1	184.5	2988.7
37	永登	3.71	1.02	0.27	2087.6	290.6	2223.9
38	皋兰	3.84	0.82	0.21	2163.0	265.1	2798.3
39	兰州	4.23	1.22	0.29	2133.9	324.7	3242.0
40	榆中	4.12	1.51	0.37	1840.4	372.2	2370.9

从表 9-1 可以看出, 本区从南向北, 从东向西, 气候资源指数、效能指数和气候利用系数基本上呈递减趋势。这一变化充分反映农业生态气候过程的地域差异现象。譬如, 气候资源指数 Cr 在本区东南部泾川、天水、甘谷高达 5.80 以上, 漳县、秦安、华亭、灵台等地也在 5.50 左右, 可以说气候潜力较大, 气候资源比较丰富; 而在本区西北部景泰、永登、皋兰等地, Cr 只有 3.70 左右, 永靖、兰州、榆中等地, Cr 也只在 4.10 左右, 即气候潜力较小, 气候资源比较贫乏。就光、热、水匹配状况而言, 由东南向西北逐渐变差, 如漳县、甘谷、天水等地, 效能指数 Ce 均在 3.00 以上, 泾川可达 3.10;

而西北部的景泰， C_e 只有 0.25，皋兰只有 0.82，永登、兰州、靖远的 C_e 值也均在 1.25 以下。

众所周知，干旱缺水是甘肃黄土高原地区农业生产的限制条件。这一点，从以上的农业气候适宜度指数也得到了反映，如景泰年降水不足 200mm，而其 C_r 值仅 3.66， C_e 值只有 0.25；皋兰年降水只有 265.1mm，而其 C_r 值只有 3.84， C_e 值只有 0.82；永登年降水 290.6mm，而其 C_r 值只有 3.71， C_e 值只有 1.02。反之，气候潜力指数 C_r 、效能指数 C_e 值较大的地区，降水也相对比较丰富。对于局部地区，热量不足，低温寒冷也是影响农业生产的重要因素之一。关于这一点，通过气候适宜度指数也得到了反映，如东乡 $\geq 10^{\circ}\text{C}$ 的年活动积温只有 1584.2 $^{\circ}\text{C}$ ，而其气候潜力指数 C_r 只有 4.25，效能指数 C_e 只有 1.87。

通过以上分析可以知道，农业生态气候适宜度指数比较全面而综合地反映了气候条件对农业生产的适宜性，同时影响农业生产的气候限制性因素通过农业生态气候适宜度指数也得到了反映。所以，农业生态气候适宜度指数可以作为我们综合评价农业生态气候条件的重要依据。依据农业生态气候适宜度指数对甘肃黄土高原区农业生态气候评价的结果见表 9-2。

表 9-2 甘肃黄土高原区农业生态气候综合评价

评价指标		评语	评价地区(台站代表区)
C_r	C_e		
5.20-6.50	2.70-3.40	气候潜力较大，光水热匹配状况比较良好，气候天然利用程度较高。	漳县、甘谷、清水、天水县、天水台、静宁、泾川、崇信、庄浪、华亭、灵台、临洮、陇西、广河、和政、康乐、宁县、临夏。
5.00-6.00	2.00-2.70	气候潜力较大，光水热匹配状况适中，气候天然利用率适中。	张家川、秦安、武山、平凉、通渭、渭源、庆阳、镇原。
< 5.00	< 2.40	气候潜力适中，但光水热匹配状况较差，气候天然利用率较低。	东乡、环县、华池、合水、正宁、靖远、会宁、定西、景泰、永登、兰州、皋兰、榆中、永靖。

第十章 灰色系统方法

客观世界，既是物质的世界又是信息的世界。它既包含大量的已知信息，也包含大量的未知信息与非确知信息。未知的或非确知的信息称为黑色信息；已知信息称为白色信息。既含有已知信息又含有未知的、非确知的信息的系统，称为灰色系统。

在现实世界中，灰色系统是普遍存在的。灰色系统理论，是由我国著名学者邓聚龙先生于 80 年代初首创的一种系统科学理论。主要包括：灰色系统建模理论、灰色系统控制理论、灰色关联分析方法、灰色预测方法、灰色规划方法、灰色决策方法等。

地理系统，是一类典型的灰色系统。因此，灰色系统方法是地理学研究中的重要方法之一。

第一节 灰色关联分析方法

在地理系统中，许多因素之间的关系是灰色的，人们很难分清哪些因素是主导因素，哪些因素是非主导因素；哪些因素之间关系密切，哪些不密切。灰色关联分析，为我们解决这类问题提供了一种行之有效的方法。

一、灰色关联分析概述

我们知道，统计相关分析是对因素之间的相互关系进行定量分析的一种有效方法。但是，我们也注意到相关系数具这样的性质： $r_{xy}=r_{yx}$ ，即因素 y 对因素 x 的相关程度与因素 x 对因素 y 的相关程度相等。暂且不去追究因素之间的相关程度究竟有多大。单就相关系数的这种性质而言，也是与实际情况不太相符的。譬如，在国民经济问题研究中，我们能将农业对工业的关联程度与工业对农业的关联程度等同看待吗？其次，由于地理现象与问题的复杂性，以及人们认识水平的限制，许多因素之间的关系是灰色的，很难用相关系数比较精确地度量其相关程度的客观大小。为了克服统计相关分析的上述种种缺陷，灰色系统理论中的灰色关联分析给我们提供了一种分析因素之间相互关系的又一种方法。

灰色关联分析，从其思想方法上来看，属于几何处理的范畴，其实质是对反映各因素变化特性的数据序列所进行的几何比较。用于度量因素之间关联程度的关联度，就是通过对因素之间的关联曲线的比较而得到的。

设 $x_1, x_2, \dots, x_N$ 为 N 个因素，反映各因素变化特性的数据列分别为 $\{x_1(t)\}, \{x_2(t)\}, \dots, \{x_N(t)\}$ ， $t=1, 2, \dots, M$ 。因素 x_j 对 x_i 的关联系数定义为

$$\xi_{ij}(t) = \frac{\Delta_{\min} + k\Delta_{\max}}{\Delta_{ij}(t) + k\Delta_{\max}} \quad t=1, 2, 3, \dots, M \quad (1)$$

(5)式中， $\xi_{ij}(t)$ 为因素 x_j 对 x_i 在 t 时刻的关联系数； $\Delta_{ij}(t) = |x_i(t) - x_j(t)|$ ， $\Delta_{\max} = \max_j \max_j \Delta_{ij}(t)$ ， $\Delta_{\min} = \min_j \min_j \Delta_{ij}(t)$ ； k 为介于 $[0, 1]$ 区间上的灰数。不难看出， $\Delta_{ij}(t)$ 的最小值是 $\Delta_{\min}$ ，

当它取最小值时，关联系数 $\xi_{ij}(t)$ 取最大值 $\max_i \xi_{ij}(t) = 1$ ； $\Delta_{ij}(t)$ 的最大

值为 $\Delta_{\max}$ ，当它取最大值时，关联系数 $\xi_{ij}(t)$ 取最小值 $\min_i \xi_{ij}(t) = \frac{1}{1+k}$

· $\left(k + \frac{\Delta_{\min}}{\Delta_{\max}}\right)$ ，即 $\xi_{ij}(t)$ 是一个有界的离散函数。若取灰数 k 的白化值为1，则有

$$\frac{1}{2} \left(1 + \frac{\Delta_{\min}}{\Delta_{\max}}\right) \leq \xi_{ij}(t) \leq 1 \quad (2)$$

在实际计算时，取 $\Delta_{\min}=0$ ，这时有

$$0.5 \leq \xi_{ij}(t) \leq 1 \quad (3)$$

作出函数 $\xi_{ij} = \xi_{ij}(t)$ 随时间变化的曲线，它就被称之为关联曲线(见图10-1)。图中的水平线，说明任何时刻的关联系数为1，它代表 x_i 与 x_i 本身的关联曲线 $\xi_{ij} \equiv 1$ ，因为自己与自己总可以认为是密切关联的。

关联曲线 $\xi_{ij}(t)$ 与 $\xi_{ji}(t)$ 与坐标轴围成的面积分别记为 S_{ij} 与 S_{ji} ，则定义 x_j 对 x_i 的关联度为

$$\gamma_{ij} = \frac{S_{ij}}{S_{ji}} \quad (4)$$

显然 $S_{ii}=1 \times M=M$ ，所以(4)式可以进一步写成

$$\gamma_{ij} = S_{ij}/M \quad (5)$$

在实际计算中，常用近似公式

$$\gamma_{ij} \approx \frac{1}{M} \sum_{i=1}^M \xi_{ij}(t) \quad (6)$$

代替式(5)或式(6)。

从以上关联度的定义可以看出，它主要取决于各时刻的关联系数 $\xi_{ij}(t)$ 的值，而 $\xi_{ij}(t)$ 又取决于各时刻 x_i 与 x_j 观测值之差 $\Delta_{ij}(t)$ 。显然， x_i 与 x_j 的量纲不同，作图比例尺就会不同，因而关联曲线的空间相对位置也会不同，这就会影响关联度 (γ_{ij}) 的计算结果。为了消除量纲的影响，增强不同量纲的因素之间的可比性，就需要在进行关联度计算之前，首先对各要素的原始数据作初值变换或均值变换，然后利用变换后所得到的新数据作关联度计算。初值变换的计算公式为

$$x'_i(t) = x_i(t) / x_i(1) \quad i = 1, 2, \dots, N; t = 1, 2, \dots, M \quad (7)$$

均值变换的计算公式为

$$x'_i(t) = x_i(t) / \bar{x}_i \quad i = 1, 2, \dots, N; t = 1, 2, \dots, M \quad (8)$$

$$\text{在(8)式中, } \bar{x}_i = \frac{1}{M} \sum_{i=1}^M x_i(t)。$$

二、实例分析

表 10-1 给出了某地区 1986—1990 年期间农业总产值及与之相关的各产业产值数据。我们用灰色关联分析方法对该地区各产业之间的相互联系作一些初步分析。

将表 10-1 中的数据作均值化变换后，在公式(1)中，取灰数 k 的白化值为 0.5，经过计算得如下的关联度矩阵：

表 10-1 某地区 1986—1990 年农业产值数据

年度	序号	农业总产值 x_1 (万元)	在农业总产值中			
			种植业产值 x_2 (万元)	林业产值 x_3 (万元)	畜牧业产值 x_4 (万元)	副业产值 x_5 (万元)
1986	1	114230	74363	1725	29851	8291
1987	2	136236	60238	2497	62593	10908
1988	3	171242	63705	5034	85581	16922
1989	4	192819	62664	4424	97696	28035
1990	5	244187	74204	2779	135769	31435

$$R = \begin{pmatrix} \gamma_{11} & \gamma_{12} & \gamma_{13} & \gamma_{14} & \gamma_{15} \\ \gamma_{21} & \gamma_{22} & \gamma_{23} & \gamma_{24} & \gamma_{25} \\ \gamma_{31} & \gamma_{32} & \gamma_{33} & \gamma_{34} & \gamma_{35} \\ \gamma_{41} & \gamma_{42} & \gamma_{43} & \gamma_{44} & \gamma_{45} \\ \gamma_{51} & \gamma_{52} & \gamma_{53} & \gamma_{54} & \gamma_{55} \end{pmatrix} = \begin{pmatrix} 1 & 0.6441 & 0.5218 & 0.6702 & 0.5461 \\ 0.7533 & 1 & 0.5529 & 0.6177 & 0.5506 \\ 0.6069 & 0.5349 & 1 & 0.6459 & 0.5999 \\ 0.7446 & 0.5938 & 0.6459 & 1 & 0.7697 \\ 0.6408 & 0.5347 & 0.6038 & 0.7752 & 1 \end{pmatrix}$$

从上述关联度矩阵，可以得到如下几点结论：

- (1) $\gamma_{14} = 0.6702 = \max_{i=1} \gamma_{1i} > \gamma_{12} > \gamma_{15} > \gamma_{13}$ ，这表明，在该地区的农业生产中，畜牧业占有最大的优势，它对农业总产值增长的贡献最大，其次是种植业，再次是副业，最后是林业。
- (2) $\gamma_{24} = 0.6177 = \max \gamma_{2i}$ ，这表明，在林、牧、副各业中，与种植业联系最为紧密的是畜牧业。
- (3) $\gamma_{34} = 0.6459 = \max \gamma_{3i}$ ，表明，在种植业、畜牧业和副业中，与林业联系最紧密的是畜牧业。
- (4) $\gamma_{45} = 0.7697 = \max \gamma_{4i}$ ，表明，在种植业、林业和副业中，与畜牧业联系最紧密的是副业。

第二节 灰色预测方法

基于灰色建模理论的灰色预测法,按照其预测问题的特征,可分为五种基本类型,即数列预测、灾变预测、季节灾变预测、拓扑预测和系统综合预测。这五种类型的预测方法,都是区域开发研究中重要而且常用的预测方法。鉴于本书篇幅所限,本节拟和只对数列预测法灾变预测法及其在地理学研究中的应用问题作简单介绍。

一、数列预测

数列预测就是对某一指标的发展变化情况所作的预测,其预测的结果是该指标在未来各个时刻的具体数值。譬如,在地理学研究中,人口数量预测、耕地面积预测、粮食产量预测、工农业总产值预测,等等,都是数列预测。

数列预测的基础,是基于累加生成数列的 GM(1, 1) 模型。

设 $x^{(0)}(1), x^{(0)}(2), \dots, x^{(0)}(M)$ 是所要预测的某项指标的原始数据。

一般而言, $\{x^{(0)}(t)\}_{t=1}^M$ 是一个不平稳的随机数列. 对于这样一个随机数列, 如果* 趋势无规律可循(如图 10-2 所示), 则无法用回归预测法对其进行预测。

如果对 $\{x^{(0)}(t)\}_{t=1}^M$ 作一次累加生成处理, 即

$$x^{(1)}(1) = x^{(0)}(1)$$

$$x^{(1)}(2) = x^{(0)}(1) + x^{(0)}(2)$$

$$x^{(1)}(3) = x^{(0)}(1) + x^{(0)}(2) + x^{(0)}(3)$$

M

$$x^{(1)}(k) = \sum_{t=1}^k x^{(0)}(t)$$

M

$$x^{(1)}(M) = \sum_{t=1}^M x^{(0)}(t)$$

则得到一个新的数列 $\{x^{(1)}(t)\}_{t=1}^M = 1$ 。这个数列与原始数列 $\{x^{(0)}(t)\}_{t=1}^M$ 相比较, 其随机性程度大大弱化, 平稳程度大大增加(如图 10-3 所示)。对于这样的新数列, 其变化趋势可以近似地用如下微分方程描述:

$$\frac{dx^{(1)}}{dt} + ax^{(1)} = u \quad (1)$$

在(1)式中, a 和 u 可以通过如下最小二乘法拟合得到:

$$\begin{bmatrix} a \\ u \end{bmatrix} = (B^T B)^{-1} B^T Y_M \quad (2)$$

在(2)式中, Y_M 为列向量 $Y_M = [x^{(0)}(2), x^{(0)}(3), \dots, x^{(0)}(M)]^T$; B 为构造数据矩阵:

$$\begin{pmatrix} -\frac{1}{2}[x^{(1)}(1) + x^{(1)}(2)] & 1 \\ -\frac{1}{2}[x^{(1)}(2) + x^{(1)}(3)] & 1 \\ \vdots & \vdots \\ -\frac{1}{2}[x^{(1)}(M-1) + x^{(1)}(M)] & 1 \end{pmatrix}$$

微分方程(1)式所对应的时间响应函数为：

$$x^{(1)}(t+1) = \left[x^{(0)}(1) - \frac{u}{a} \right] e^{-at} + \frac{u}{a} \quad (3)$$

(3)式就是数列预测的基础公式，由(3)式对一次累加生成数列的预测值 $\bar{x}^{(1)}(t)$ 可以求得原始数的还原值：

$$\bar{x}^{(0)}(t) = \bar{x}^{(1)}(t) - \bar{x}^{(1)}(t-1) \quad (4)$$

在(4)式中， $t = 1, 2, \dots, M$ ，并规定 $\bar{x}^{(1)}(0) = 0$ 。原始数据的还原值与其观测值之间的残差值 $\epsilon^{(0)}(t)$ 和相对误差值 $q(t)$ 如下：

$$\begin{cases} \epsilon^{(0)}(t) = x^{(0)}(t) - \bar{x}^{(0)}(t) \\ q(t) = \frac{\epsilon^{(0)}(t)}{x^{(0)}(t)} \times 100\% \end{cases} \quad (5)$$

对于预测公式(3)，我们所关心的是它的预测精度。这一预测公式是否达到精度要求，可按下述方法进行精度检验。

首先计算：

$$\begin{aligned} \bar{x}^{(0)} &= \frac{1}{M} \sum_{t=1}^M x^{(0)}(t) \\ s_1^2 &= \frac{1}{M} \sum_{t=1}^M (x^{(0)}(t) - \bar{x}^{(0)})^2 \\ \overline{\epsilon^{(0)}} &= \frac{1}{M-1} \sum_{t=2}^M (\epsilon^{(0)}(t) - \overline{\epsilon^{(0)}})^2 \end{aligned}$$

其次计算：方差比 $c = s_2/s_1$

及小误差概率： $p\{ |\epsilon^{(0)}(t) - \overline{\epsilon^{(0)}}| < 0.6745s_1 \}$

一般地，预测公式(3)的精度检验可由表 10-2 给出。如果 p 和 c 都在允许范围之内，则可以计算预测值。否则，需要通过对残差序列 $\{\epsilon^{(0)}(t)\}_{t=2}^M$ 的分析对(3)式进行修正，灰色预测常用的修正方法有残差序列建模法和周斯分析法两种。

表 10-2 灰色预测精度检验等级标准

检验指标 等 级	P	C
好	> 0.95	< 0.35
合格	> 0.80	< 0.5
勉强	> 0.70	< 0.65
不合格	≤ 0.70	≥ 0.65

笔者曾对解放以来西北地区(包括陕、甘、宁、青、新等五省(区))耕地面积的变化情况作了分析,发现自 70 年代以来,本区耕地面积呈下降趋势。用 1971—1985 年的耕地面积数据建立耕地面积预测的 $G_M(1, 1)$ 模型,其预测公式如下:

$$x^{(1)}(t+1) = -4566239e^{-0.00409t} + 4584978 \quad (6)$$

公式(6)的精确度检验值为:

$$c = 0.2993 \quad (\text{好})$$

$$p = 0.9333 \quad (\text{合格})$$

根据上述公式计算,得到 1988—2000 年西北地区耕地面积的预测值(见表 10-3)。

二、灾变预测

一般地,如果表征系统行为特征的指标超出了某个阈值(临界值),则称发生了灾害。因

表 10-3 西北地区耕地面积预测值(单位:万亩)

年 份	1988	1989	1990	1991	1992	1993	1994
耕地面积	17399.00	17328.00	17257.00	17186.50	17116.26	17046.80	16976.75
年 份	1995	1996	1997	1998	1999	2000	
耕地面积	16907.25	16838.50	16769.50	16701.25	16639.00	16565.00	

此,所谓灾变是相对于所研究的问题的表征变量而言的。是否发生灾变要依据有关的表征变量的数值大小而定。譬如,旱灾和涝灾是相对于农作物生长过程中,作物需水与大气降水的差值大小而言的。如果以降水量作为旱涝灾害表征指标,则只有当降水量小于(或大于)某一阈值时,才认为发生了旱(或涝)灾。灾变预测就是指对灾变发生的年份的预测。

对于表征系统行为的指标数列:

$$\{x^{(0)}(1), x^{(0)}(2), \dots, x^{(0)}(N)\} \quad (7)$$

规定一个灾变阈值 ξ , $x^{(0)}(i)$ 中那些 $\leq \xi$ (或 $\geq \xi$) 的点被认为是具有异常值的点(灾变发生点),把它们按原来的编序挑选出来组成一个新的数据序列

$$\{x_{\xi}^{(0)}(i')\} = \{x^{(0)}(q) \mid x^{(0)}(q) \leq \xi\} \quad (8)$$

则式(8)称之为下限(或上限)灾变数列。

作灾变映射

$$p: \{i'\} \rightarrow \{q\} \quad (9)$$

则灾变预测就是按灾变日期序列

$$p = \{p(1'), p(2'), \dots, p(n')\} \quad (10)$$

建立 GM(1, 1) 预测模型所进行的灾变日期预测。

譬如，某地区连续 17 年的降水量数据如表 10-4 所示。若规定降水量 $\xi \leq 320\text{mm}$ 的年份为旱灾年份，试用灾变预测法预测下次旱灾发生的年份。

表 10-4 某地区年降水量(单位：mm)

序号(i)	1	2	3	4	5	6	7	
降水量 $x^{(0)}(i)$	390.6	412	320	559.2	380.8	542.4	553	3
序号(i)	10	11	12	13	14	15	16	
$x^{(0)}(i)$	300	632	540	406.2	313.8	576	587.6	3

(1) 首先作灾变映射，建立 GM(1, 1) 模型。

$$\begin{aligned} x_{\xi}^{(0)}(i') &= \{x_{\xi}^{(0)}(1'), x_{\xi}^{(0)}(2'), x_{\xi}^{(0)}(3'), x_{\xi}^{(0)}(4'), x_{\xi}^{(0)}(5')\} \\ &= \{320, 310, 300, 313.8, 318.5\} \\ &= \{x^{(0)}(3), x^{(0)}(8), x^{(0)}(10), x^{(0)}(14), x^{(0)}(17)\} \end{aligned}$$

作映射

$$p: \{i'\} \rightarrow \{q\}$$

对灾变日期序列

$$\begin{aligned} p &= \{p(1'), p(2'), p(3'), p(4'), p(5')\} \\ &= \{3, 8, 10, 14, 17\} \end{aligned}$$

建立 GM(1, 1) 模型

为了书写方便，不妨将 $p(i')$ 记为 $p(i)$ ($i=1, 2, 3, 4, 5$) 将 p 中的数据作一次累加处理：

$$\begin{aligned} p^{(1)}(1) &= p(1) = 3 \\ p^{(1)}(2) &= p(1) + p(2) = 11 \\ p^{(1)}(3) &= p(1) + p(2) + p(3) = 21 \\ p^{(1)}(4) &= p(1) + p(2) + p(3) + p(4) = 35 \\ p^{(1)}(5) &= p(1) + p(2) + p(3) + p(4) + p(5) = 52 \end{aligned}$$

■ $p^{(1)}(t)$ 可用下述微分方程拟合：

$$\frac{dp^{(1)}}{dt} + ap^{(1)} = u \quad (11)$$

而系统辨识参数为

$$\begin{bmatrix} a \\ u \end{bmatrix} = (B^T B)^{-1} B^T Y_M \quad (12)$$

(12) 式中：

$$B = \begin{pmatrix} -\frac{1}{2}[p^{(1)}(2) + p^{(1)}(1)] & 1 \\ -\frac{1}{2}[p^{(1)}(3) + p^{(1)}(2)] & 1 \\ -\frac{1}{2}[p^{(1)}(4) + p^{(1)}(3)] & 1 \\ -\frac{1}{2}[p^{(1)}(5) + p^{(1)}(4)] & 1 \end{pmatrix} = \begin{pmatrix} -7 & 1 \\ -16.5 & 1 \\ -28 & 1 \\ -43.5 & 1 \end{pmatrix}$$

$$Y_M = [p(2), p(3), p(4), p(5)]^T$$

$$= [8, 10, 14, 17]^T$$

经计算得: $\begin{bmatrix} a \\ u \end{bmatrix} = (B^T B)^{-1} B^T Y_M = \begin{bmatrix} -0.25361 \\ 6.288339 \end{bmatrix}$

因此(5)式就为:

$$\frac{DP^{(1)}}{DT} - 0.25361p^{(1)} = 6.288339 \quad (13)$$

(13)式的时间响应为:

$$p^{(1)}(i+1) = 27.677e^{-0.25361i} - 24.677 \quad (14)$$

(2)误差分析: 灾变日期数列的预测计算值与实际值的相对误差计算如下:

下:

计算值	实际值	相对误差
p(2)=7.999	p(2)=8	q(2)=0.125%
p(3)=10.286	p(3)=10	q(3)=-2.86%
p(4)=13.268	p(4)=14	q(4)=5.1%
p(5)=17.099	p(5)=17	q(5)=-0.582%

显然, 最大相对误差为 5.1%。所以上述模型(14)式可用于预测。

(3)预测:

将 i=5, 和 i=6 分别代入(14)式得:

$$p^{(1)}(5) = 51.662, p^{(1)}(6) = 73.342$$

$$\text{因此: } p(6) = p^{(1)}(6) - p^{(1)}(5) = 21.68$$

由于从 n=17 算起, 21.68 与 17 之差为 4.68, 所以从现在算起将在 4 年左右发生下一次旱灾。

第三节 灰色线性规划及其应用

我们在前面探讨了线性规划方法在物资调运、资源利用、合理下料、农业生产结构调整等规划方面的应用问题。然而，线性规划模型是静态模型，无动态性可言，而且还常常无解，不能适宜自然环境、技术条件和社会经济状况发展变化的要求。因此，我们用它来研究问题时，就难免存在这样或那样的困难。而其它的一些规划方法在解决实际问题时，都因存在算法实现上的困难。因而其应用的广泛性也受到了一定的限制。鉴于这种情形，有必要探讨一种理论上可行，算法实现比较简便的规划方法。为农业生产结构的优化服务。在这方面，灰色预测与含有灰元的线性规划相结合的方法——灰色线性规划方法为我们提供了一种可供尝试的途径。

一、灰色线性规划的数学描述

灰色线性规划的一般形式为：

$$\text{求 max(min)} Z = CX = \sum_{i=1}^n c_i x_i \quad (1)$$

$$\text{满足：} \begin{cases} \otimes(A)X \leq (=, \geq) b & (2) \\ X \geq 0 & (3) \end{cases}$$

其中， $X=[x_1, x_2, \dots, x_n]^T$ 为模型变量向量； $C=[c_1, c_2, \dots, c_n]$ 的目标函数的价值系数向量， $c_j (j=1, 2, \dots, n)$ 可以是灰数； $\otimes(A)$ 为约束条件的系数矩阵， A 为 $\otimes(A)$ 的白化矩阵，即

$$\otimes(A) = \begin{pmatrix} \otimes_{11} & \otimes_{12} & \cdots & \otimes_{1n} \\ \otimes_{21} & \otimes_{22} & \cdots & \otimes_{2n} \\ M & M & M & M \\ \otimes_{m1} & \otimes_{m2} & \cdots & \otimes_{mn} \end{pmatrix} \quad (4)$$

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ M & M & M & M \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix} \quad (5)$$

$\otimes_{ij} \in [\underline{a}_{ij}, \bar{a}_{ij}]$ ， $\bar{a}_{ij}$ 和 $\underline{a}_{ij}$ 分别为 $\otimes_{ij}$ 的下限值和上限值； $b=[b_1, b_2, \dots, b_m]^T$ 为约束向量。

在对灰色线性规划求解时，要首先将灰色参数白化，一般可以指定 A 中元素的值，即将 $\otimes(A)$ 中的元素白化。首先按最小白化值 $\underline{a}_{ij}$ 作一次计算，如果无解，则用最大白化值 $\bar{a}_{ij}$ 作一次计算，或者取区间 $[\underline{a}_{ij}, \bar{a}_{ij}]$ 中其它白化值进行计算，如果白化参数选取合适，则往往能够使规划从无解变为有解。

约束值 b 可以这样得到，若对第 i 个约束值 $b_i (i=1, 2, \dots, m)$ ，有一组白化值序列：

$$b_i^{(0)} = \{b_i^{(0)}(1), b_i^{(0)}(2), \dots, b_i^{(0)}(N)\} \quad (6)$$

若对 $b_i^{(0)}$ 作 AGO (一次累加) 后得 $b_i^{(1)}$ ，以 $b_i^{(1)}$ 数据建立 GM(1, 1) 预测模型，则

可以通过模型求出预测值 $\bar{b}_i^{(0)}(k)(k>N)$ 。作规划时，可按约束条件

$$\otimes(A)X \leq \begin{pmatrix} \bar{b}_1^{(0)}(k) \\ \bar{b}_2^{(0)}(k) \\ M \\ \bar{b}_m^{(0)}(k) \end{pmatrix} \quad (7)$$

求出 k 时刻的灰色线性规划值，在 $k>N$ 的条件下取不同值，可以得到未来发展时刻的各种规划值，这样不但可以知道现在条件下的最优结构，而且还可以知道最优结构的发展变化趋势。

二、应用实例

笔者曾研究了西北地区(包括陕西、甘肃、宁夏、新疆、青海五省区)种植业结构，运用灰色线性规划方法建立了本区种植结构的优化模型。现简述如下。

(一)模型变量

本模型由下述 14 个变量描述：

- x_1 ：水稻种植面积(万亩)；
- x_2 ：水地小麦播种面积(万亩)；
- x_3 ：旱地小麦播种面积(万亩)；
- x_4 ：水地玉米播种面积(万亩)；
- x_5 ：旱地玉米播种面积(万亩)；
- x_6 ：水地高粱播种面积(万亩)；
- x_7 ：旱地高粱播种面积(万亩)；
- x_8 ：其它杂粮播种面积(万亩)；
- x_9 ：水地复种面积(万亩)；
- x_{10} ：旱地复种面积(万亩)；
- x_{11} ：棉花播种面积(万亩)；
- x_{12} ：油料播种面积(万亩)；
- x_{13} ：大麻播种面积(万亩)；
- x_{14} ：其它经济作物播种面积(万亩)。

(二)约束条件

本模型共有下述 15 个约束条件：

- (1)总耕地面积限制；
- (2)水田面积限制；
- (3)水地面积限制；
- (4)旱地面积限制；
- (5)水地玉米播种面积约束；
- (6)旱地玉米播种面积约束；
- (7)水地高粱播种面积约束；
- (8)旱地高粱播种面积约束；
- (9)杂粮播种面积约束；
- (10)水地复种面积约束；
- (11)旱地复种面积约束；

- (12)棉花需求量约束；
- (13)油料播种面积约束；
- (14)大麻播种面积约束；
- (15)经济收入要求。

(三)目标函数：

使全区种植业总产值达到最大值(以 1980 年不变价格计算)。

(四)模型参数的白化

笔者运用灰色预测方法(详见上节内容)，建立了 GM(1, 1)预测模型群，对各类耕地面积以及主要农作物的单位面积产量作了预测，由此确定了模型中有关灰色参数的白化值(如表 10-5 所示)。

(五)白化模型及其计算结果

1990 年的规划模型为：

求： $\max Z = 156.326x_1 + 133.702x_2 + 66.536x_3 + 103.616x_4 + 51.808x_5 + 45x_6 + 22.5x_7 + 9.839x_8 + 19.678x_9 + 9.839x_{10} + 88.774x_{11} + 74.46x_{12} + 59.166x_{13} + 180x_{14}$ 满足：

$$\begin{aligned} \sum_{i=1}^8 x_i + \sum_{i=11}^{14} x_i &\leq 17257 & x_1 &= 500 \\ \sum_{i=1}^3 x_{2i} + 0.6x_{14} &\leq 8169.5 & \sum_{i=1}^3 x_{2i} + 1 + x_8 + \sum_{i=11}^{13} x_i + 0.4x_{14} &\leq 8587.5 \\ x_4 &\geq 816.95 & x_5 &\geq 858.75 \\ x_6 &\geq 100 & x_7 &\geq 150 \end{aligned}$$

表 10-5 模型参数白化值参数

年 度 参数项目	1990	1995	2000
耕地总面积(万亩)	17257	16907.25	16565
水地面积(万亩)	8169.5	8486.75	8749.375
旱地小麦亩产(千克/亩)	166.34	186.315	204
旱地玉米亩产(千克/亩)	259.04	295.175	327.68
旱地高粱亩产(千克/亩)	225	250	275
棉花亩产(千克/亩)	63.41	76.02	87.895
油料亩产(千克/亩)	74.46	83.81	94.445
大麻亩产(千克/亩)	31.305	36.315	45.36
水稻亩产(千克/亩)	390.815	431.135	466.36
其它杂粮亩产(千克/亩)	98.39	109.41	121.67

$$\begin{aligned} x_3 &\geq 120 & 0.06x_2 - x_9 &\geq 0 \\ 0.06x_3 - x_{10} &\geq 0 & x_{11} &\geq 95 \\ x_{12} &\geq 162 & x_{13} &\geq 192.5 \\ x_{14} &= 2000 & x_i &\geq 0 (i=1, 2, \dots, 14) \end{aligned}$$

1995 的规划模型为：

求：

$$\max Z = 172.45x_1 + 149.052x_2 + 74.526x_3 + 118.07x_4 + 59.035x_5 + 50x_6 + 25x_7 + 10.94x_8 + 21.88x_9 + 10.94x_{10} + 106.43x_{11} + 83.81x_{12} + 68.64x_{13} + 230x_{14} \text{ 满足：}$$

$$\sum_{i=1}^8 x_{1i} + \sum_{i=11}^{14} x_{1i} \leq 16907.25, \quad x_1 \leq 530,$$

$$\sum_{i=1}^3 x_{2i} + 0.6x_{14} \leq 8486.75, \quad \sum_{i=1}^3 x_{2i} + 1 + x_8 + \sum_{i=11}^{13} x_{1i} + 0.4x_{14} \leq 7890.5,$$

$$x_4 \geq 848.67, \quad x_5 \geq 789.05,$$

$$x_6 \geq 100, \quad x_7 \geq 140,$$

$$x_8 \geq 120, \quad 0.07x_2 - x_9 \geq 0,$$

$$0.07x_{13} - x_{10} \geq 0, \quad x_{11} \geq 87.5,$$

$$x_{12} \geq 158.7, \quad x_{13} \geq 183.17,$$

$$x_{14} = 2500, \quad x_i \geq 0 (i=1, 2, \dots, 14)。$$

2000 年的规划模型为：

$$\max Z = 186.54x_1 + 163.2x_2 + 81.6x_3 + 131.072x_4 + 65.54x_5 + 55x_6 + 27.5x_7 + 12.167x_8 + 24.334x_9 + 12.167x_{10} + 123.05x_{11} + 84.445x_{12} + 85.73x_{13} + 280x_{14} \text{ 满足：}$$

$$\sum_{i=1}^8 x_{1i} + \sum_{i=11}^{14} x_{1i} \leq 16565, \quad x_1 \leq 550,$$

$$\sum_{i=1}^8 x_{2i} + 0.6x_{14} \leq 8749.375, \quad \sum_{i=1}^3 x_{2i} + 1 + x_3 + \sum_{i=11}^{13} x_{1i} + 0.4x_{14} \leq 7265.625,$$

$$x_4 \geq 874.94, \quad x_5 \geq 726.56,$$

$$x_6 \geq 100, \quad x_7 \geq 140,$$

$$x_8 \geq 120, \quad 0.08x_2 - x_9 \geq 0,$$

$$0.08x_3 - x_{10} \geq 0, \quad x_{11} \geq 80,$$

$$x_{12} \geq 150, \quad x_{13} \geq 161.91,$$

$$x_{14} = 3000, \quad x_i \geq 0 (i=1, 2, \dots, 14)。$$

上述各时刻白化模型的计算结果见表 10-6。

表 10-6 西北地区种植业结构优化结果优化值

年度	1990	1995	2000
x_1 (万亩)	500.00	530.00	550.00
x_2 (万亩)	6052.55	6038.08	5974.436
x_3 (万亩)	6209.50	5412.08	4687.155
x_4 (万亩)	816.95	848.67	874.94
x_5 (万亩)	858.75	726.56	
x_6 (万亩)	100.00	100.00	100.00
x_7 (万亩)	150.00	140.00	140.00
x_8 (万亩)	120.00	120.00	120.00
x_9 (万亩)	363.15	422.67	477.95
x_{10} (万亩)	372.56	378.85	374.97
x_{11} (万亩)	95.00	87.00	80.00
x_{12} (万亩)	162.00	158.70	150.00
x_{13} (万亩)	192.50	183.17	161.91
x_{14} (万亩)	2000.00	2500.00	3000.00
maxZ(万元)	1837 619.00	2174 903.00	2525787.00
经济作物产值(万元)	391885.51	610186.06	876 9991.29
粮食总产量(万公斤)	4086186.17	4451510.765	4731757.37

第四节 灰色局势决策方法

灰色局势决策，是灰色系统理论中一种重要的决策方法之一，它是将事件、对策、效果、目标等决策四要素综合考虑的一种决策分析方法。这种方法的最大特点是它适用于处理数据中含有灰元，即信息不完备的决策问题。在区域开发活动中，许多问题的解决是在信息不完备的情况下作出决策的。因此，灰色局势决策是地理学研究中常用的决策分析方法之一。

一、灰色局势决策的数学模型与基本步骤

(一)灰色局势决策的数学模型

决策，一般都包括如下四个基本要素：

- (1)事件，即需要处理的事物；
- (2)对策，即处理某一事物的措施；
- (3)效果，即用某个对策对付某个事件的效果；
- (4)目标，即用来评价效果的准则。

所谓决策就是指，对于某个(或某些)事件，考虑许多对策去对付，不同对策效果不同，然后用某种(或某几种)目标去衡量，从这些对策中选择一个(或一批)效果最佳者。

灰色局势决策，是一种将事件、对策、效果、目标等决策四要素综合考虑的一种决策分析方法。灰色局势决策的数学模型，实质上是运用有关的数学语言对决策四要素之间的相互关系所作的一种综合性描述。这种描述主要包括如下几个方面的基本内容。

1.决策元、决策向量与决策矩阵

(1)决策元。在灰色局势决策中，事件 a_i 和对策 b_j 的二元组合 $s_{ij}=(a_i, b_j)$ 称为局势，它表示用第 j 个对策(b_j)去对付第 i 个事件(a_i)的局势。

若局势 s_{ij} 的效果测度为 r_{ij} ，则称

$$\frac{r_{ij}}{s_{ij}} = \frac{r_{ij}}{(a_i, b_j)}$$

为决策元。它表示用第 j 个对策(b_j)去对付第 i 个事件(a_i)这一局势的效果为 r_{ij} 。

(2)决策向量。若某一类决策问题有 n 个事件 $a_1, a_2, \dots, a_n$ 和 m 个对策 $b_1, b_2, \dots, b_m$ ，且对于每一个事件 $a_i (i=1, 2, \dots, n)$ 都可以用 $b_1, b_2, \dots, b_m$ 等 m 个对策去对付。那么，对于每一个事件 $a_i (i=1, 2, \dots, n)$ ，就存在有 m 个局势：

$$(a_i, b_1), (a_i, b_2) \cdots, (a_i, b_m)$$

这些局势相应的决策元可排成一行，便构成了一个决策行向量：

$$\delta_i = \left[\frac{r_{i1}}{s_{i1}}, \frac{r_{i2}}{s_{i2}}, \dots, \frac{r_{im}}{s_{im}} \right] \quad (1)$$

(1)式中， r_{ij} 为局势 $s_{ij}=(a_i, b_j)$ 的效果测度。

同样，对于每一个对策 $b_j (j=1, 2, \dots, m)$ ，可以用事件 $a_1, a_2, \dots, a_n$ 去匹配，其相应的决策元可排成一列，便构成了一个决策列向量：

$$\theta_j = \begin{pmatrix} \frac{r_{1j}}{s_{1j}} \\ \frac{r_{2j}}{s_{2j}} \\ M \\ \frac{r_{nj}}{s_{nj}} \end{pmatrix} \quad (2)$$

(3)决策矩阵。将每一个决策行向量 $\delta_i (i=1, 2, \dots, n)$ 或每一个决策列向量 $\theta_j (j=1, 2, \dots, m)$ 依次排列起来，便构成了一个 $n \times m$ 的局势决策矩阵：

$$M = \begin{pmatrix} \frac{r_{11}}{s_{11}} & \frac{r_{12}}{s_{12}} & \dots & \frac{r_{1m}}{s_{1m}} \\ \frac{r_{21}}{s_{21}} & \frac{r_{22}}{s_{22}} & \dots & \frac{r_{2m}}{s_{2m}} \\ M & M & & M \\ \frac{r_{n1}}{s_{n1}} & \frac{r_{n2}}{s_{n2}} & \dots & \frac{r_{nm}}{s_{nm}} \end{pmatrix} \quad (3)$$

2.效果测度 效果测度就是对于局势所产生的实际效果,在不同目标之间进行比较的量度。

对于时间序列来说，就是比较两个序列在同一时刻的关联系数，其计算公式为：

$$r_{ij}(t) = \frac{\Delta_{\min} + K\Delta_{\max}}{\Delta_{ij}(t) + K\Delta_{\max}} \quad (4)$$

(4)式中， $\Delta_{ij}(t)$ 为两序列在 t 时刻的绝对差； $\Delta_{\min}$ 和 $\Delta_{\max}$ 分别是两序列绝对差的最小值和最大值； K 是在 $[0, 1]$ 区间上取值的灰数。

作为时间序列的效果测度，其被比较的母线，一般应为规划的目标效益曲线。

对于单点效果测度，可分为以下几种情形：

(1)上限效果测度，其计算公式为：

$$r_{ij} = \frac{u_{ij}}{u_{\max}} \quad (5)$$

(5)式中， u_{ij} 为局势 s_{ij} 的实际效果； $u_{\max}$ 为所有局势 s_{ij} 实际效果的最大值。由于 $u_{ij} \leq u_{\max}$ ，所以效果测度 $r_{ij} \leq 1$ 。

(2)下限效果测度，其计算公式为：

$$r_{ij} = \frac{u_{\min}}{u_{ij}} \quad (6)$$

(6)式中， u_{ij} 的意义同(5)式， $u_{\min}$ 为所有 u_{ij} 中的最小者。由于 $u_{ij} \geq u_{\min}$ ，显然 $r_{ij} \leq 1$ 。

(3)适中效果测度，其计算公式为：

$$r_{ij} = \frac{\min\{u_{ij}, u_0\}}{\max\{u_{ij}, u_0\}} \quad (7)$$

(7)式中, u_{ij} 的意义同(5)式, u_0 是一个指定的适中值。由(7)式容易知道, $r_{ij} \leq 1$ 。

如果 u_0 是以几何中心为参考点的数值, 则适中效果测度的计算公式为:

$$r_{ij} = \frac{u_0}{u_0 + |u_{ij} - u_0|} \quad (8)$$

在实际应用中, 究竟采用哪种效果测度, 应依据目标的性质而定。如产值、效益之类应该是越大越好, 可采用上限效果测度; 如投资、灾害之类应该是越小越好, 可采用下限效果测度; 而对于降水量、施肥量等应以适量为宜, 可采用适中效果测度。

此外, 对于局势 s_{ij} 有效益时间序列, 则需求稳态效果测度。即对时间序列 $\{u_{ij}(t)\}$ 建立 GM(1, 1) 模型, 解得灰色参数 $a = [a, u]^T$ 。当以 u 为输入时, 则稳态增益为:

$$\lim_{t \rightarrow \infty} \frac{u_{ij}(t)}{u} = \frac{1}{a} \quad (9)$$

因此, $\frac{1}{a}$ 被称为稳态效果测度, 记为:

$$r_{ij} = \frac{1}{a} \quad (10)$$

3. 多目标综合决策矩阵 当有 l 个决策目标时, 记局势 s_{ij} 在第 p 个目标下的效果测度为 $r_{ij}^{(p)}$, 则其相应的决策元为 $\frac{r_{ij}^{(p)}}{s_{ij}}$ 。如前所述, 同样可以得到相

应的行决策向量 $\delta_i^{(p)}$ ($i = 1, 2, \dots, n$) 和列决策向量 $\theta_j^{(p)}$, 以及决策矩阵

$$M^{(p)} = \begin{pmatrix} \frac{r_{11}^{(p)}}{s_{11}} & \frac{r_{12}^{(p)}}{s_{12}} & \dots & \frac{r_{1m}^{(p)}}{s_{1m}} \\ \frac{r_{21}^{(p)}}{s_{21}} & \frac{r_{22}^{(p)}}{s_{22}} & \dots & \frac{r_{2m}^{(p)}}{s_{2m}} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{r_{n1}^{(p)}}{s_{n1}} & \frac{r_{n2}^{(p)}}{s_{n2}} & \dots & \frac{r_{nm}^{(p)}}{s_{nm}} \end{pmatrix} \quad (11)$$

如果第 p 个决策目标的权重值为 a_p ($p=1, 2, \dots, l$), 则对于局势 s_{ij} , 可以得到如下的综合效果测度:

$$r_{ij}^{(\Sigma)} = \sum_{p=1}^l a_p T_{ij}^{(p)} \quad (12)$$

当各目标的权重值相等, 即 $a_1 = a_2 = \dots = a_l = \frac{1}{n}$ 时, (12)式就变为:

$$r_{ij}^{(\Sigma)} = \frac{1}{l} \sum_{p=1}^l r_{ij}^{(p)} \quad (12')$$

这样, 我们就得到如下的多目标综合决策矩阵:

$$M^{(\Sigma)} = \left\{ \begin{array}{cccc} \frac{r_{11}^{(\Sigma)}}{s_{11}} & \frac{r_{12}^{(\Sigma)}}{s_{12}} & \dots & \frac{r_{1m}^{(\Sigma)}}{s_{1m}} \\ \frac{r_{21}^{(\Sigma)}}{s_{21}} & \frac{r_{22}^{(\Sigma)}}{s_{22}} & \dots & \frac{r_{2m}^{(\Sigma)}}{s_{2m}} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{r_{n1}^{(\Sigma)}}{s_{n1}} & \frac{r_{n2}^{(\Sigma)}}{s_{n2}} & \dots & \frac{r_{nm}^{(\Sigma)}}{s_{nm}} \end{array} \right\} \quad (13)$$

4. 决策原则决策就是选择效果最佳的局势。这种选择可以有两种方式：

- (1) 由事件选择最好的对策，即行决策；
- (2) 由对策匹配最适宜的事件，即列决策。

行决策的原则是，对于综合决策矩阵 $M^{(\Sigma)}$ ，在行决策向量 $\delta_i^{(\Sigma)}$ 中，选取效果测度最大的决策元，即：

$$r_{ij^*} = \max \gamma_{ij}^{(\Sigma)} = \max \{\gamma_{i1}^{(\Sigma)}, \gamma_{i2}^{(\Sigma)}, \dots, \gamma_{im}^{(\Sigma)}\} \quad (14)$$

则 $\frac{r_{ij^*}^{(\Sigma)}}{s_{ij^*}}$ 称为行决策元， s_{ij^*} 为最优局势，即表示 b_{j^*} 是对付事件 a_i 的最优

对策。

列决策的原则是，对于综合决策矩阵 $M^{(\Sigma)}$ ，在列决策向量 $\theta_j^{(\Sigma)}$ 中，选取效果测度最大的决策元，即：

$$r_{i^*j}^{(\Sigma)} = \max_i r_{ij}^{(\Sigma)} = \max \{r_{1j}^{(\Sigma)}, r_{2j}^{(\Sigma)}, \dots, r_{nj}^{(\Sigma)}\} \quad (15)$$

则 $\frac{r_{i^*j}^{(\Sigma)}}{s_{i^*j}}$ 称为列决策元， s_{i^*j} 为最优局势，即表示 a_{i^*} 是对策 b_j 最适宜的

事件。

(二) 灰色局势决策的基本步骤

灰色局势决策方法求解问题的过程，一般可以按下述步骤进行：

- (1) 给出事件与对策；
- (2) 构造局势；
- (3) 确定目标；
- (4) 给出不同目标的白化值；
- (5) 计算不同目标的局势效果测度，写出决策矩阵；
- (6) 计算多目标的局势综合效果测度，写出多目标综合决策矩阵；
- (7) 按照行决策或列决策原则，选择最佳局势。

二、应用实例

假设某县有三个经济区，各区农、林、牧各业的单位土地面积产值(即地均产值)和单位产值所消耗的物质水平(即物耗水平)的样本值如表 10-7 所示，试用灰色局势决策方法确定各经济区农业发展的方向。

表 10-7 某农业县地均产值和物耗水平

区号	地均产值(百元/km ²)			物耗水平(%)		
	农业	林业	牧业	农业	林业	牧业
I	3902	549	776	0.37	0.33	0.52
II	7790	490	1303	0.31	0.29	0.45
III	8118	460	1092	0.33	0.35	0.51

(一)给出事件与对策

1.事件该决策问题共含三个事件：

(1) a_1 ：I 区；

(2) a_2 ：II 区；

(3) a_3 ：III区。

2.对策 该决策问题共有三个对策：

(1) b_1 ：农业；

(2) b_2 ：林业；

(3) b_3 ：牧业。

(二)构造局势

各个事件与对策相互匹配可得如下局势：

(1) $s_{11}=(a_1, b_1)=(I \text{ 区, 农业})$ ；

(2) $s_{12}=(a_1, b_2)=(I \text{ 区, 林业})$ ；

(3) $s_{13}=(a_1, b_3)=(I \text{ 区, 牧业})$ ；

(4) $s_{21}=(a_2, b_1)=(II \text{ 区, 农业})$ ；

(5) $s_{22}=(a_2, b_2)=(II \text{ 区, 林业})$ ；

(6) $s_{23}=(a_2, b_3)=(II \text{ 区, 牧业})$ ；

(7) $s_{31}=(a_3, b_1)=(III \text{ 区, 农业})$ ；

(8) $s_{32}=(a_3, b_2)=(III \text{ 区, 林业})$ ；

(9) $s_{33}=(a_3, b_3)=(III \text{ 区, 牧业})$ 。

(三)确定目标

本决策问题主要考虑如下两个目标：

(1)单位土地面积产值尽可能大(目标 1)；

(2)单位产值的物质消耗水平尽可能低(目标 2)。

(四)给出不同目标的白化值

由表10-7可知，目标1的白化值分别为： $u_{11}^{(1)} = 3902$ ， $u_{12}^{(1)} = 549$ ，

$u_{13}^{(1)} = 776$ ， $u_{21}^{(1)} = 7790$ ， $u_{22}^{(1)} = 499$ ， $u_{23}^{(1)} = 1303$ ， $u_{31}^{(1)} = 8118$ ， $u_{32}^{(1)} = 460$ ，

$u_{33}^{(1)} = 1092$ ；目标2的白化值分别为： $u_{11}^{(2)} = 0.37$ ， $u_{12}^{(2)} = 0.33$ ， $u_{13}^{(2)} = 0.52$ ，

$u_{21}^{(2)} = 0.31$ ， $u_{22}^{(2)} = 0.29$ ， $u_{23}^{(2)} = 0.45$ ， $u_{31}^{(2)} = 0.33$ ， $u_{32}^{(2)} = 0.35$ ， $u_{33}^{(2)} = 0.51$ 。

(五)计算各目标的效果测度，写出决策矩阵

(1)对于目标 1，按上限效果测度计算，可得如下决策矩阵：

$$M^{(1)} = \begin{pmatrix} \frac{r_{11}^{(1)}}{s_{11}} & \frac{r_{12}^{(1)}}{s_{12}} & \frac{r_{13}^{(1)}}{s_{13}} \\ \frac{r_{21}^{(1)}}{s_{21}} & \frac{r_{22}^{(1)}}{s_{22}} & \frac{r_{23}^{(1)}}{s_{23}} \\ \frac{r_{31}^{(1)}}{s_{31}} & \frac{r_{32}^{(1)}}{s_{32}} & \frac{r_{33}^{(1)}}{s_{33}} \end{pmatrix} = \begin{bmatrix} 0.4810 & 0.0676 & 0.0956 \\ 0.9596 & 0.0553 & 0.1605 \\ 1 & 0.0567 & 0.1345 \end{bmatrix}$$

(2)对于目标 2，按下限效果测度公式计算，可得如下决策矩阵：

$$M = \begin{pmatrix} \frac{r_{11}^{(2)}}{s_{11}} & \frac{r_{12}^{(2)}}{s_{12}} & \frac{r_{13}^{(2)}}{s_{13}} \\ \frac{r_{21}^{(2)}}{s_{21}} & \frac{r_{22}^{(2)}}{s_{22}} & \frac{r_{23}^{(2)}}{s_{23}} \\ \frac{r_{31}^{(2)}}{s_{31}} & \frac{r_{32}^{(2)}}{s_{32}} & \frac{r_{33}^{(2)}}{s_{33}} \end{pmatrix} = \begin{bmatrix} 0.7838 & 0.8788 & 0.5577 \\ 0.9355 & 1 & 0.6444 \\ 0.8788 & 0.8286 & 0.5686 \end{bmatrix}$$

(六)计算多目标的局势综合效果测度，写出综合决策矩阵

取目标 1 和目标 2 的权重值为 $\alpha_1 = \alpha_2 = 0.5$ ，将其代入公式(12)计算，可得多目标的局势综合效果测度 $r_{ij}^{(\Sigma)}$ ($i, j = 1, 2, 3$)，其相应的综合决策矩阵为：

$$M^{(\Sigma)} = \begin{pmatrix} \frac{r_{11}^{(\Sigma)}}{s_{11}} & \frac{r_{12}^{(\Sigma)}}{s_{12}} & \frac{r_{13}^{(\Sigma)}}{s_{13}} \\ \frac{r_{21}^{(\Sigma)}}{s_{21}} & \frac{r_{22}^{(\Sigma)}}{s_{22}} & \frac{r_{23}^{(\Sigma)}}{s_{23}} \\ \frac{r_{31}^{(\Sigma)}}{s_{31}} & \frac{r_{32}^{(\Sigma)}}{s_{32}} & \frac{r_{33}^{(\Sigma)}}{s_{33}} \end{pmatrix} = \begin{bmatrix} 0.6324 & 0.4732 & 0.32665 \\ 0.94755 & 0.52765 & 0.40245 \\ 0.9394 & 0.44265 & 0.35155 \end{bmatrix}$$

(七)进行决策

(1)按行决策，可得最优局势： s_{11} ， s_{21} ， s_{31} 。这表明，对于每一个经济区来说，都具有发展农业的优势。

(2)按列决策，可得最优局势： s_{21} ， s_{22} ， s_{23} 。表明，最适合发展农业、林业和牧业的都是 II 区。

综合上述结果可知，I 区和 III 区的发展方向应该是农业，II 区的发展方向应该是农、林、牧各业综合发展。

第五节 灰色去余控制理论及其应用

灰色去余控制理论，是我国著名的控制论专家邓聚龙先生针对灰色系统的控制问题提出的一种优化控制方法。目前，这一方法已被广泛地应用于各类系统的控制分析之中。下面介绍灰色去余控制理论及其在地理系统调控中的应用。

一、灰色去余控制理论概述

(一)系统的结构模型

由第八章第二节的介绍，我们知道，一个系统的动态过程，可以用一定的传递函数表示。如果 U 与 X 分别表示系统的输入与输出的拉普拉斯变换，系统的传递函数记为 W ，则有：

$$\frac{X}{U} = W \quad (1)$$

有时更习惯于用

$$W^{-1}X = U \quad (2)$$

表示系统的输入-输出关系。

若系统中有反馈 Z ，将 X 反馈回输入端，则其动态关系为：

$$\frac{X}{U} = \frac{W}{1 + WZ} \quad (3)$$

(3)式可以被写成：

$$\frac{1 + WZ}{W} X = U$$

或者

$$W^{-1}X = U - ZX \quad (4)$$

上述关系式称为反馈系统的结构范式，或者称为系统的结构模型。

对于上述的系统结构模型，特作以下几点说明：

(1)等号(=)是系统输入与输出的比较环节。等号左端的 X 为系统的输出，右端的 U 为系统的输入。

(2)从系统的输入端 U 经过 W 到输出端 X 的通道称为主通道， W 称为主通道的传递函数；将 X 经过 Z 反馈到输入端的通道称为反馈通道， Z 称为反馈通道的传递函数， ZX 称为反馈项。

(3)结构模型等号左端 X 的系数 W^{-1} 是主通道传递函数的倒数。

(4)反馈项 ZX 前面的符号代表反馈极性，“+”表示正反馈，“-”表示负反馈。

(5)结构模型中的 X 与 U 可以是单一变量，也可以是多个变量构成的向量；当 X 与 U 为向量时， W 与 Z 为矩阵。

(6)称 $Z=0$ 的结构模型

$$W^{-1}X = U$$

为开环系统，称

$$W^{-1}X = U - ZX$$

为闭环系统。但是，这种叫法是相对的，事实上，带有反馈项的结构模型同样可以化为无反馈的形式，譬如令：

$$W_{\Sigma}^{-1} = \frac{1 + WZ}{W}$$

则有：

$$W_{\Sigma}^{-1}X = U$$

上式在形式上是无反馈的开环系统，而实际上却是有反馈的闭环系统。

(7)系统的结构模型不是唯一的。

(二)灰色去余控制的基本思想

灰色去余控制理论认为，系统的动态品质的好坏，主要反映在闭环系统传递函数 G 的结构与参数上，因此，要改善系统的动态品质，实现系统的优化控制，就需要改变闭环系统的传递函数。

设原系统的结构模型为：

$$G^{-1}X=U \quad (5)$$

记预期的(优化的)系统为

$$G_*^{-1}X = U \quad (6)$$

记两系统的逆传递函数矩阵 G^{-1} 与 G_*^{-1} 之差为 Q ，即：

$$Q = G_*^{-1} - G^{-1} \quad (7)$$

则原系统可以改写成：

$$G_*^{-1}X = U + QX \quad (8)$$

(8)式代表了主通道传递函数为 G_* ，反馈项为 QX 的系统。这说明原来的系统 $G^{-1}X = U$ ，可以看作是由预期系统 $G_*^{-1}X = U$ 与反馈项 QX 所组成。从预期系统的动态品质来看， QX 项是多余的，故 QX 称为系统的多余项，它以“虚内反馈”的形式作用在系统上，恶化了系统的动态品质。

当系统的多余项被分离出来后，则只要加入传递函数相等，反馈的输入与输出相同，而极性相反的控制项抵消多余项，系统就可以得到令人满意的品质。这种用外反馈抵消多余项的控制方法称为系统的去余控制。其去余过程，可以用结构关系：

多余项 去余项

$$\begin{array}{c} G_*^{-1}X = U + QX - QX \\ \downarrow \\ G^{-1}X = U \\ \text{预期系统} \end{array}$$

及图 10-4 所示的框图来表示：

(三)灰色去余控制的途径、方式和准则

制定控制决策，以改善系统的动态品质，实现系统的优化控制，可以有不同的途径、不同的方式和不同的准则。

1.控制决策的实施途径

(1)参数调整决策。这是指对系统结构参数的大小、符号进行改变，以改善系统的动态品质。

(2)结构改变决策。通过从外部引入控制结构。譬如，加去余控制部分以改善系统动态。

2.控制决策的方式

(1)通过系统的结构模型分解出多余项，然后采取去余控制措施，这种去余控制常称为频域去余，或者结构去余。

(2)通过系统的状态模型实现去余，称为状态去余，或者时域去余。

3.控制策略的准则

(1)输入无偏准则。记系统的输入为 U ，反馈为 z ，则动态无偏准则为：

$$J=U-z=\min \quad (9)$$

这里 z 可以是向量，譬如：

$$z=ZX \quad (10)$$

在(10)式中，

$$Z = \begin{pmatrix} z_{11} & z_{12} & \cdots & z_{1n} \\ z_{21} & z_{22} & \cdots & z_{2n} \\ M & M & & M \\ z_{n1} & z_{n2} & \cdots & z_{nn} \end{pmatrix}$$

(2)动态无偏准则。记 G_u 及 G_* 分别表示接受控制后以及预期的系统传递函数，则动态无偏准则可以表示为：

$$J = G_u^{-1} - G_*^{-1} = \min(s) \quad (11)$$

(11)式中， $\min(s)$ 表示某一个系数尽可能小的 s 多项式的分式。

(3)状态无偏准则。记 A_u 及 A_* 分别为接受了控制及预期的状态矩阵，则状态无偏准则为：

$$J=A_u-A_*= \min \quad (12)$$

二、人工草地生态经济系统的优化调控模型

黄土高原曾是中华民族的摇篮，但长期以来，由于不合理的人类活动，滥垦，滥殖，破坏了森林和草原，从而导致了干旱、风沙灾害、水土流失等。日趋严重的环境问题，这些问题都直接或间接地影响着人类的生存。解放以后，国家对黄土高原的治理工作非常重视。在科技人员与广大人民群众的努力探索下，找到了一系列切实可行的治理措施，如植树造林、种草养畜、小流域综合治理等。这些措施都对黄土高原农业生态系统的调控起到了积极的作用。特别是种草养畜这一措施的推广实施，不但改善了生态环境，而且取得了良好的经济效益。譬如，地处黄土丘陵区的山西省柳林县李家垣村，自1979年开始种草养兔，到1983年已形成了产业优势。阎文瑛、王学萌先生曾根据该村五年种草养兔过程中所积累的有关数据(表10-8)运用灰色建模方法建立了系统的动态模型，并运用灰色去余控制理论提出了该人工草地生态系统的优化调控策略，现在将他们的工作介绍如下。

(一)系统的动态模型

显然，售兔收入是系统的输出变量，而投资则是系统的输入变量或控制变量。但售兔收入受诸多因素影响，除市场因素外，在生产过程中首先取决于可供繁殖用的种兔及饲草来源的制约，而这两个因素又取决于投资。从投资——生产——产品输出(即养兔收入)，有如下几个环节(见图10-5)：

表 10-8 种草养兔生产数据

年 度	1979	1980	1981	1982	1983
序 号(i)	1	2	3	4	5
售兔收入 $x_1^{(0)}(i)$ (元)	4834	7625	10500	11316	17818
种兔只数 $x_2^{(0)}(i)$ (只)	83	131	180	195	306
种草面积 $x_3^{(0)}(i)$ (亩)	146	212	233	259	404
种兔投资 $x_4^{(0)}(i)$ (元)	251	396	545	590	926
种草投资 $x_5^{(0)}(i)$ (元)	1218	1762	1941	2151	3360

在图 10-5 中, x_1 : 售兔收入; x_2 : 繁殖用种兔; x_3 : 种草面积; x_4 : 种兔投资; x_5 : 种草投资; u : 初投资。

为此, 我们需建立整个系统的模型, 以便分析各环节间的动态关系以及系统的优化调控决策。

1. 种兔与其投资环节以数据 $\{x_2^{(0)}(i)\}$ 与 $\{x_4^{(0)}(i)\}$ 建立 GM(1, 2) 模型, 构造数据矩阵为:

$$X(2, 4) = \begin{pmatrix} -\frac{1}{2} \left[\sum_{i=1}^2 x_2^{(0)}(i) + \sum_{i=1}^1 x_2^{(0)}(i) \right], \sum_{i=1}^2 x_4^{(0)}(i) \\ -\frac{1}{2} \left[\sum_{i=1}^3 x_2^{(0)}(i) + \sum_{i=1}^2 x_2^{(0)}(i) \right], \sum_{i=1}^3 x_4^{(0)}(i) \\ -\frac{1}{2} \left[\sum_{i=1}^4 x_2^{(0)}(i) + \sum_{i=1}^3 x_2^{(0)}(i) \right], \sum_{i=1}^4 x_4^{(0)}(i) \\ -\frac{1}{2} \left[\sum_{i=1}^5 x_2^{(0)}(i) + \sum_{i=1}^4 x_2^{(0)}(i) \right], \sum_{i=1}^5 x_4^{(0)}(i) \end{pmatrix} = \begin{bmatrix} -148.5 & 647 \\ -304 & 1192 \\ -491.5 & 1782 \\ -742 & 2708 \end{bmatrix}$$

$$Y_5 = [x_2^{(0)}(2), x_2^{(0)}(3), x_2^{(0)}(4), x_2^{(0)}(5)]^T = [131, 180, 195, 306]$$

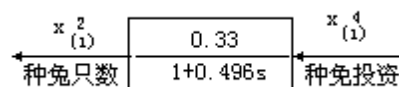
按最小的乘法解系统辨识系数 a, b 为:

$$\begin{bmatrix} a \\ b \end{bmatrix} = [X(2,4)]^{-1} \cdot X(2,4)^T \cdot Y_5 = \begin{bmatrix} 2.0078 \\ 0.6632 \end{bmatrix}$$

得该环节的微分方程为:

$$\frac{dx_2^{(1)}}{dt} + 2.0078x_2^{(1)} = 0.6632x_4^{(1)} \quad (13)$$

记 s 为 Laplace 算子, 则从投资 $x_4^{(1)}$ 到种兔 $x_2^{(1)}$ 的动态环节框图及传递函数为:



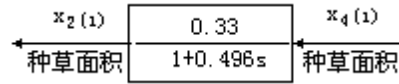
2. 种草及其投资环节 建立种草面积数据列 $\{x_3^{(0)}(i)\}$ 与其投资数据 $\{x_5^{(0)}(i)\}$ 的 GM(1, 2) 模型, 系统辨识系数为:

$$\begin{bmatrix} a \\ b \end{bmatrix} = \begin{bmatrix} 1.9917 \\ 0.2395 \end{bmatrix}$$

该环节的微分方程为：

$$\frac{dx_3^{(1)}}{dt} + 1.9917x_3^{(1)} = 0.2395x_5^{(1)} \quad (14)$$

该动态环节框图及传递函数为：



3. 养兔收入与种兔、种草环节 根据数据列 $\{x_1^{(0)}(i)\}$, $\{x_2^{(0)}(i)\}$ 及 $\{x_3^{(0)}(i)\}$ 建立 GM(1, 3) 型, 构造数据矩阵为：

$$X(1, 3) = \begin{pmatrix} -\frac{1}{2} \left[\sum_{i=1}^2 x_1^{(0)}(i) + \sum_{i=1}^0 x_1^{(0)}(i) \right], \sum_{i=1}^2 x_2^{(0)}(i), \sum_{i=1}^2 x_3^{(0)}(i) \\ -\frac{1}{2} \left[\sum_{i=1}^3 x_1^{(0)}(i) + \sum_{i=1}^2 x_1^{(0)}(i) \right], \sum_{i=1}^3 x_2^{(0)}(i), \sum_{i=1}^3 x_3^{(0)}(i) \\ -\frac{1}{2} \left[\sum_{i=1}^4 x_1^{(0)}(i) + \sum_{i=1}^3 x_1^{(0)}(i) \right], \sum_{i=1}^4 x_2^{(0)}(i), \sum_{i=1}^4 x_3^{(0)}(i) \\ -\frac{1}{2} \left[\sum_{i=1}^5 x_1^{(0)}(i) + \sum_{i=1}^4 x_1^{(0)}(i) \right], \sum_{i=1}^5 x_2^{(0)}(i), \sum_{i=1}^5 x_3^{(0)}(i) \end{pmatrix}$$

$$= \begin{bmatrix} -8646.5 & 214 & 358 \\ -17709 & 394 & 591 \\ -28617 & 589 & 850 \\ -43134 & 895 & 1254 \end{bmatrix}$$

$$Y_5 = [x_1^{(0)}(2), x_1^{(0)}(3), x_1^{(0)}(4), x_1^{(0)}(5)]^T = [7625, 10500, 11316, 17818]^T$$

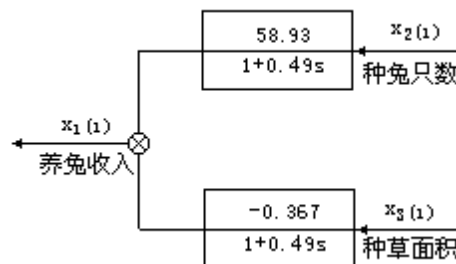
系统识别参数为：

$$\begin{bmatrix} a \\ b_1 \\ b_2 \end{bmatrix} = [X(1,3)^T \cdot X(1,3)]^{-1} \cdot X(1,3)^T \cdot Y_5 = \begin{bmatrix} 2.0399 \\ 119.3953 \\ -0.749 \end{bmatrix}$$

得其微分方程为：

$$\frac{dx_1^{(1)}}{dt} + 2.0399x_1^{(1)} = 119.3953x_2^{(1)} - 0.749x_3^{(1)} \quad (15)$$

相应的动态框图及传递函数为：



综合上述结果, 可得到系统框图及各个环节的传递函数为(图 10-6)。

在图10-6中, $\otimes_3$ 为加入系统中的一个灰色反馈参数。以 $x_1^{(1)}$ 为输出, $u^{(1)}$ 为输入, 系统的传递函数为:

$$\Phi(s) = \frac{x_1^{(1)}(s)}{u^{(1)}(s)} = \frac{19.27(1+0.502s)}{0.1225s^3 + 0.74s^2 + (1.49 - 9.674\otimes_3)s + 1 - 19.27\otimes_3} \quad (16)$$

(二)系统的动态特征分析及优化调控模型

1. 系统的动态特征分析 传递函数 $\Phi(s)$ 的分母为系统的特征多项式, 通过对它的分析, 可以了解该系统的主要动态特征, 考虑到经济生态系统中 3 阶变量的实际意义不大, 故令 $0.1225s^3=0$, 则系统的特征方程为:

$$0.74s^2 + (1.49 - 9.674\otimes_3)s + 1 - 19.27\otimes_3 = 0 \quad (17)$$

系统的特征根为:

$$s_{1,2} = \frac{-(1.49 - 9.674\otimes_3) \pm \sqrt{(1.49 - 9.674\otimes_3)^2 - 4 \times 0.74(1 - 19.27\otimes_3)}}{2 \times 0.74}$$

①若 $(1.49 - 9.674\otimes_3)^2 > 4 \times 0.74(1 - 19.27\otimes_3)$, 则方程有两个不相等的实根, 系统动态一方面受正指数作用, 另一方面受负指数的作用, 因而其发展是有限度的, 会饱和的, 即达到一定限额后, 就会稳定在这一水平上, 不再发展。

②若 $(1.49 - 9.674\otimes_3)^2 < 4 \times 0.74(1 - 19.27\otimes_3)$, 则方程有两个复数根, 系统将出现摆动, 效益时高时低。

③令 $1.49 - 9.674\otimes_3 = 0$, 则 $\otimes_3 = \frac{1.49}{9.674} = 0.154$, 当 $\otimes_3 > 0.154$ 时, 系统的效益将不断增长, 永无止境。不难看出, 这时 $\otimes_3$ 是有下界而无上界的灰色参数调节这个参数, 将使系统保持良好发展。

④ 系统增长的快慢与系数 a 有关:

$$a = \frac{-(1.49 - 9.674\otimes_3)}{2 \times 0.74} = \frac{9.674\otimes_3 - 1.49}{1.48}$$

因为 a 是特征根的实数部分, 若系统的动态有正指数成分, a 必是其中的一部分, 且 a 越大增长越快。

⑤ 特征方程式的根, 若全为负实根, 且 a 是其中绝对值最小者, 则系统的整个动态过程, 也就是由初态发展到稳定态的限额过程, 其时间长短可用

$$\tau = 3 \sim 4 \frac{1}{a} \text{ 进行估计。}$$

2. 系统的优化调控模型 灰色去余控制理论告诉我们, 要实现系统的优化调控, 只有将影响系统理想品质的多余项去掉, 才能达到目的。为此, 需要在原系统的结构模型中加入一个与多余项的传递函数相等, 符号相反的去余项, 以抵消多余项的作用。在原系统的结构模型中, 容易看出, 这个多余项就是在前面加入系统中的灰色参数 $\otimes_3$ 。

通过前面的分析, 我们知道, 系统的传递函数为:

$$\frac{X_1^{(1)}}{U^{(1)}} = \frac{19.27(1+0.502s)}{0.74s^2 + (1.49 - 9.674\otimes_3)s + 1 - 19.27\otimes_3} \quad (18)$$

因为种草养兔业生产属于生物生产过程, 系统的变化较慢, 系统 2 阶变

量也可以暂不考虑，即令 $0.74s^2=0$ ，则系统微分方程可以写成：

$$\begin{aligned} (1.49 - 9.674 \otimes_3) \frac{dx_1^{(1)}}{dt} + (1 - 19.27 \otimes_3)x_1^{(1)} \\ = 19.27 \mu^{(1)} + 9.67 \mu^{(0)} \end{aligned} \quad (19)$$

时间响应函数为：

$$x_1^{(1)}(t) = (x_1^{(0)}(1) - c_1 u^{(1)} - c_2 u^{(0)}) \frac{1 - 19.27 \otimes_3}{e^{1.49 - 9.674 \otimes_3}} t + c_1 u^{(1)} + c_2 u^{(0)} \quad (20)$$

考虑 e 的正指数因数，以保持系统的稳定增长，则：

$$\begin{aligned} 19.27 \otimes_3 > 1 \quad \otimes_3 > \frac{1}{19.27} = 0.052 \\ 9.674 \otimes_3 < 1 \quad \otimes_3 < \frac{1}{9.674} = 0.103 \end{aligned}$$

即，将 $\otimes_3$ 调节在 $0.052 \leq \otimes_3 \leq 0.103$ 的范围内。就可保持系统能够满足我们希望的品质。这一决策表明，经过优化，去余项 $\otimes_3$ 由大于 0.154，且只有下界而无上界的灰色参数，可进一步控制在 0.052—0.103 的灰色区间内，从而为决策者提供了便于择优控制的范围。事实上，从养兔收入中提取 5.2% 至 10.3% 的资金用于系统再生产的投资，在实际生产中是完全可以做到的，而且也清楚地反映出养兔是一种投资小收益大的生产事业，在经济条件尚有困难的地区也是容易推广的富民事业。

第十一章 系统动力学方法

系统动力学(System Dynamics)是美国麻省理工学院福瑞斯特(Jay W. Forrester)教授于 1956 年首创的一种运用结构、功能、历史相结合的方法,通过建立 DYNAMO 模型并借助于计算机仿真而定量地研究高阶次、非线性、多重反馈复杂时变系统的系统分析技术。目前,这一技术已被广泛地应用于自然科学、社会科学以及工程技术研究的各个领域。本章,我们将结合有关实例,介绍和探讨系统动力学方法在地理学研究中的应用问题。

第一节 系统动力学基本原理与方法简介

一、系统动力学的基本观点

(一)系统动力学关于系统的基本观点

1.系统组成 系统动力学认为,系统是由单元、单元的运动及信息组成的。单元是系统存在的现实基础,而信息在系统中发挥着关键的作用,赖以信息的单元才形成系统的结构,单元的运动才形成系统统一的行为与功能。这就是说,系统是结构与功能的统一体。系统动力学所研究的系统,其范围与规模可大可小,其种类可以是各种自然系统、社会系统、经济系统、技术系统、思维系统等以及这些系统相互作用所构成的复合高阶时变系统。

2.系统结构 所谓系统的结构,是指组成系统的各个单元之间相互作用与相互关系的秩序。

在系统动力学中,系统的基本单元是反馈回路。反馈回路是耦合系统的状态、速率(或称行为)与信息的一条回路。它们对应于系统的三个组成部分:单元、运动与信息。任何复杂系统都是由这些相互作用的反馈回路所构成的,这些回路之间相互作用、相互耦合形成了系统的总体功能。其中,构成系统的任何一条反馈回路又都包含了多个反馈环节。按照反馈过程的特点,又可以将这些反馈分为正反馈和负反馈。正反馈的特点是能够产生自我强化的作用机制,负反馈的特点则是能够产生自我抑制的作用机制。具有正反馈特性的回路称为正反馈回路;具有负反馈特性的回路称为负反馈回路。正、负反馈回路的交叉作用机制决定着复杂的系统行为。图 11-1 给出了某土地人口承载力系统中的两条基本反馈回路的正、负反馈作用机制。

3.系统功能 所谓系统的功能,是指系统中各单元活动的秩序,或指单元本身的运动或单元之间相互作用的总体效益。系统动力学是以定性定量相结合的方法研究系统结构,模拟系统的功能的。它从系统的微观构造入手,通过建造反映系统基本结构的模型,进而对系统随时间变化的行为进行模拟研究。建立系统动力学模型的过程,也就是剖析系统的结构与功能之间对立统一关系的过程。

(二)系统动力学关于系统特性的基本观点

1.关于系统的一般特性

(1)总体性与相关性:总体性是系统最基本的特性之一。系统总体不简单地等于其各个组成部分之和。一般而言,系统总体大于部分之和,然而一个失去组织的系统的总体也可能小于部分之和。

相关性是指系统总体与部分、部分与部分、系统与环境之间的普遍相互关系以及单元、运动、信息之间相互关系。在复杂系统中,存在着一因多果、一果多因,甚至多因多果相互交叉的因果关系链。系统动力学采用反馈因果关系代替已往的单向因果关系,这无疑是对系统相关性的进一步认识。

(2)系统的层次与等级:系统动力学强调系统中各单元、各子系统之间的相互联系、相互影响的关系。然而,各单元、各子系统之间还存在着相对的独立性,即系统结构具有层次与等级性。系统结构的层次性、等级性决定了系统功能的层次性、等级性。根据系统的这种层次性与等级性,研究者可以将系统加以划分,从而使无从着手解决的问题,按系统的层次与等级逐级分解。

(3)系统的类似性：系统动力学认为，在自然界与人类社会等不同领域里，各种类型的系统都存在着结构与功能上的类似性，即系统是相似的。这就是说可以用类似的规律和行为模式来描述看来似乎属于截然不同领域内的事物与现象。系统的类似性，决定了不同的系统之间存在着相同的研究模式与方法，这就是结构-功能模拟方法。这也是系统动力学用建立规范化模型的方法去研究和模拟真实系统的一个基本依据。

2. 关于复杂系统的特性

系统动力学方法，主要用于研究与处理那些具有高阶次、非线性和多回路特点的复杂系统。它认为，这些复杂系统具有下述几个方面的特性。

(1)反直观性：所有复杂系统，都毫无例外地表现出反直观的特性。在人们的日常生活思维过程中，所遇到的大多数是关于一阶负反馈系统的经验，人们了解事物的因果关系总是紧密地与时空相关。然而，在复杂系统中，这种简单的因果关系已不复存在，原因与结果的联系在时空上往往是分离的，因而比简单系统复杂得多，也往往被诱入歧途，使人们把系统的某些症候与某一种在时空上贴近的原因联系在一起，但事实上它们并无因果关系。

(2)对变动参数的不敏感性：由于非线性的存在，使得即便是把复杂反馈系统模型的大多数参数加以变动，甚至使部分参数变动数倍，其模型模式也可能无多大变化。复杂系统的这一特性，使得即便是在缺乏严密的基础数据的情况下，有关研究者与决策者也可以通过“会诊”而估计参数，使复杂系统的研究成为可能，从而克服了资料不足给研究工作所带来的困难。

(3)长短期效果的矛盾性：一般而言，由于非线性的作用，使得变更复杂系统内部结构与参数所引起的短期与长期的影响往往是彼此相反的。譬如，在国民经济问题研究中，当涉及到投资时，往往就需要研究积累率问题，积累率是否适当，其影响甚大，持续的高积累率必然会导致国民经济短期的高速增长，但也会导致积累与消费比例的严重失调，从而影响消费，以至制约生产，最终使整个国民经济陷入困境。系统动力学方法在处理复杂系统的这种长短期效果的矛盾方面，有着其它方法无法比拟的独到优点。

二、系统动力学解决问题的过程与步骤

(一)系统分析

(1)了解问题。即回答要解决什么问题？拟达到什么目的？完成此项任务需要哪些条件？现已具备哪些条件？还需要准备哪些条件？等等。

(2)分析系统的基本问题与主要问题、基本矛盾与主要矛盾以及矛盾的主要方面。

(3)初步划定系统的边界，确定内生变量、外生变量和输入变量。一般而言，系统的范围取决于研究目的，系统边界的划定一般是把与建模目的有关的内容圈入系统内部，使其与外界环境隔开。那么，如何才算确定了系统的范围？系统的边界又应划在何处呢？按照系统动力学的观点，划定系统边界的一条基本准则是：应将系统中的反馈回路考虑成闭合的回路。应该力图把那些与建模目的关系密切、变量值较为重要的都划入系统内部。由此可见，划定系统边界之前应首先明确研究的目的。没有目的就无法确定系统的边界。

(4)确定系统行为的参考模式。即用图形表示出系统中的主要变量，并由此引出与这些变量有关的其它重要的变量，通过各方面的定性分析，勾绘出有待研究的问题的发展趋势。由于系统动力学所研究的对象大多数是复杂系

统，其发展趋势很难准确地预测，需要会同各方面专家，集思广益地“会诊”或运用专家咨询法予以解决。一旦参考模式确立，在整个建模过程中，构模者就要反复地参考这些模式，以防研究偏离方向。

(5)调查、搜集有关资料。系统动力学模型被认为是真实系统的“实验室”，要想通过模型模拟和剖析真实系统，获取更丰富、更深刻的信息，进而寻求解决问题的途径，“实验室”的建立是至关重要的。而要建好“实验室”，就必须在认真调查研究的基础上，花大力气搜集、完备各种资料。毫无疑问，为使模型更真实地反映系统，收集的资料应越多越好。但是，要强调的是，资料搜集工作必须紧紧围绕着研究目的进行，如果偏离了研究目的，即使资料再多也是徒劳的，而且还会给资料的筛选带来许多困难。

(二)构建模型

模型的构建，是系统动力学研究、解决问题的关键性的一个步骤。系统动力学模型的建造，一般包括如下两个相互联系的工作环节。

1.分析系统结构 在需要研究的问题已经明确，系统中的重要变量与参考模式已经确定，资料搜集工作也已基本完成之后，就要研究系统及其组成部分之间的相互关系，系统中的主要变量与其它有关变量之间的关系，分析系统的结构。为了使建模工作一开始就能把握整个研究过程的方向，建模者首先要分析系统整体与局部的关系，然后分析变量与变量之间的关系(正关系、负关系、无关系)，最后把这些关系转绘成反映系统结构的因果关系图和流程图。

因果关系图，是反映变量与变量之间因果关系的示意图。(如图 11-1 所示)，其中，变量之间相互影响作用的性质用因果关系键来表示。因果关系键中的正、负极性分别表示了正、负两种不同的影响作用。

正因果关系键 $A \xrightarrow{+} B$ ，表明A的变化使B在同一方向上发生变化，即箭头指向的变量 B 将随着箭头源发的变量 A 的增加(减少)而增加(减少)；负因果关系键 $A \xrightarrow{-} B$ ，表明A的变化使B在相反方向上发生变化，即变量B将随着变量 A 的增加(减少)而减少(增加)。

因果关系键把若干个变量串联后又折回源发变量，这样便形成了一个反馈回路。对于反馈回路，也有正、负极性之区别。如果沿着某一反馈回路绕行一周后，各因果关系键的累计效应为正，则该回路为正反馈回路，反之则为负反馈回路。正反馈具有自我强化的作用机制，负反馈则具有自我抑制的作用机制。

因果关系图虽然能够描述系统反馈结构的基本方面，但不能反映不同性质变量的区别。譬如，状态变量是系统动力学中最重要的变量，它具有积累效应。正是由于状态变量的积累效应，才使系统动力学模型的计算机模拟成为可能。为了进一步揭示系统变量的区别，分别用不同的符号代表不同的变量，并把有关的代表不同变量的各类符号用带箭头的线联结起来，便形成了反映系统结构的流图，如图 11-2 所示。

有了流图，便可以根据一定的规则建立 DYNAMO 模型。在 DYNAMO 模型中，常用的流图符号如图 11-3 所示。

2.建立 DYNAMO 方程

在 DYNAMO 模型中，主要有六种方程，其标志符号分别为：

L 状态变量方程

R 速率方程
A 辅助方程
C 赋值予常数
T 赋值予表函数中 Y 坐标
N 计算初始值

在这些方程中，C，T 与 N 方程都是为模型提供参数值的，并且这些值在同一次模拟中保持不变。L 方程是积累方程，R 与 A 方程是代数运算方程。下面我们将重点介绍 L，R 与 A 方程。

(1) 状态变量方程。在 DYNAMO 模型中，计算状态变量的方程称为状态方程，该方程的基本形式为：

$$LFVEL(\text{现在}) = LEVEL(\text{过去}) + DT(\text{输入速率} - \text{输出速率}) \quad (1)$$

譬如，以图 11-2 所示的流图为例，其方程为：

$$L \text{ LEVEL.K} = LEVEL.J + DT \times (\text{INF LOW.JK} - \text{OUTFLOW.JK}) \quad (2)$$

在(2)式中，LEVEL 表示状态变量，INFLOW 表示输入速率，OUTFLOW 表示输出速率，DT 表示计算时间间隔，亦称时间步长。+、-、 $\times$ 、/ 分别为加、减、乘、除的代数运算符号，J、K、L 作为时间下标主要用以区别时间的先后顺序。K 表示现在，J 表示刚过去的那一时刻，L 表示即将到来的未来那一时刻。DT 表示 J 与 K 以及 K 与 L 之间的时间步长(见图 11-4)。

(2) 速率方程。在状态变量方程中，代表输入(INFLOW)与输出(OUTFLOW)的变量称为速率变量，计算速率变量的代数方程称为速率方程。譬如，人口数量(状态变量)的输入速率(出生率)方程可以写成：

$$R \text{ BIRTHS.KL} = \text{BRF} \times \text{POP.K} \quad (3)$$

在(3)式中，BIRTHS 代表出生率(人/年)，BRF 代表出生率系数(人/年·人)，POP 代表人口数(人)。

(3) 辅助变量方程。在 DYNAMO 模型中，附加的代数运算方程称为辅助方程。“辅助”的涵义就是帮助建立速率方程。一般而言，辅助方程没有统一的标准格式，但是其下标总是 K。辅助变量的值可由现在时刻的其它变量，如状态变量、变化率、其它辅助变量和常量求得。譬如，土地占用率 LF0 的辅助方程式可以写成如下形式：

$$A \text{ LF0.K} = \text{BLDNGS.K} \times \text{LPB} \quad (4)$$

$$A \text{ BLDNGS.K} = \text{BIRTHS.K} \times \text{PBL} \quad (5)$$

在(4)与(5)式中，LF0 代表土地占用率(亩/年)，BLDNGS 代表新建建筑物(座/年)，LPB 代表平均每座建筑物占用土地(亩/座)，BIRTHS 代表每年新增人口数(人/年)，PBL 代表人均占用建筑物(座/人)。

在建立系统动力学模型时，为了使方程书写得井井有条，往往先把方程按照各子块(子系统)书写，书写顺序一般是沿流图按顺时针方向进行。

3. 参数的确定与赋值

DYNAMO 模型中的参数，主要有表函数、初始值、常数、转换系数、调节时间与参考数值等。在运用 DYNAMO 模型对真实系统进行模拟之前，首先应对以上参数赋值。

(三)模型的模拟与评估

当系统动力学模型建构完成以后，经过反复检查各个方程，发现准确无误后，便可将其输入计算机进行调试运行。当模型调试运行通过后，研究者就可以根据其研究的目的，设计不同的方案，运用模型进行模拟运算，对真实系统进行仿真。然而，仿真结果是否可信，其关键是模型本身是否真实、有效。由此可见，对模型的真实性和有效性检验也还是系统动力学仿真研究工作中一个十分重要的环节。

一般而言，在系统动力学研究中，对模型的真实性和有效性检验主要包括如下四个方面。

1. 模型结构适合性检验

(1)量纲检验：系统动力学模型与其它模型一样，决不允许量纲不一致的情况出现。量纲的一致性检验是模型检验的一个最为基本的方面。量纲检验的要求是各变量必须有正确的量纲，而且各个方程式左右两端的量纲必须相同。

(2)方程式极端条件的检验。即检验模型中每一个方程式在其变量的可能变化的极端条件下是否仍有意义。只要在极端条件下，方程运行仍然合理，那么就能确定方程确具有强壮性。

(3)模型边界检验。即主要检验模型所包含的变量与反馈回路是否足以描述所面向的问题和是否符合预定的研究目的。系统边界不宜过大，也不宜过小。如果边界划得过大，就会使模型变得过于复杂，反而模糊了系统结构与动态行为之间的主要关系；而当边界划得过小时，则意味着模型可能忽略了某些重要的变量，或者忽略了富有活力的反馈链。

2. 模型行为适合性检验

(1)结构灵敏度检验。模型结构是决定其行为的主要因素。一般来说，变动模型的结构会对其行为产生较大的影响，模型结构的最大变动即意味着改变系统的边界。但对于系统动力学模型来说，模型行为对结构与相应的方程式的合理变动也不是过于敏感的，而是表现出一定的强壮性，如果模型行为对结构的合理变动过于敏感，则模型不宜作仿真分析之用。这里所谓的“合理变动”是指在某些模型中使参数取极值或变函数为常数时的情况。这些改变意味着这些参数或变函数代表的因果关系键被取消，系统的结构当然也就改变了。即便是在这种情况下，强壮性较好的模型的行为仍然不会有大的变化。

(2)参数灵敏度检验。改变参数对模型行为的影响没有象改变模型结构带来的影响那样大。改变某一参数而不影响模型行为的情况是常见的。系统动力学模型对参数变化是不敏感的。究其原因有二：一是变动某参数时，可能在一段时间内对进行过程中的一部分起作用，但随着主反馈回路的转移，在其余时间或对其余部分不发生任何影响，这时若改变反馈回路中的参数，对系统行为的影响是微乎其微的。二是反馈回路的补偿作用。当改变某一参数时，固然可以加强或削弱某一回路，但由于系统动力学的多回路特点，与此同时将自然而然地加强或削弱其它回路去补偿前述的相反作用，其最后结果则是对模型行为影响甚微或毫无影响。因此，在对系统动力学模型进行参数灵敏度检验时，不应把其它类型的定量模型的高参数灵敏度强加给系统动力学模型，并把灵敏度的高低作为衡量模型精确性的主要标准。与此相反的是，如果系统动力学模型对参数变动很敏感，则只能说明此模型没有实用性。

3.模型结构与真实系统一致性检验 这一工作,主要是请熟悉真实系统的人员参与判定模型结构是否与真实系统相象。如果模型的结构从“外观”上来看与实际系统毫无相似之处,那么即使模型的行为被判定是合适的,也不能认为模型是可信的。

4.模型行为与真实系统一致性检验 该环节检验首先应判断模型行为是否再现最初确立的那些参考模式,如果模型行为与参考模式差别较大,则这种模型再“好”也是无用的。但是,我们必须具体问题具体分析,切勿一遇模型行为与参考模式不符就对模型予以否定。因为模型与参考模式不符的原因有两种:一是模型有误,需要修改完善;二是模型出现的“奇特”行为很有可能是对真实系统的本质反映,而对这种“反映”人们以前从未注意到过。对此必须严格加以证实,如果“反映”的确有意义,而且产生“奇特”行为的机制是真实的,那么模型更是有效的。

系统动力学方法研究问题的过程是一个分解综合,循环反复,逐步实现研究目的的过程。这一过程的繁简及长短与研究对象的复杂程度有关,也与研究目的有关。但是,无论进行何种研究,建模不可能一次成功,即便是一次成功,也需要反复地修改、调试和改进,直至达到满足研究目的要求的模型。如此循环往复的过程,也正是系统动力学对于系统内部结构及其行为关系的认识不断深入的过程。

第二节 应用实例——绿洲型城市发展过程的 S.D. 仿真研究

为了揭示绿洲型城市的发展规律,为绿洲型城市生态经济系统调控提供科学的决策依据,笔者曾与中国科学院新疆地理研究所牛达奎、罗格平等同志合作,从限制与影响绿洲型城市发展的主要环境因素水、土资源出发,运用系统动力学方法,通过仿真模拟实验,把绿洲型城市发展的过程和规模定量地体现出来,并对绿洲型城市发展的模式进行理论上的探讨。该仿真模型由四个状态变量,近 130 个方程式描述,包括城市水资源及污水利用、城市用地、城市工业、城市人口等四个子块。在仿真研究过程中,我们选择了典型的绿洲型城市新疆维吾尔自治区奎屯市,对模型作了真实性和有效性评价,并在此基础上对奎屯市未来发展作了多种方案的仿真实验,为奎屯市的发展和城市总体规划的制定提供了决策依据。

一、仿真模型的建立

(一)系统反馈结构

系统动力学认为,系统的基本结构是反馈回路,反馈回路是系统状态、速率与信息的耦合回路。它们对应系统的三个组成部分是单元、运动与信息。一个复杂系统的结构由若干相互作用的反馈回路耦合而成。反馈回路的交叉、相互作用形成了系统的总功能,反映了系统的动态行为特性。

构造系统的反馈结构,首先需分析系统的整体与局部的关系,将系统划分为若干子块,然后分析诸子块内部及子块之间的因果关系,绘制系统的因果与相互关系图,依据因果关系图,作出系统流图。系统的反馈结构在系统的因果关系图和流图中得以充分体现。

1.系统子块划分 根据系统理论的分解协调原理(即整体与局部关系原理),可将绿洲型城市系统划分成四个子块:

- (1)城市水资源及城市污水处理利用子块;
- (2)城市工业子块;
- (3)城市土地利用子块;
- (4)城市人口子块

2.因果关系图 因果关系是构成系统动力学仿真模型的基础。系统中诸要素(变量)之间的因果关系是体现系统结构和系统整体性的关键环节。

在绿洲型城市仿真模型中,所选择的变量共有 50 多个,各个因素(变量)之间相互作用所形成的因果反馈环路构成了模型的基本结构。限于篇幅,这里只给出其中几个最基本的要素(变量)之间的因果关系回路图(见图 11-5)。图中有 9 个主要反馈环,其中正反馈环有 3 个:

I: 城市人口规模 $\xrightarrow{+}$ 城市污水 $\xrightarrow{+}$ 水资源 $\xrightarrow{+}$ 城市人口规模。

II: 工业规模 $\xrightarrow{+}$ 城市污水 $\xrightarrow{+}$ (未利用)水资源 $\xrightarrow{+}$ 工业规模。

III: 水资源 $\xrightarrow{+}$ 人口规模 $\xrightarrow{-}$ 土地资源 $\xrightarrow{+}$ 工业规模 $\xrightarrow{-}$ 水资源。

反馈有 6 个:

IV: 城市人口规模 $\xrightarrow{-}$ 水资源 $\xrightarrow{+}$ 城市人口规模;

V: 工业规模 $\xrightarrow{-}$ 水资源 $\xrightarrow{+}$ 工业规模;

VI：城市人口规模 $\xrightarrow{-}$ 土地资源 $\xrightarrow{+}$ 城市人口规模；

VII：城市人口规模 $\xrightarrow{+}$ 绿地 $\xrightarrow{+}$ 绿化用水 $\xrightarrow{-}$ 水资源
 $\xrightarrow{+}$ 城市人口规模；

VIII：工业规模 $\xrightarrow{-}$ 土地资源 $\xrightarrow{+}$ 工业规模。

IX：城市人口规模 $\xrightarrow{+}$ 绿地 $\xrightarrow{-}$ 土地资源 $\xrightarrow{+}$ 城市人口规模。

环 I 和 II 反映了城市污水利用增加了水资源的数量进而扩大了城市规模；环 III 说明了城市水土资源和城市规模之间的关系；环 IV 和 V 反映了水资源对人口和工业规模的抑制作用；环 VI 和 VII 说明了土地资源对人口和工业规模扩大的负作用；环 VIII 和 IX 反映了水土资源通过人口规模控制城市绿地规模。

在系统因果关系中，正环使系统振荡；负环使系统趋于稳定，主要的正负环相互耦合形成了系统的基本动态行为。

3. 系统结构流图 系统因果关系图只能描述系统反馈结构的基本方面，不能表示不同性质的变量区别，如哪些是状态变量，哪些是速率变量等。另外，也不能直接利用因果关系图建立数学规范模型，这是因果关系图的弱点。克服这个弱点的方法，是将因果关系能转化成系统结构流图，转化之前最重要的工作是确定哪些是状态变量（即随时间变化的积累量，也称积累变量，它是系统动力学最重要的变量，是结构流图的核心部分），哪些是速率变量（即状态变量中输入输出的变量）哪些是辅助变量（即表述系统内信息传递的变量）等。

考察绿洲型城市系统因果关系图中所有的变量，确定了四个状态变量：城市供水量、工业固定资产、城市人口。城市用地。控制各状态变量相应的速率变量有 6 个：水资源开发利用增长速率；工业总投资和工业折旧；人口增长速率和人口减少速率；土地开发速率。

根据系统子块的划分和因果关系图可分别作出各个子块相应的结构流图（分别如图 11-6、11-7、11-8 和 11-9 所示）。

流图中各变量及参数意义见本节附录。

(二) DYNAMO 方程

在系统动力学中，仿真模型是依据系统结构流图编写的 DYNAMO 方程式。在绿洲型城市系统仿真模型中，共包含近 130 个有效方程，下面我们分别给出各个子块中最重要方程式

1. 城市水资源及污水处理利用子块 该子块是绿洲城市模型最重要的子块，它描述了水资源开发利用状况及城市各类用水的供需状况。水资源的开发利用速率受水资源利用潜力和城市用水供需差控制。水资源利用潜力愈高和用水供需差愈大，开发速率就愈大，否则就愈小。城市污水处理、重复利用，又使城市用水供需减小。水资源利用潜力可用水资源利用程度表示（即城市用水/城市潜在可利用的水资源）。利用程度愈高，潜力愈小。利用程度为 100%，理论上讲，水资源利用潜力为 0。

该子块主要 DYNAMO 方程如下：

L WS.K=WS.J+DT*YIMQWR.JK 城市供水

R YIMQWR.KL=DSDW.K $\times$ EDSWR.K 水资源开发利用速率

A UPW.K=INWS.K+DWSU.K 城市污水

A DSDW.K=DSDDW.K+DSDGLW.K+DSDIW.K 城市用水供需差

A EDUWR.K=WS.K/UWR 水资源开发利用程度

2.城市土地利用子块 该子块反映了城市土地开发利用的状况及控制城市土地开发利用的主要因子，同时体现了城市各项用地供需关系。土地资源的开发利用速率受土地资源利用潜力和城市用地供需差的控制。它们对土地资源开发速率的控制作用与水资源利用潜力和城市用水供需差对水资源开发利用速率的控制作用相似。

该子块主要 DYNAMO 方程如下：

L UL.K=UL.J+DT0 $\times$ RDL.JK 城市用地

R RDL.KL=DSDUL.K $\times$ EDULR.K 土地开发利用速率

A DSDUL.K=DSDULP.K+DSDNP.K+DSDGL.K 土地供需差

A DULU.K=UL.K/ULU 土地开发利用程度

3.城市人口子块 该子块体现了绿洲型城市人口增长的规律，城市人口的自然增长率相对于人口的机械增长率来说是一个相对稳定的值，而人口的机械增长率变化却很大，它主要受就业率、生活条件及城市环境的影响。当就业率低，生活条件和城市环境差的情况下，净迁入人口少，甚至为 0 或为负值，这时城市人口的增长就基本上停止，人口变化渐趋稳定。因此，可以认为人口增长的规律符合 S 形增长趋势见图 11-10。

该子块主要 DYNAMO 方程式如下：

L POP.K=POP.J+DT $\times$ (INC JK-DBC.JK)城市人口

R INC.KL=BR.K+IM.K 人口增长速率

R DEC.KL=DE.K+EM.K 人口减少速率

A IM.K=POP.L $\times$ IMCOE $\times$ EECRIM.K $\times$ ERWRIM.K $\times$ ERNPIM.K

人口迁入

A EM.K=POP.K $\times$ EMCOE $\times$ EECREM.K $\times$ ERWREM.K $\times$ ERNP EM.K

人口迁出

4.工业子块 该子块主要说明了工业固定资产的增长规律、工业用水供需状况及工业用地如何影响工业投资的实施和利润转化问题。速率变量工业总投资转化为固定资产及固定资产折旧均有延迟现象，在模型中用 DELAY1 或 DELAY3 表示。这就是系统的重要特性之一——时滞性。同人口增长规律相似，工业规模也体现了 S 型增长趋势，这些结论在后面的仿真结果中可得以说明。

该子块主要的 DYNAMO 方程

L FIX.K=FIX.J+DT $\times$ DGRINV.JK-DDEPRE.JK 工业固定资产

R DGRINV.KL=DELAY1(GROINV.JK, DEL1) 工业投资延迟

R GROINV.KL=DELAY3(DEPRE.JK, DEL2) 工业投资

R DDEPRE.KL=DELAY3(DEPRE.JK, DEL2) 工业折旧延迟

R DEPRE.KL=FIX.K $\times$ RDEPRE 工业折旧

A GVIOP.K=FIX.K $\times$ YPUEA 工业总产值

二、模型模拟调试、模型评价与政策分析

在调试模拟仿真模型之前，首先应对模型中所有的参数(包括常数、表函

数、初始方程)赋值,为此我们选定了奎屯市,模型中的参数均源自于这个具有现实意义的典型绿洲型城市,模型真实性、有效性评价及决策分析也以奎屯市为依据。

(一)奎屯市简介

奎屯市位于天山北麓,准噶尔盆地西南缘,地貌上,西部属于奎屯河冲洪积扇东半部分的中下部和奎屯河冲积平原的中上部、东部属于安集海河古冲积扇西半部分的中下部。整个城市规划面积约 200km^2 。奎屯是新疆工业生产较为发达的城市,欧亚大陆桥贯通后,它是北疆铁路最大的区段编组站,未来将成为我国通向西亚和苏欧的桥头堡,成为北疆地区的交通枢纽和重要的物资集散地。新疆维吾尔自治区人民政府已拟设奎屯市为经济开发区,将来可望成为我国对外贸易的重要基地。所有这些都将为奎屯市的发展开辟广阔的前景。因此开展对奎屯市的研究具有突出的现实意义。

1989年,全市有非农业人口 5.74 万人,人口的自然增长率(1980—1989 年)平均为 6.8% 。净迁入率(1980—1989 年) 15% ,1989 年工业总投资 2107 万元,固定资产 1989 年净值为 13894 万元,工业总产值为 21080 万元(80 年不变价)、34260 万元(现价),建成区面积 18 平方公里,约有 40—50%的土地面积未利用,职工总数 89 年为 3.43 万人,城市规划范围 200 平方公里,城市规划范围内地下水的允许开采量约为 0.9 亿立方米,目前地下水开采量为 2950 万立方米(城市用水 980 万立方米,余为农业用水)。

(二)模型调试、评价与决策分析

仿真运算分两步进行,第一步仿真 1980—1989 年奎屯市发展状况,检验模型的真实性和有效性;第二步,用经过检验的模型,在不同决策方案下,实施各种政策,仿真 1990—2010 年或 2025 年奎屯市的发展状况,对各种政策作出可行性的论证。

1.模型调试和修改及参数确定 综合分析奎屯市 1980—1989 年 10 年统计资料,确定各类参数的合理范围,对模型进行调试和改进。模型的改进指的是修改模型中不合理的反馈结构和参数取值,例如将部分模型的参数修正为表函数,使参数更符合实际意义。一般来说,模型的结构是决定模型行为的主要因素。对修改后的奎屯市模型进行调试,如果不改变模型的结构,当参数在合理范围内变化时,发现模型的行为模式或图形基本上不变,变化的只是仿真结果的数值,也就是说模型行为模式对参数的变化不敏感,那么,就可认为模型的强壮性很好,对模型的调试和改进也就可以停止了。

第一步仿真(即模型调试)过程中的参数估计,主要是依据 80—89 年奎屯市的统计资料,主要的参数确定如下:

(1)奎屯市可利用的地下水资源 0.9 亿立方米(保守的数字);

(2)水资源利用程度对水资源开发速率的影响(表函数):

A $\text{EDUWR.K} = \text{TABHL}(\text{TEDUWR}, \text{DUWR.K}, 0, 1.6, 0.2)$

T $\text{TEDUWR} = 1/1/1/1/0.7/0.01/0.2/0.5/-0.8$

(3)城市生活用水、工业用水、绿化用水各自占城市总用水的比例分别为 0.2, 0.35, 0.45;

(4)年人均用水 32m^3 、绿化亩均用水 800m^3 ,工业万元产值新鲜用水 140m^3 ;

(5)工业重复利用率 86 前为 3%,后为 8%;

(6)可供城市利用的土地面积为 $1.333 \times 10^4 \text{ha}$;

(7)土地利用程度对土地开发利用速率的影响(表函数)

A EDULR.K=TABHL(TEDULR, DULU.K, 0, 1.6, 0.2)

T TEDULR=1/1/1/1/0.95/0.8/0.001/0/0

(8)城市非生产性用地、生产性用地(工业用水)和绿化用地分别占已利用土地比例为 0.38、0.32、0.30；

(9)平均每百元固定资产占地 0.033ha，人均拥有非生产性用地 0.006ha；人均需绿地 0.006ha；

(10)单位产值利税率 0.16；

(11)单位固定资产(元)工业产值 2.62 元；

(12)折旧率 0.03(因为有一部分折旧资金作为工业投资又重新转化成了固定资产，故纯折旧率很低)；

(13)城市人口自然增长率为 6.8‰，机械增长率系数 15‰，(不同于机械增长率)；

(14)劳动力系数为 0.57；

(15)职工人均占有固定资产(表函数给出)：

A B.K=TABHL(TB, TIME.K, 1980, 1989, 1)

T TB=1140/1270/1420/1600/1350/2200/2470/2950/3200/4010

1980—1989 年期间的仿真结果见表 11-1。

2.模型真实性、有效性的评价 系统动力学模型真实性评价，首先要和需要解决的问题和建模的目的联系起来；其次看模型是否和预定要描述的绿洲型城市系统一致。从模型结构及仿真结果输出的形式可以看出。模型是能够较为圆满地解决所提出的问题，并达到了预定

表 11-1 1980—1989 年奎屯市系统仿真结果

TIME	WS	INWS	INWD	DWSU	GLWS
E00	E06	E03	E03	E03	E03
1980	4.900	1715	804	980	2217
1985	4.931	1725	1034	986	2231
1990	5.205	1821	1265	1041	2355
1983	5.551	1942	1533	1110	2511
1984	5.953	2083	1847	1190	2693
1985	6.425	2248	2223	1284	2907
1986	7.123	2493	2581	1424	3252
1987	7.672	2685	3110	1534	3503
1988	8.477	2966	3723	1695	3870
1989	9.427	3290	4355	1885	4304
TIME	EIX	GV10P	GROINV	INDPRO	
E00	E06	E06	E06	E06	
1980	22.40	58.69	7.099	9.390	
1981	28.83	75.53	9.136	12.084	
1982	35.25	92.36	10.745	14.778	
1983	42.70	111.37	12.888	17.899	
1984	51.46	134.82	15.531	21.571	
1985	61.94	162.28	18.695	25.966	
1986	74.59	195.43	20.977	31.269	
1987	89.89	235.50	19.711	37.681	
1988	107.61	281.93	18.386	45.109	
1989	125.86	329.74	15.331	52.759	
TIME	POP	LF		PE	
E00	E03	E03		E03	
1980	39.800	22.686		19.649	
1981	40.216	22.923		22.698	
1982	41.329	23.558		24.827	
1983	42.637	24.303		26.687	
1984	44.154	25.168		38.117	
1985	47.264	26.940		28.155	
1986	49.402	28.159		30.199	
1987	52.084	29.688		30.470	
1988	55.378	31.565		33.627	
1989	59.147	33.714		31.385	
TIME	UL	ULP	ULNP	LDNP	GLS
E00	E02	E02	E02	E02	E02
1980	5.34	2.00	1.68	2.38	1.58
1981	5.52	2.09	1.76	2.41	1.65
1982	5.78	2.19	1.85	2.48	1.73
1983	6.13	2.33	1.96	2.55	1.84
1984	6.54	2.48	2.09	2.64	1.96
1985	7.01	2.66	2.24	2.85	2.10

的目的。至于模型和绿洲型城市系统的一致性检验要从多种角度进行检验，最有效的方法就是看仿真结果(数据)和实际数据是否一致，因为系统动力学认为检验模型与实际数据一致性是检验模型真实性工作中的重要方面。经验证，仿真结果和实际数据基本一致，相对误差较小，例如，城市人口 80—89 年仿真值与实际数据相对误差为 2.3%，工业固定资产为 4.7%。

评价模型的最终标准是其实用性与有效性，真实性较好的模型不一定实用，但真实性差的模型不具有有效性和实用性。对于模型有效性的检验是要面向政策，看政策的实施在实际中的客观效果是否与模型分析(仿真结果)基本一致。

3. 借助模型进行仿真实验研究——决策分析 绿洲型城市系统动态模型用于决策分析主要包括面向问题的绿洲型城市动态发展方案的确定、方案仿真模拟和可行性分析，为城市总体规划提供依据。

城市总体规划是一定期限内，城市发展的目标和计划，是城市建设的综合部署和城市建设的依据。城市总体规划中最重要的内容有二，一是从城市发展的特点出发，提出长远发展设想和区域规划的各项指标，确定城市的性质和发展规模；二是合理选择城市各项建设用地，确定城市结构。考虑城市发展的方向。

在绿洲型城市仿真模型中，上述城市规划内容都得以反映，依据模型作绿洲型城市发展的决策分析，我们认为可很好地为城市的总体规划服务，我们同样选择奎屯市作为模型用于决策分析的实例。

确定仿真方案之前，我们给出最重要的政策参数，它们取不同的值，可以认为是不同的方案。

主要政策参数：城市可利用的水资源数量、工业用水比例、生活用水比例、绿化用水比例、每万元产值用水量、工业水重复利用率、城市可利用的土地资源、生产性用地比例、非生产性用地比例、绿地比例、人均需绿地面积、人口迁入系数、职工人均占有工业固定资产等。

仿真方案 I：按奎屯市现状仿真 1980—2010 年奎屯市的发展过程和规模。结果见表 11-2。现状参数取值同第一步仿真中的参数。

表 11-2 方案 I 仿真结果

TIME	WS	INWS	INWSU	DWSU	GLWS
E00	E06	E06	E06	E06	E06
1980	4.90	1.71	0.80	0.98	2.21
1985	6.42	2.24	2.22	1.28	2.90
1990	10.44	3.65	4.97	2.08	4.76
1995	16.18	5.66	7.47	3.23	7.38
2000	24.23	8.48	11.20	4.84	11.06
2005	36.40	12.74	16.91	7.28	16.62
2010	54.66	19.13	25.39	10.93	24.95
TIME	FIX	GVIOF	GROINV	INDPRO	
E00	E06	E06	E06	E06	
1980	22.40	58.7	7.090	9.39	
1985	61.94	162.3	18.695	25.97	
1990	143.64	376.3	13.959	60.21	
1995	216.03	566.0	25.997	90.56	
2000	326.27	854.8	37.004	136.77	
2005	488.76	1280.6	56.493	204.89	
2010	733.87	1922.7	84.870	307.64	

仿真方案 II：依据奎屯发展的总趋势，参照其它绿洲城市的资料，确定政策参数变化情况如下：

(1)工业用水比例逐年递增至 2010 年达 48%，工业用水重复利用率由 1989 年的 8%提高到 42%；

(2)绿化用地随着城市的发展，城市人口密度的增加，人均拥有绿地面积具有逐年减小的特性，由 1989 年的 61m²/人，减少至 2010 的 40m²/人。

(3)职工人均占有工业固定资产由 1989 年的 4010 元提高到 5900 元/人以上指标的逐渐变化可用表函数表示。上述参数渐变必然引起水土资源分配上的变化，我们同样以表函数的形式，给出相应的变化，以适应和满足上述变化和要求。

(4)生活用水比例仍为 20%。

(5)绿化用水比例：

A PGLRW.K=TABHL(TPALRW, TIME.K, 1980.2010, 5)

T TPGLRW=0.45/0.45/0.45/0.42/0.38/0.35/0.32

(6)工业用水比例：

A PIRW.K=TABHL(TPIRW, TIME.K, 1980, 2010, 5)

T TPIRW=0.30/0.35/0.35/0.35/0.32/0.42/0.45/0.41

(7)生产性用地比例：

A PULP.K=TABHL(TPULP, TIME.K, 1980, 2010, 5)

T TPULP=0.30/0.30/0.30/0.315/0.33/0.345/0.36

(8)绿化用地比例：

A PGLS.K=TABHL(TPGLS, TIME.K, 1980, 2010, 5)

T TPGLS=0.30/0.30/0.30/0.285/0.255/0.24

(9)非生产性用地比例为 40%。

仿真结果见表 11-3。

表 11-3 方案 II 仿真结果

TIME	WS	TNWS	INWD	DWSU	GLWS
E00	E06	E06	E03	E06	
1980	4.90	1.71	0.79	980	2.21
1985	6.33	2.21	2.16	1267	2.87
1990	10.12	3.54	4.73	2025	4.63
1995	14.76	5.61	7.09	2953	6.39
2000	20.87	8.76	10.71	4175	8.35
2005	29.86	13.44	16.33	5973	11.22
2010	40.94	19.65	24.26	8189	14.37
TIME	FIX	GVIOP	GROINV	INDPRO	
E00	E06	E06	E06	E06	
1980	22.4	58.7	7.090	9.39	
1985	61.9	162.4	18.708	25.98	
1990	143.2	375.4	15.638	60.06	
1995	228.0	597.6	30.574	95.61	
2000	365.8	958.5	48.837	153.35	
2005	583.7	1529.3	72.305	244.69	
2010	900.1	2358.3	97.912	377.34	
TIME	POP	LF	PE		
E00	E03	E03	E03		
1980	39.80	22.69	19.65		
1985	45.79	26.10	28.17		
1990	62.10	35.35	35.73		
1995	90.46	51.56	50.68		
2000	133.02	75.82	73.16		
2005	193.49	110.29	106.13		
2010	274.06	156.20	152.56		
TIME	UL	ULP	ULNP	LDNP	GLS
E00	E02	E02	E02	E02	E02
1980	5.26	1.58	2.10	2.07	1.58
1985	7.20	2.16	2.88	3.41	2.16
1990	10.95	3.28	4.38	4.63	3.28
1995	16.28	5.13	6.51	6.75	4.64
2000	23.80	7.85	9.52	9.93	6.42
2000	35.06	12.09	14.02	14.44	8.94
2010	47.14	18.10	20.11	20.46	12.06

仿真方案III：在上述两方案中，奎屯市利用的水资源主要是地下水，如果能利用部分地表水(假设数量为 1000 万立方米)，另外奎屯市 20 世纪末至 21 世纪初是人口增长的高峰期，主要表现是人口的大量迁入，人口的迁入系

数用表函数给出：

A IMCOE.K=TABHL(TIMCOE, TIME.K, 1980, 2010, 5)

T TIMCOE=0.029/0.029/0.030/0.033/0.039/0.035/0.033

而迁出系数为 0.015，则 1980—2010 年仿真结果见表 11-4。

仿真方案IV 由于奎屯市可利用的地下水资源约为 0.9 亿立方米(保守的数字)，依方案III仿真结果，至 2010 年预测城市总用水量为 5600 万立方米，若农村用水(农业用水)和现在持平仍为 2000 万立方米,则真正能供城市利用的地下水只有 0.7 亿立方米，再加上利用的 1000

表 11-4 方案III仿真结果

TIME	WS	INWS	DNWD	DWSU	GIWS
E00	E06	E06	E06	E06	E06
1980	4.90	1.71	0.79	0.98	2.21
1985	6.58	2.30	2.35	1.81	2.99
1990	11.43	4.00	5.32	2.28	5.23
1995	17.50	6.65	8.40	3.50	7.58
2000	26.13	10.97	13.42	5.22	10.45
2005	39.15	17.61	21.46	7.33	14.71
2010	56.94	27.33	33.37	11.38	19.98
TIME	FIX	GVIOP	GROIINV	INDPRO	
E00	E06	E06	E06	E06	
1980	22.4	58.7	7.87	9.39	
1985	67.6	177.0	21.64	28.32	
1990	161.2	422.4	19.99	67.59	
1995	270.2	707.9	40.65	113.27	
2000	458.3	1200.7	67.78	192.12	
2005	766.7	2008.8	104.50	321.40	
2010	1237.9	3243.4	153.96	518.94	
TIME	POP	LF	PE		
E00	E03	E03	E03		
1980	39.80	22.69	19.65		
1985	47.76	27.22	30.70		
1990	71.36	40.67	40.21		
1995	109.55	62.44	60.04		
2000	169.24	96.47	91.66		
2005	257.72	146.90	139.40		
2010	383.95	218.85	209.82		
TIME	UL	ULP	ULNP	LDNP	GLS
E00	E02	E02	E02	E02	E02
1980	5.26	1.53	2.10	2.97	1.58
1985	7.45	2.23	2.98	3.56	2.23
1990	12.38	3.71	4.95	5.32	3.71
1995	19.36	6.10	7.74	8.18	5.52
2000	29.84	9.84	11.93	12.63	8.05
2005	46.02	15.87	18.40	19.24	11.73
2010	69.19	24.91	27.67	28.66	16.60

万立方米的地表下，共计 0.8 亿立方米，也就是说在 2010 年，城市水资源的供给有盈余，若将时间延迟到 2025 年，仿真结果见表 11-5 和图 11-11。由表 11-5 可知，在 2015—2020 年期间，水资源供需处于平衡状态，从理论上讲，此时城市发展规模基本停止了，而此时城市用地为 $1.0378 \times 10^4 \text{ha}$ ，小于城市可利用土地。也就是说，土地对绿洲型城市发展的限制作用仍然很小。

4. 奎屯市仿真方案评价 尽管我们对模型的真实性和有效性进行了评

价，但对各种仿真方案还应进行评价，对模型模拟的政策进行可行性分析，避免制定无法在实际系统中实施的政策。

方案Ⅰ为现状仿真方案，城市无论怎样发展，未来不同于现在，因此，模型指标的变化不可能等同于现状参数值，如工业用水比例，在城市发展的初期(80 年代)较小，随着城市的

表 11-5 方案Ⅳ仿真结果

TIME	WS	INWS	INWD	DWSU	GLWS
E00	E06	E06	E06	E06	E06
1980	4.90	1.71	0.79	0.98	2.21
1985	6.58	2.30	2.35	1.31	2.99
1990	11.43	4.00	5.32	2.28	5.23
1995	17.50	6.65	8.40	3.50	7.58
2000	26.13	10.97	13.42	5.22	10.45
2005	39.15	17.61	21.46	7.83	14.71
2010	56.94	27.33	33.37	11.38	19.98
2015	77.22	37.07	50.39	15.44	27.10
2020	80.06	38.43	63.45	16.01	28.10
2025	80.07	38.43	60.91	16.01	28.10
TIME	FIX	GVIPP	GR0INV	INDPRO	
E00	E06	E06	E06	E06	
1980	22.4	58.7	7.87	9.39	
1985	67.6	177.0	21.64	28.32	
1990	161.2	422.4	19.99	67.59	
1995	270.2	707.9	40.65	112.27	
2000	458.3	1200.7	67.78	192.12	
2005	766.7	2008.8	104.50	321.40	
2010	1237.9	3243.4	153.96	518.94	
2015	1869.2	4897.2	144.24	783.55	
2020	2353.4	6165.9	9.07	986.55	
2025	2259.3	5919.3	61.13	947.09	
TIME	POP	LF	PE		
E00	E03	E03	E03		
1980	39.80	22.69	19.65		
1985	47.76	27.22	30.70		
1990	71.36	40.67	40.21		
1995	109.55	62.44	60.04		
2000	169.24	96.47	91.66		
2005	257.72	146.90	139.40		
2010	383.95	218.85	209.82		
2015	570.99	325.47	316.81		
2020	749.01	426.93	398.88		
2025	812.52	463.14	382.93		
TIME	UL	ULP	ULNP	LDNP	GLS
E00	E02	E03	E03	E03	E03
1980	5.26	1.58	2.10	2.97	1.58
1985	7.45	2.23	2.98	3.56	2.23
1990	12.38	3.71	4.95	5.32	3.71
1995	19.36	6.10	7.74	8.18	5.52
2000	29.84	9.84	11.93	12.63	8.05
2005	46.02	15.87	18.40	19.24	11.73

发展,工业规模的扩大,工业用水比例必然逐渐提高。因此,现状仿真中,工业规模(如固定资产、工业总产值)因受工业用水供需矛盾突出的制约,使得仿真结果偏小。方案Ⅰ可作为决策人员的参考方案。

奎屯是新兴城市、工业基础较薄弱,以轻工业为主,奎屯市政策又比较重视节水工作,使得奎屯市工业万元产值用新鲜水量较低为 140m^3 (乌鲁木齐市为 380m^3)。随着工业的发展工业结构也将有所改变,会出现一部分重工业企业。这样将提高万元产值用水量,使得工业,用水增加。因此我们认为方案Ⅱ、Ⅲ、Ⅳ仿真结果中城市总用水量仿真值偏小,而且工业规模偏大,城市用水供需矛盾将提前来临。作为奎屯市政府每年能够利用多少地表水,这是一个未知数,也许在城市用水紧张的情况下,会增加地表水量的利用,在这种情况下,城市规模,还可继续扩大,这就另当别论了。

对于影响奎屯市人口迁移的因素,在模型中,只从就业率,生活条件,城市环境三方面进行分析,很少考虑决策人的行为因素。在制定迁入系数时,也只是依据奎屯市的现状,考虑奎屯市未来人口的发展趋势,参考城市人口迁移的规律确定的,是否适合,目前还很难断定。从仿真结果看,方案Ⅱ,2010年为27.4万人;方案Ⅲ,2010年为38.4万人,从城市人口的自然增长率和机械增长率系数来看,两方案差别很小,但方案Ⅲ的迁入系数或迁出系数均高于方案Ⅱ8—9个千分点,说明就业率、生活条件、城市环境对人口的迁入迁出影响很大。现列出方案Ⅱ和方案Ⅲ迁入迁出系数,以作比较。

方案Ⅱ, 迁入系数

A IMCOE.K=TABHL(TIMCOE, TIME.K, 1980, 2010, 5)

T TIMCOE=0.021/0.021/0.021/0.024/0.030/0.026/0.024/

迁出系数 EMCOE=0.006

方案Ⅲ, 迁入系数

A IMCOE.K=TABHL(TIMCOE, TIME.K, 1980, 2010, 5)

T TIMCOE=0.029/0.029/0.033/0.039/0.035/0.033

迁出系数 CMCOE=0.015

不管怎样,对仿真方案进行全面的评价还较困难,但时间会对仿真方案作出全面的评价,本文只对奎屯市发展提出了四套方案,不可能给出所有的方案,当然如果城市发展的决策人员提出城市发展的其它仿真方案,还可以继续依据模型进行仿真模拟,其目的是为了说明模型如何明了地为决策者在制定城市规划时提供依据、提供参考服务,当然也是对我们所研制的模型全面的检验和考验。

三、结论与建议

(一)基本结论

通过对绿洲型城市的系统分析及奎屯市的动态发展的仿真模拟,我们得到有关绿洲型城市基本结论。

(1)新疆绿洲城市生态环境脆弱,维持城市生态系统的良性循环是绿洲型城市系统存在发展的前提和基础。

(2)水土资源始终是限制城市发展规模最主要的因子,尤其是水资源,表现为一是各时期水资源供给与需水量对比关系的制约限制,二是受水资源开

发潜力的制约，目前乌鲁木齐的城市发展已经印证了这一点。水资源的开发一旦达到潜力时，城市发展就趋于和缓，城市规模的扩展就基本上停止了。

(3)从仿真结果来看，绿洲型城市的发展速度取决于社会经济条件。社会经济条件(如投资环境、工业基础等)优越，则城市发展的速度就快，城市发展的最终规模就提前来临。

(4)人口的发展对绿洲城市的运行起着重要的作用。一方面，人口的数量体现了城市实际负荷状况(即城市规模)，另一方面城市人口直接参与系统运行，成为诸多反馈环上的重要环节。当资源开发潜力为0时，理论上说，城市人口规模的发展就停止了。

(二)建议

(1)必须高度重视城市水资源的保护、管理和开发，在合理开发水资源的基础上，积极发展节水型工业，以轻工业为主，尽量不建或少建耗水量大的企业，而且某些用水量大的轻工企业也应尽量避免建设。同时，还要降低工业用水定额，不断提高水的重复利用率，使有限的水资源发挥其应有经济效益。

(2)研究城市生态经济系统表现的属性及其评价指标体系。做好城市环境的评价工作，不断地为城市发展的决策者提供依据，制定切实可行的相应政策和措施。例如做好新建项目优化选址和环境影响评价，为工业及其它行业的合理布局提供服务。

(3)城市应广泛地实行人口计划政策，适当控制城市人口的机械迁移速率，使人口的增长适应且满足城市的需求，避免造成人满为患的局面，避免由此而引起的各类环境问题。

附录：模型变量及参数释义

(1)状态变量(state variable)

WS：城市供水量(立方米)

UL：城市用地(公顷)

POP：城市人口(人)

FIX：工业固定资产(元)

(2)速率变量(Rate variable)

YIMQWR：年新增开采量(或年供水新增量)(立方米/年)

RDL：土地开发利用速率(公顷/年)

INC：城市人口增长速率(人/年)

DEC：城市人口减少速率(人/年)

GROINV：工业总投资(元/年)

DGRINV：工业总投资一阶延迟(元/年)

DEPRE：工业纯折旧(元/年)

DDEPRE：工业折旧三阶延迟(元/年)

(3)辅助变量

DUWR：水资源利用程度

INWS：工业供水(立方米/年)

INWD：工业需水(立方米/年)

DWSU：城市生活供水(立方米/年)

DWDE：城市生活需水(立方米/年)

GLWS：城市绿化供水(立方米/年)

GLWD：城市绿化需水(立方米/年)
DSDW：城市用水供需差(立方米/年)
DSDIW：工业用水供需差(立方米/年)
DSDDW：生活水供需差(立方米/年)
DSDGLW：绿化用水供需差(立方米/年)
GLRW：绿化配给新鲜水(立方米/年)
UPW：城市污水
RUSW：重复利用的污水
URSW：污水重复利用率
QDFSW：外排环境的处理污水(立方米/年)
USWIN：工业利用的污水(立方米/年)
URIW：工业用水重复利用率
UWPUOP：单位产值新鲜用水(立方米/万元)
USWGL：城市绿化利用的处理污水(立方米/年)
DULU：城市土地利用程度
ULP：生产性用地(限于工业用地)(公顷)
LOP：生产性需地(限于工业)(公顷)
LDNP：非生产性需地(公顷)
ULNP：非生产性用地(公顷)
GLS：绿化用地(公顷)
GLD：绿化需地(公顷)
DSDULP：生产性用地供需差(公顷)
DSDNP：非生产性用地供需差(公顷)
DSDGL：绿化用地供需差(公顷)
DSDUL：城市用地供需差(公顷)
LF：适龄劳动力(人)
PE：就业人口(人)
EC：就业率
IM：迁出率(人/年)
EM：迁入率(人/年)
RSDDW：生活用水供需比
BR：出生率(人/年)
DE：死亡率(人/年)
RSUPL：生产性用地供需比
GVIOP：工业总产值(元/年)
RSDIW：工业用水供需比
INDPRD：工业利税(元/年)
INDINV：工业投资(元/年)
EXTINV：外部投资(元/年)
(4)表函数(Table function)
EDUWR：水资源利用程度对其开发的影响
EURIWU：工业用水重复利用率对工业单位产值需水量的影响
EDULR：土地利用程度对土地开发速率的影响
ERWRIM：生活用水供需比对人口迁入的影响

ERWREM：生活用水供需比对人口迁出的影响
EECRIM：就业率对人口迁入的影响
EECREM：就业率对人口迁出的影响
ERNPIM：生活性用地供需比对人口迁入的影响
ERNPEM：生活性用地供需比对人口迁出的影响
ERLPIN：生产性用地供需比对工业投资的影响
ERSDIW：工业用水供需比对工业投资的影响

(5)参数

UWR：城市可利用的水资源(立方米)
PIRW：工业用水比例(1/年)
PISW：工业利用污水比例(1/年)
UCP：城市污染物容量(立方米)
PGLRW：绿化用水比例(1/年)
QUWPU：工业单位产值需水定额(立方米/万元)
PGLSWE：绿化利用的污水(立方米/年)
WDPUGL：单位绿地面积需水量(立方米/公顷·年)
PDWSU：生活供水比例(1/年)
ADDWEP：人均年需水量(立方米/人·年)
ULU：城市可利用的土地资源(公顷)
PGLS：绿化用地比例(1/年)
GLDPP：人均需绿地(公顷/人·年)
PULNP：非生产性用地比例(1/年)
LDPP：人均需生活用地(公顷/人·年)
PULP：生产性用地比例(1/年)
LDPUFA：单位固定资产需地(公顷/元)
BRCOE：出生率系数(1/年)
DECOE：死亡率系数(1/年)
IMCOE：额定迁入系数(1/年)
EMCOE：额定迁出系数(1/年)
PLF：劳动力比例
YPUFA：单位固定资产产值(元/元·年)
PR：工业利税率
PROII：工业投资比例
DEL1：工业投资一阶延迟时间(年)
DEL2：工业折旧三阶延迟时间(年)
RDEPRE：工业折旧率(1/年)

第十二章 投入产出分析方法

投入产出分析，又称“部门平衡”法，或称“产业联系”分析，是由美国经济学家瓦·列昂捷夫(Wassily, W. Leontief)在本世纪 30 年代最早提出来的。它主要通过编制投入产出表及建立相应的数学模型，反映经济系统各个部门(产业间)的关系。这一方法提出以来，特别是自本世纪 50 年代以来，已被广泛地应用于区域产业构成分析、区域之间的相互联系分析、资源利用及环境保护研究等各个方面。目前，投入产出分析方法已成为经济地理学研究的重要方法之一。

第一节 投入产出表与投入产出分析的基本原理

投入产出模型，按照分析时期的不同，可分为静态模型和动态模型两大类。静态模型主要分析、研究某一个时期的再生产过程；动态模型则分析、研究若干时期的再生产过程，并研究各个时期再生过程的相互联系。本章主要阐述静态投入产业模型的原理及其在地理系统分析中的应用。

静态投入产出模型，按照计量单位的不同，可分为实物型和价值型两种型式。在投入产出表中，前者是按实物单位(如吨、担、米等)计量的，而后者则是按货币单位(如元)计量的。这两种类型是最基本而又最能反映投入产出分析特征的模型。

一、实物型投入产出模型

实物型投入产出表，是以各种产品为对象，以不同的实物计量单位编制出来的。表 12-1 是一个简化的实物型的投入产出表。在表 12-1 中，共有几类产品。表的横行，反映了各类产品的生产与分配使用情况，它的一部分作为中间产品供其它产品生产过程中使用，另一部分作为最终产品供积累、消费和出口，两部分之和就是最终一定时间内各类产品的生产总量 q_i 。表的纵列，反映了各类产品生产过程中，消耗的其它产品(包括自己本身与劳动消耗的数量)，由于各类产品的计量单位不一致，所以无法相加。表的第一栏中，表示中间产品之间的流量， q_{ij} 表示第 i 类产品流向第 j 类产品的数量，或者说是第 j 类产品生产过程中消耗的第 i 类产品的数量， q_{ii} 表示各类产品的自身消耗量。

显然，在表 12-1 中，按每一行可以建立一个方程，这样，我们共有 $n+1$ 个方程：

$$\begin{aligned} q_{11} + q_{12} + \cdots + q_{1n} + y_1 &= q_1 \\ q_{21} + q_{22} + \cdots + q_{2n} + y_2 &= q_2 \\ M \quad M \quad \quad M \quad M \quad M \\ q_{n1} + q_{n2} + \cdots + q_{nn} + y_n &= q_n \\ q_{01} + q_{02} + \cdots + q_{0n} &= L \end{aligned}$$

以上方程式可以写成：

表 12-1 简化的实物型投入产出表

产出 ↓投入	中间产品				最终产品	总产品
	1	2	...	n		
1	q_{11}	q_{12}	...	q_{1n}	y_1	q_1
2	q_{21}	q_{22}	...	q_{2n}	y_2	q_2
M	M	M		M	M	M
n	q_{n1}	q_{n2}	...	q_{nn}	y_n	q_n
劳 动	q_{01}	q_{02}	...	q_{0n}	l	L

$$\sum_{j=1}^n q_{ij} + y_i = q_i \quad (i = 1, 2, \dots, n) \quad (1)$$

$$\sum_{i=1}^n q_{0j} = L \quad (2)$$

令

$$a_{ij} = \frac{q_{ij}}{q_j} \quad (i, j = 1, 2, \dots, n) \quad (3)$$

则 a_{ij} 表示每生产单位 j 类产品需要消耗的 i 类产品的数量，它被称为产品的直接消耗系数。同理，劳动的直接消耗系数为：

$$a_{0j} = \frac{q_{0j}}{q_j} \quad (j = 1, 2, \dots, n) \quad (4)$$

将(3)式和(4)式分别代入(1)式和(2)式，可以得到：

$$\sum_{j=1}^n a_{ij} q_j + y_i = q_i \quad (i = 1, 2, \dots, n) \quad (5)$$

$$\sum_{j=1}^n a_{0j} q_j = L \quad (6)$$

若令：

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{pmatrix}$$

$Q = [q_1, q_2, \dots, q_n]^T$, $Y = [y_1, y_2, \dots, y_n]^T$ ，则关于几类产品的生产与分配使用的方程(5)，就可以写成矩阵形式

$$AQ + Y = Q \quad (7)$$

即

$$(I - A)Q = Y \quad (8)$$

(8)式中， I 为 n 阶单位矩阵。矩阵 $(I - A)$ 称为列昂捷夫矩阵，其具体形式为：

$$(I - A) = \begin{pmatrix} 1 - a_{11} & -a_{12} & \cdots & -a_{1n} \\ -a_{21} & 1 - a_{22} & \cdots & -a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ -a_{n1} & -a_{n2} & \cdots & 1 - a_{nn} \end{pmatrix}$$

(8)式表明了总产量与最终产品之间的关系，若已知各类产品的总产量，则可以通过(8)式求出各类产品的最终产品需求量；若已知各类产品的最终产品需求量，则可以通过求解(8)式得到各类产品的总产量：

$$Q = (I - A)^{-1} Y \quad (9)$$

(9)式中， $(I - A)^{-1}$ 称为列昂捷夫逆矩阵，它反映了最终产品与总产量之间的关系。

实物型投入产出模型，建立了经济系统中各类主要产品的生产和分配使用之间的平衡关系。在模型中，直接消耗系数矩阵 A 反映了生产过程的技术结构。模型通过列昂捷夫矩阵 $(I - A)$ 建立了总产品与最终产品之间的联系，通过列昂捷夫矩阵 $(I - A)^{-1}$ 建立了最终产品与总产品之间的联系。

二、价值型投入产出模型

价值型投入产出模型，是根据价值型投入产出表而建立的。它将整个国民经济系统划分为若干子系统——生产部门，并以货币为计量单位。由于它不仅能反映各部门产品的实物运动过程，而且能够精确地描述各部门产品的价值流动过程，因而它的实用性与实用范围比实物型更广。表 12-2 为一简化的价值型投入产出表。

由于价值型投入产出表是以货币(一般用不变价格)作为计量单位的，因此我们可以按行或列来建立数学模型。

根据表 12-2，按横行建立数学模型与实物型是类似的，它也是反映了各部门产品的生产与分配使用情况，描述了最终产品与总产品之间的平衡关系。对于几个部门，有 n 个方程：

$$x_{11} + x_{12} + \cdots + x_{1n} + y_1 = x_1$$

$$x_{21} + x_{22} + \cdots + x_{2n} + y_2 = x_2$$

$$\begin{matrix} M & M & M & M & M \end{matrix}$$

$$x_{n1} + x_{n2} + \cdots + x_{nn} + y_n = x_n$$

以上方程可以写成：

$$\sum_{j=1}^n x_{ij} + y_i = x_i \quad (i = 1, 2, \cdots, n) \quad (10)$$

记各部门的直接消耗系数为：

$$a_{ij} = \frac{x_{ij}}{x_j} \quad (i, j = 1, 2, \cdots, n) \quad (11)$$

将(11)式代入(10)式，得：

$$\sum_{j=1}^n a_{ij} x_j + y_i = x_i \quad (i = 1, 2, \cdots, n) \quad (12)$$

(12)式叫做产品分配方程组，它表明某一部门的总产品等于从该部门流向其它部门的产品及最终产品之和。若记 $X = [x_1, x_2, \cdots, x_n]^T$, $Y = [y_1, y_2, \cdots, y_n]$ 及

表 12-2 价值型投入产出表

↓投入 →产出		中间使用					最终产品	总产值
		部门1	部门2	...	部门n	小计		
物质消耗	部门1	x_{11}	x_{12}	...	x_{1n}	E_1	y_1	x_1
	部门2	x_{21}	x_{22}	...	x_{2n}	E_2	y_2	x_2
	⋮	⋮	⋮		⋮	⋮	⋮	⋮
	部门n	x_{n1}	x_{n2}	...	x_{nn}	E_n	y_n	x_n
	小计	C_1	C_2		C_n	C	y	x
新创造价值	劳动报酬	V_1	V_2	...	V_n	V		
	纯收入	m_1	m_2	...	m_n	m		
	小计	N_1	N_2	...	N_n	N_0		
总产值		x_1	x_2	...	x_n	x		

$$A = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \cdots & \cdots & & \cdots \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{bmatrix}$$

则方程组(12)可以换写成矩阵形式：

$$AX+Y=X$$

即

$$(I-A)Y \quad (13)$$

若假设 $|I-A| \neq 0$ ，则有：

$$X=(I-A)^{-1}Y \quad (14)$$

对于价值型的投入产出表，亦可以按纵列建立模型，反映各部门产品价值形成的过程，反映生产与消耗之间的平衡关系。对于 n 个部门，按纵列同样可以建立 n 个方程：

$$\begin{aligned} x_{11} + x_{21} + \cdots + x_{n1} + V_1 + m_1 &= x_1 \\ x_{12} + x_{22} + \cdots + x_{n2} + V_2 + m_2 &= x_2 \\ M \quad M \quad \quad M \quad M \quad M \quad M \\ x_{1n} + x_{2n} + \cdots + x_{nn} + V_n + m_n &= x_n \end{aligned}$$

以上方程也可以写成：

$$\sum_{i=1}^n x_{ij} + V_j + m_j = x_j (j=1, 2, \cdots, n) \quad (15)$$

(15)式叫做费用平衡方程组，它反映出物质消耗费用、新创造价值与产品总价值之间的关系。将(11)式代入(15)式，并利用 $N_j=V_j+m_j$ ，则(15)式可以写成：

$$\sum_{i=1}^n a_{ij}x_j + N_j = x_j (j=1, 2, \cdots, n) \quad (16)$$

(16)式中， $\sum_{i=1}^n a_{ij}$ 为生产单位 j 部门产品的全部物质消耗系数。

若记物质消耗系数矩阵为：

$$C = \begin{pmatrix} \sum_{i=1}^n a_{i1} & 0 & \cdots & 0 \\ 0 & \sum_{i=1}^n a_{i2} & \cdots & 0 \\ M & M & & M \\ 0 & 0 & \cdots & \sum_{i=1}^n a_{in} \end{pmatrix}$$

并记 $N=[N_1, N_2, \cdots, N_n]^T$ ，则(16)式就可以写成矩阵形式：

$$CX+N=X$$

即： $(I-C)X=N \quad (17)$

(17)式建立了总产值与新创造价值之间的联系。若 $|I-C| \neq 0$ ，则可以建立新创造价值与总产值之间的联系：

$$X=(I-C)^{-1}N \quad (18)$$

利用(17)式，可以在已知物质消耗系数与各部门总产值的情况下，求出各部门的新创造价值；利用(18)式，则可以通过各部门的新创造价值来计算各部门的总产值。

与实物型模型相比，价值型投入产出模型具有以下几个特点：

1. 价值型模型可以包括国民经济系统中的全部物质生产部门，它与实物型只包括大类产品相比，范围更广，因此更能反映国民经济系统的全貌。而且价值型模型具有统一的计量单位，可以根据分析问题的需要与资料取得之可能，将部门进行合并和分解，显得更为灵活一些。

2. 由于计量单位统一，对于价值型投入产出表，既可以按横行建立模型来反映各部门产品的产生与分配使用情况，亦可以按照列建立模型来反映各部门产品价值的形成过程。

3. 价值型投入产出模型，可以同时从产品的使用价值与价值两个方面反映国民经济系统各个部门产品的再生产过程。

4. 价值型投入产出表中的部门是“纯部门”，是根据同类产品的原则来划分的，而不是按行政和企业来划分的。因此，在应用本模型解决有关实际问题时，统计资料的收集和处理一定要注意这一点。

第二节 区域经济活动分析的投入产出模型

区域经济活动分析，是区域地理系统分析的主要内容之一。只要对投入产出分析的基本模型作一些改造、加工，就可以将它应用于区域经济活动分析之中，用以解决区域内与区域外以及区域间的产业联系。

一、区域模型

一个较大的地理区域(如，一个国家、一个地区等)是由许多小的地理区域(如省、市、县)构成的。对于每一个小的地理区域，都可以看成是较大的区域地理系统(即较大的地理区域)的一个分系统。区域投入产出模型就是研究这个较大的区域地理系统的分系统——小的地理区域的经济活动的投入与产出的关系的计量模型。

区域投入产出模型具有如下特点：

(1)部门分类不完整。一个区域，由于受自然资源(如，气候、土地、生物、矿产、能源等等)和历史条件的限制，不一定能生产自己本区域所需的全部产品。

(2)一个区域往往有一个或若干个主导产业部门，例如我国山西的煤炭，山东的石油，甘肃的有色金属工业部门等。这些部门在该区域的经济活动中占有重要的地位。

(3)来自区域之外的输入和区域向外界的输出，在区域经济活动中占有重要的地位。这是因为，第一，区域经济是整个区域地理系统的有机组成部分，各区域之间有着密切的政治和经济联系；第二，区域范围较小，部门不完整。所以，区域模型在结构上的一个重大特点是把输入与输出详细划分，形成模型中的单独部分。

(4)区域的国民收入生产额与使用额可以长期存在很大的差额。例如，在新建工业区中，国民收入的生产额不大，但国民收入的使用额(基建投资)可以很大。

综合以上特点，区域模型的结构如表 12-3 所示。

表 12-3 区域投入产出表

		中间产品					最终产品		
		1	2	...	n	合计	消费	投资...输出	合计
区域生产部门	1	X ₁₁	X ₁₂	...	X _{1n}				Y ₁
	2								
	M	X ₂₁	X ₂₂	...	X _{2n}				Y ₂
		M	M		M				M
	n	X _{n1}	X _{n2}	...	X _{nn}				Y _n
	合 计								
外地输入产品	1	u ₁₁	u ₁₂	...	u _{1n}				w ₁
	2								
	M	u ₂₁	u ₂₂	...	u _{2n}				w ₂
		M	M		M				M
	m	u _{n1}	u _{n2}	...	u _{mn}				w _n
	合 计								
新创造价值	劳动报酬	v ₁	v ₂	...	v _n				
	纯收入	m ₂	m ₂	...	m _n				
	合 计								
总 产 品		x ₁	x ₂	...	x _n				

从表 12-3 的水平方向来看，有如下两种平衡关系式。

(1)本区域生产的产品，其生产与使用平衡方程式为

$$\sum_{i=1}^n x_{ij} + y_i = x_i \quad (i = 1, 2, \dots, n) \quad (1)$$

上式也可以写成：

$$\sum_{j=1}^n a_{ij} x_j + y_i = x_i \quad (i = 1, 2, \dots, n) \quad (2)$$

(2)式中， a_{ij} 表示区域的直接消耗系数。

(2)来自区域以外输入产品使用的平衡方程式：

$$\sum_{j=1}^n u_{ij} + w_i = u_i \quad (i = 1, 2, \dots, m) \quad (3)$$

这里， u_{ij} 表示本区域第 j 部门对来自区域以外的第 i 种产品的消耗量， w_i 表示第 i 种输入产品作为本区域最终产品的数量； u_i 表示第 i 种输入产品的输入总量。令：

$$d_{ij} = \frac{u_{ij}}{x_j} \quad \begin{pmatrix} i = 1, 2, \dots, m \\ j = 1, 2, \dots, n \end{pmatrix} \quad (4)$$

表示对输入产品的直接消耗系数。于是，(3)式可以写成：

$$\sum_{j=1}^n d_{ij}x_j + w_i = u_i \quad (i = 1, 2, \dots, m) \quad (5)$$

从表 12-3 的垂直方向看，有如下关系式：

$$\sum_{i=1}^n x_{ij} + \sum_{i=1}^m u_{ij} + v_j + m_j = x_j \quad (j = 1, 2, \dots, n) \quad (6)$$

(6)式反映了产品的价值构成情况，它可以进一步改写为：

$$\sum_{i=1}^n a_{ij}x_j + \sum_{i=1}^m d_{ij}x_j + v_j + m_j = x_j \quad (j = 1, 2, \dots, n) \quad (7)$$

如果令： $D = \begin{pmatrix} d_{11} & d_{12} & \cdots & d_{1n} \\ d_{21} & d_{22} & \cdots & d_{2n} \\ M & M & & M \\ d_{m1} & d_{m2} & \cdots & d_{mn} \end{pmatrix}$

$$\dot{D} = \begin{pmatrix} \sum_{i=1}^m d_{i1} & 0 & \cdots & 0 \\ 0 & \sum_{i=1}^m d_{i2} & \cdots & 0 \\ 0 & 0 & \cdots & \sum_{i=1}^m d_{in} \end{pmatrix}$$

$$W = [w_1, w_2, \dots, w_m]^T$$

$$U = [u_1, u_2, \dots, u_m]^T$$

$$V = [v_1, v_2, \dots, v_n]^T$$

$$M = [m_1, m_2, \dots, m_n]^T$$

则(2)式、(5)式、(7)式可分别表示成如下的矩阵形式：

$$AX + Y = X \quad (8)$$

$$DX + W = U \quad (9)$$

$$(C + \dot{D})X + V + M = X \quad (10)$$

如果已知本区域的最终产品向量 Y ，那么由(8)式求解得：

$$X = (I - A)^{-1}Y \quad (11)$$

将其代入方程(9)，可求得该区域输入产品向量

$$U = D(I - A)^{-1}Y + W \quad (12)$$

二、区域间模型

如前所述，一个较大的区域地理系统可以分为若干区域，一个区域往往又有若干部门。一个部门的产品除了满足本区域的需要外，还满足其它区域的需要；而这个区域另一些产品也可能是依靠其它区域的供应。区域间的投入产出模型就是研究区域之间的经济联系，发挥各区域优势的一种方法。区域间投入产出模型的结构见表 12-4。

表 12-4 区域间的投入产出表

			中 间 产 品								总 产 品					
			区域1		…	区域m						区 域 1	…	区 域 m	大 区 (m+1)	合 计
			部 门 1	部 门 n		部 门 1	部 门 n									
补 偿 价 值	区 域 1	部门1	$x_{11}^{11} \dots x_{1n}^{11}$	…	$x_{11}^{1m} \dots x_{1n}^{1m}$	y_{11}^{11} … y_{1n}^{11}	…	y_{1m}^{1m} … y_{nm}^{1m}	$y_{1,m+1}^{1,m+1}$ … $y_{n,m+1}^{1,m+1}$	y_{1n}^{10} … y_{nn}^{10}	x_1^1 … x_n^1					
		部门n	$x_{n1}^{11} \dots x_{nn}^{11}$		$x_{n1}^{1m} \dots x_{nn}^{1m}$											
		…														
新 价 值	区 域 m	部门1	$x_{11}^{m1} \dots x_{1n}^{m1}$	…	$x_{11}^{mm} \dots x_{1n}^{mm}$	y_{11}^{m1} … y_{1n}^{m1}	…	y_{1m}^{mm} … y_{nm}^{mm}	$y_{1,m+1}^{m,m+1}$ … $y_{n,m+1}^{m,m+1}$	y_{1n}^{m0} … y_{nn}^{m0}	x_1^m … x_n^m					
		部门n	$x_{n1}^{m1} \dots x_{nn}^{m1}$		$x_{n1}^{mm} \dots x_{nn}^{mm}$											
新 价 创 造 值	劳动报酬 纯收入		$v_1^1 \dots v_n^1$ $m_1^1 \dots m_n^1$	…	$v_1^m \dots v_n^m$ $m_1^m \dots m_n^m$											
	总产品		$x_1^1 \dots x_n^1$		$x_1^m \dots x_n^m$											

在表 12-4 中，假设区域地理系统包含了 m 个区域，每一个区域有 n 个部门，表中记号的上标表示区域，下标表示部门，如：

x_{ij}^{pq} 表示 p 区域供应 q 区域的第 i 部门产品用于第 j 部门生产消耗的数量；

y_i^{pq} 表示 p 区域供应 q 区域的第 i 部门产品用作 q 地区最终产品的数量。

当 $q = m + 1$ 时， $y_i^{p,m+1}$ 就表示 p 区域生产的第 i 部门产品满足整个大区域(即区域地理系统)最终需求的数量。

y_i^{p0} 表示 p 区域生产的第 i 部门产品用作各个区域及大区的最终产品数量之和，即：

$$y_i^{p0} = \sum_{q=1}^{m+1} y_i^{pq}$$

y_i^{qp} 表示 q 区域从各个区域得到的第 i 部门最终产品数量之和，即：

$$y_i^{qp} = \sum_{p=1}^m y_i^{pq}$$

v_j^q ， m_j^q ， x_j^q 分别表示 q 区域 j 部门的劳动报酬，纯收入及总产出。

从表 12-4 的水平方向来看，有如下的平衡关系：

$$\sum_{q=1}^m \sum_{j=1}^n x_{ij}^{pq} + y_i^{p0} = x_i^p \quad \left(\begin{array}{l} p = 1, 2, \dots, m \\ i = 1, 2, \dots, n \end{array} \right) \quad (13)$$

(13) 式反映了各区域各部门产品的生产与分配使用情况。

从表 12-4 的垂直方向看，有如下平衡关系：

$$\sum_{p=1}^m \sum_{i=1}^n x_{ij}^{pq} + v_j^q + m_j^q = x_j^q \quad \left(\begin{array}{l} q = 1, 2, \dots, m \\ j = 1, 2, \dots, n \end{array} \right) \quad (14)$$

照前面的作法，我们引入分区产品直接消耗系数 a_{ij}^{pq} 的概念，它表示 q 区域生产单位 j 种产品消耗的 p 区域供应的第 i 种产品的数量，即：

$$a_{ij}^{pq} = \frac{x_{ij}^{pq}}{x_j^q} \quad \left(\begin{array}{l} p, q = 1, 2, \dots, m \\ i, j = 1, 2, \dots, n \end{array} \right) \quad (15)$$

将(15)式分别代入(13)式和(14)式得：

$$\sum_{q=1}^m \sum_{j=1}^n a_{ij}^{pq} x_j^q + y_i^{p0} = x_i^p \quad \begin{pmatrix} p = 1, 2, \dots, m \\ i = 1, 2, \dots, n \end{pmatrix} \quad (16)$$

和

$$\sum_{p=1}^m \sum_{i=1}^n a_{ij}^{pq} x_j^q + v_j^q + m_j^q = x_j^q \quad \begin{pmatrix} q = 1, 2, \dots, m \\ j = 1, 2, \dots, n \end{pmatrix} \quad (17)$$

这两式若用矩阵表示就是：

$$\sum_{q=1}^m A^{pq} X^q + Y^{p0} = X^p \quad (p = 1, 2, \dots, m) \quad (18)$$

和

$$B^q X^q + V^q + M^q = X^q \quad (q = 1, 2, \dots, m) \quad (19)$$

其中：

$$A^{pq} = \begin{pmatrix} a_{11}^{pq} & a_{12}^{pq} & \dots & a_{1n}^{pq} \\ a_{21}^{pq} & a_{22}^{pq} & \dots & a_{2n}^{pq} \\ M & M & & M \\ a_{n1}^{pq} & a_{n2}^{pq} & \dots & a_{nn}^{pq} \end{pmatrix} \quad (p, q = 1, 2, \dots, m)$$

称为分区直接消耗系数矩阵，即区域 q 的各部门产品生产过程中对区域 p 各部门产品的直接消耗系数矩阵，而 B^q 为对角矩阵：

$$B^q = \begin{pmatrix} \sum_{p=1}^m \sum_{i=1}^n a_{i1}^{pq} & 0 & \dots & 0 \\ 0 & \sum_{p=1}^m \sum_{i=1}^n a_{i2}^{pq} & \dots & 0 \\ M & M & & M \\ 0 & 0 & \dots & \sum_{p=1}^m \sum_{i=1}^n a_{in}^{pq} \end{pmatrix}$$

而 X^p ， Y^{p0} 分别表示区域 p 各部门总产品列向量和最终产品列向量； V^q 和 M^q 分别表示 q 区域的劳动报酬和纯收入向量。

如果我们再引入分块矩阵：

$$A = \begin{pmatrix} A^{11} & A^{12} & \dots & A^{1m} \\ A^{21} & A^{22} & \dots & A^{2m} \\ M & M & & M \\ A^{m1} & A^{m2} & \dots & A^{mm} \end{pmatrix}$$

$$\bar{B} = \begin{pmatrix} \bar{B} & 0 & \dots & 0 \\ 0 & \bar{B}_2 & \dots & 0 \\ M & M & & M \\ 0 & 0 & \dots & \bar{B}^m \end{pmatrix}$$

和列向量

$$X = \begin{pmatrix} X^1 \\ X^2 \\ M \\ X^m \end{pmatrix}, \quad Y = \begin{pmatrix} Y^{10} \\ Y^{20} \\ M \\ Y^{m0} \end{pmatrix}$$

$$V = \begin{pmatrix} V^1 \\ V^2 \\ M \\ V^m \end{pmatrix}, \quad M = \begin{pmatrix} M^1 \\ M^2 \\ M \\ M^m \end{pmatrix}$$

那么, (18)式和(19)式就可以分别用更简洁的矩阵形式表示如下:

$$AX + Y = X \quad (20)$$

$$\bar{B}X + V + M = X \quad (21)$$

三、区域间的相互作用: 引力模型

如果运用投入产出模型对区域之间产品的流动情况进行分析, 则产生了区域间相互作用的引力模型, 它是由列昂捷夫和斯特鲁脱(Alan Strout)于1961年提出来的。区域间引力模型由以下三个方程组组成:

$$x_i^{oq} = \sum_{j=1}^n a_{ij}^q x_j^{qo} + y_i^{oq} \quad \begin{pmatrix} q = 1, 2, \dots, m \\ i = 1, 2, \dots, n \end{pmatrix} \quad (22)$$

$$x_i^{po} = \sum_{q=1}^m x_i^{pq} \quad \begin{pmatrix} p = 1, 2, \dots, m \\ i = 1, 2, \dots, n \end{pmatrix} \quad (23)$$

$$\text{和} \quad x_i^{oq} = \sum_{p=1}^m x_i^{pq} \quad \begin{pmatrix} q = 1, 2, \dots, m \\ i = 1, 2, \dots, n \end{pmatrix} \quad (24)$$

方程(22)中, x_i^{oq} 表示q区域对第i种产品的总使用量, x_j^{qo} 表示q区域第j种产品的产量, a_{ij}^q 表示q区域生产单位第j种产品对第i种产品的直接消耗系数, y_i^{oq} 表示q区域对第i种产品的最终需求量。该方程的意义是, 第i种产品在q区域的总使用量等于q区域的生产消耗与最终产品之和。

方程(23)表示p区域运给所有区域的第i种产品的数量之和等于p区域的产量。

方程(24)表示所有区域运给p区域的第i种产品数量之和等于q区域的总使用量。

将方程(23)和(24)按区域相加, 得:

$$\sum_{p=1}^m \sum_{q=1}^m x_i^{pq} = \sum_{p=1}^m x_i^{po} = \sum_{q=1}^m x_i^{oq} = x_i^{oo} \quad (i = 1, 2, \dots, n) \quad (25)$$

(25)式说明, 对于第i种产品, 所有区域的供应量之和等于这种产品的总使用量, 也等于这种产品大区的总产量 x_i^{oo} 。

为了研究区域间产品流动情况, 我们作如下假设:

第一, p区域供应q区域的第i种产品的产量成正比;

第二, p区域供应q区域的第i种产品的数量与q区域的第i种产品的使用量成正比;

第三，p 区域供应 q 区域的第 i 种产品的数量与第 i 种产品的大区产量成反比。

这三条假设用数学公式来表达就是：

$$x_i^{pq} = \frac{x_i^{p0} x_i^{0q}}{x_i^{00}} Q_i^{pq} \quad (p, q = 1, 2, \dots, m; i = 1, 2, \dots, n) \quad (26)$$

(26)式称为区域间产品流量引力方程，其中 Q_i^{pq} 为比例常数，它由以下四个辅助参数决定：

$$Q_i^{pq} = (c_i^p + k_i^q) d_i^{pq} \delta_i^{pq} \quad (27)$$

在(27)式中， d_i^{pq} 是第i种产品由区域p到区域q运输费用的倒数，为已知常数； δ_i^{pq} 取1或0，如果区域p向区域q运输i种产品，则取1，否则取0； c_i^p 和 k_i^q 是在报告期统计资料基础上，利用最小二乘法计算出来的，其计算过程如下：

设 $\bar{x}_i^{pq}$ 是报告期第i种产品的实际供应量，而

$$\begin{aligned} \bar{x}_i^{pq} &= \frac{\bar{x}_i^{p0} \bar{x}_i^{0q}}{\bar{x}_i^{00}} Q_i^{pq} = \frac{\bar{x}_i^{p0} \bar{x}_i^{0q}}{\bar{x}_i^{00}} (c_i^p + k_i^q) d_i^{pq} \delta_i^{pq} \\ &= b_i^{pq} (c_i^p + k_i^q) \end{aligned} \quad (28)$$

为第i种产品的估计供应量，(28)式中， $\bar{x}_i^{p0}$ ， $\bar{x}_i^{0q}$ ， $\bar{x}_i^{0q}$ 和 $\bar{x}_i^{00}$ 分别为报告期区域p第i种产品的产量、q区域对第i种产品的总使用量和大区的总产量，式中，

$$b_i^{pq} = \frac{\bar{x}_i^{p0}}{\bar{x}_i^{00}} d_i^{pq} \delta_i^{pq} \quad (29)$$

构造目标函数 Φ ，它具有形式：

$$\begin{aligned} \Phi &= \sum_{p=1}^m \sum_{q=1}^m [\bar{x}_i^{pq} - b_i^{pq} (c_i^p + k_i^q)]^2 \\ &= \sum_{p=1}^m \sum_{q=1}^m [\bar{x}_i^{pq} - b_i^{pq} (c_i^p + k_i^q)]^2 \quad (i = 1, 2, \dots, n) \end{aligned} \quad (30)$$

显然，根据报告期已知的统计资料， Φ 仅为 c_i^p 和 k_i^q 的函数。求(30)式的极小值，根据函数求极值的必要条件，令：

$$\begin{cases} \frac{\partial \Phi}{\partial c_i^p} = 2 \sum_{q=1}^m [\bar{x}_i^{pq} - b_i^{pq} (c_i^p + k_i^q)] (-b_i^{pq}) = 0 \\ \frac{\partial \Phi}{\partial k_i^q} = 2 \sum_{p=1}^m [\bar{x}_i^{pq} - b_i^{pq} (c_i^p + k_i^q)] (-b_i^{pq}) = 0 \end{cases} \quad (31)$$

(l=1,2,⋯,n;q=1,2,⋯,m)

上式对于每一种产品i，共有2m个方程和2m个未知量，由此可以求出 c_i^p (p=1, 2, ⋯, m)和 k_i^q (q=1, 2, ⋯, m)的数值。

列昂捷夫曾利用这个方法计算了美国 1954 年九个地区的钢材流动方程，确定了钢材在各个地区的流动情况。

第三节 资源利用与环境保护的投入产出分析

地理学研究的根本目的，就是揭示地理系统的发展变化规律，为实现资源的合理利用与环境保护提供科学的决策依据。因此，资源利用与环境保护不但是当代应用地理学的重大研究课题，而且也是地理科学研究的根本宗旨。

一、资源合理利用的投入产出分析与线性规划模型

关于资源利用问题的研究，人们往往是在资源调查与评价的基础上，运用优化数学方法(如线性规划，非线性规划，多目标规划等)建立资源合理利用的优化模型，然后，通过对优化模型的求解寻求资源合理利用的最佳决策。表面上看来，这种研究方法是科学的、合理的，是无可指责的。但是，这一研究过程却忽视了资源利用系统各个子系统(即各个部门或产业)之间的相互联系。为了克服上述研究过程的这一缺点，一个很自然的想法就是将资源合理利用的优化建模过程与资源利用系统各个子系统之间的相互联系分析结合起来，即将资源合理利用的优化建模与资源利用系统的投入产出分析结合起来。基于这种思想，以下我们来探讨资源合理利用的数学模型——投入产出分析与线性规划模型。

(一)资源利用系统的投入产出分析

在传统的投入产出表中，没有资源项目。在研究资源利用系统各个部门(或产业)之间的相互联系时，我们加入了资源项目，从而得到如下的投入产出表(见表 12-5)。

表 12-5 资源利用系统的投入产出表

		资源利用部门 (生产部门)						最终产品 (值)	总产品 (值)
		1	2	...	j	...	n		
资源利用部门 (生产部门)	1	x_{11}	x_{12}	...	x_{1j}	...	x_{1n}	y_1	x_1
	2	x_{21}	x_{22}	...	x_{2j}	...	x_{2n}	y_2	x_2
	M							M	M
	i	x_{i1}	x_{i2}	...	x_{ij}	...	x_{in}	y_i	x_i
	M	M	M		M		M	M	M
	n	x_{n1}	x_{n2}	...	x_{nj}	...	x_{nn}	y_n	x_n
资源	1	c_{11}	c_{12}	...	c_{1j}	...	c_{1n}		
	2	c_{21}	c_{22}	...	c_{2j}	...	c_{2n}		
	M	M	M		M		M		
	k	c_{k1}	c_{k2}	...	c_{kj}	...	c_{kn}		
	M	M	M		M		M		
	m	c_{m1}	c_{m2}	...	c_{mj}	...	c_{mn}		

在表 12-5 中， x_i 为资源利用各部门的总产值或总产品； y_i 为资源利用各部门的最终产值或最终产品； x_{ij} 为 j 部门生产中所需要的 i 部门的产品(或转

化为产值)数量； c_{kj} 为 j 部门生产中所需消耗的 k 种资源的数量。

如用矩阵形式表示，则表 12-5 的上半部分可写成

$$AX+Y=X \quad (1)$$

或者 $(I-A)X=Y \quad (2)$

式(1)或(2)即综合平衡方程。其中 A 为直接消耗系数矩阵，其元素 $a_{ij} = \frac{x_{ij}}{x_j}$

($i, j=1, 2, \dots, n$) 为直接消耗系数，其意义为第 j 部门生产单位产品(产值)所需消耗的第 i 部门产品(产值)的数量。

同样，在表12-5的下半部分中，令 $d_{kj} = \frac{c_{kj}}{x_j}$ ($i=1, 2, \dots, m; j=1,$

$2, \dots, n$)，则表中的 $c_{kj}=d_{kj}x_j$ 。 d_{kj} 的意义表示 j 部门生产单位产品(产值)所需消耗的 k 种资源的数量，称之为资源消耗系数。记 b_k 为第 k 种资源的拥有量，引入矩阵记号： $B=[b_1, b_2, \dots, b_m]^T$ 及

$$D = \begin{pmatrix} d_{11} & d_{12} & \cdots & d_{1n} \\ d_{21} & d_{22} & \cdots & d_{2n} \\ \cdots & \cdots & \cdots & \cdots \\ d_{m1} & d_{m2} & \cdots & d_{mn} \end{pmatrix}$$

则表 12-5 的下半部分可以写成矩阵形式：

$$DX \leq B \quad (3)$$

(二)资源合理利用的线性规划模型

作为资源合理利用的优化数学模型——线性规划模型，它必须包含目标函数与约束条件两大部分的内容。

1. 目标函数的确定。资源合理利用的线性规划模型，其目标函数的确定，一般可以从以下几个方面去考虑：

(1)使计划期内资源利用所创造的国民收入达到最大，即取

$$\max Z = \sum_{i=1}^n (x_i - \sum_{j=1}^n a_{ij} x_j) \quad (4)$$

(2)使计划期内资源利用所创造的社会总产品(产值)数量达到最大，即取

$$\max Z = \sum_{i=1}^n x_i \quad (5)$$

(3)使计划期内资源利用所创造的最终产品(产值)数量达到最大，即取

$$\max Z = \sum_{i=1}^n y_i \quad (6)$$

(4)使计划期内资源利用所创造的净产值达到最大，即取：

$$\max Z = \sum_{i=1}^n p_i y_i \quad (7)$$

(7)式中， p_i 表示第 i 个资源利用部门产品的单价。

2. 约束条件。本模型最重要的约束条件有三类，即部门联系约束(亦称综合平衡约束)、资源拥有量约束和非负约束。结合以上关于资源利用的投入产出分析，这三类约束可用矩阵形式表示为：

$$\begin{cases} (I-A)X=Y \\ DX \leq B \\ X \geq 0, Y \geq 0 \end{cases} \quad (8)$$

此外，还可以考虑其它约束条件，譬如资源利用部门的资金约束、劳动力约束等等。例：某地有甲、乙两个资源利用部门(生产部门)，它们利用煤炭(燃料)和矿石(原料)分别生产甲、乙两类产品，经投入产出分析得其投入产出系数如表 12-6 所示。若煤炭拥有量为 360 个单位；原料拥有量为 200 个单位；劳动力拥有量为 300 个单位；甲、乙两类产品的单价分别为 700 万元和 1200 万元。试问如何安排生产计划，才能使净产值达到最大？

表 12-6 直接消耗系数

		资源利用部门(生产部门)	
		部门甲	部门乙
资源利用(生产)部门	部门甲	0.1	0.2
	部门乙	0.2	0.3
	煤炭	9	4
	矿石	4	5
劳动力		3	10

设甲乙两个部门总的生产计划量分别为 x_1 和 x_2 ，由表 12-6 可以列出下列约束条件：

(1)综合平衡约束：

$$\begin{aligned} (1-0.1)x_1-0.2x_2 &= y_1 \\ -0.2x_1+(1-0.3)x_2 &= y_2 \end{aligned}$$

(2)资源拥有量约束：

$$\begin{aligned} 9x_1+4x_2 &\leq 360 \\ 4x_1+5x_2 &\leq 200 \end{aligned}$$

(3)劳动力约束

$$3x_1+10x_2 \leq 300$$

(4)非负约束：

$$x_1, x_2, y_1, y_2 \geq 0$$

目标函数为：

$$\max Z = 700y_1 + 1200y_2$$

上述模型，利用单纯形方法求解可以得到：

$$x_1 = 20 \text{ 个单位}$$

$$x_2 = 24 \text{ 个单位}$$

$$\max Z = 24600 \text{ (万元)}$$

此时，矿石资源被完全利用，而煤炭资源尚节余 84 个单位。甲部门利用煤炭、矿石两种资源向社会提供的最终产品(甲类产品)为 3.2 个单位；乙部门利用煤炭、矿石两种资源向社会提供的最终产品(乙类产品)为 12.8 个单位。

二、环境保护的投入产出分析

本世纪 70 年代以来,随着后工业化社会的发展,环境污染已成为人类面临的三大难题(环境污染、能源短缺、人口膨胀)之一。然而,从本质上讲,环境污染与人类生产活动密切相关,环境中的污染物(包括各种物理污染物和各种化学污染物)大都来自人类生产活动。人类在利用各种资源创造大量物质财富的同时,也排出了大量的污染物,从而造成了环境的污染。因此我们应该将环境污染问题的研究与人类生产活动联系在一起考虑。而投入产出分析则是对人类生产活动的各个部门之间相互联系的分析。所以投入产出分析方法应该成为环境污染和环境保护问题研究的有效方法之一。事实上,在本世纪 70 年代初期,列昂捷夫就曾经运用投入产出模型,对环境保护问题作了研究。现在,我们对列昂捷夫环境保护的投入产出分析模型作一简单介绍。

列昂捷夫的环境保护投入产出模型的基本结构如表 12-7 所示。在 12-7 中,除了通常的几个生产部门外,还增加了 m 个污染部门(污染物质的种类)。表中,各项记号的意义如下:

表 12-7 环境保护的投入产出表

		中 间 产 品								最终产品及最 终产品领域所 产生的污染	总 污 染 物
		生产部门				消除污染部门					
1	2	⋯	n	1	2	⋯	m				
生 产 部 门	1	x ₁₁	x ₁₂	⋯	x _{1n}	E ₁₁	E ₁₂	⋯	E _{1m}	y ₁	
	2	x ₂₁	x ₂₂	⋯	x _{2n}	E ₂₁	E ₂₂	⋯	E _{2m}	y ₂	
	M	M	M	M	M	M	M	M	M	M	
	n	x _{n1}	x _{n2}	⋯	x _{nn}	E _{n1}	E _{n2}	⋯	E _{nm}	y _n	
污 染 种 类	1	P ₁₁	P ₁₂	⋯	P _{1n}	F ₁₁	F ₁₂	⋯	F _{1m}	R ₁	
	2	P ₂₁	P ₂₂	⋯	P _{2n}	F ₂₁	F ₂₂	⋯	F _{2m}	R ₂	
	M	M	M	M	M	M	M	M	M	M	
	m	P _{n1}	P _{n2}	⋯	P _{nn}	F _{n1}	F _{n2}	⋯	F _{nm}	R _m	
固定资产折旧		d ₁	d ₂	⋯	d _n	$\overline{d}_1$	$\overline{d}_2$	⋯	$\overline{d}_n$		
新创造 价值	劳动报酬 社会纯收入	v ₁	v ₂	⋯	v _n	$\overline{v}_1$	$\overline{v}_2$	⋯	$\overline{v}_m$		
		m ₁	m ₂	⋯	m _n	$\overline{m}_1$	$\overline{m}_2$	⋯	$\overline{m}_m$		
总产品及污染物消除总量		x ₁	x ₂	⋯	x _n	S ₁	S ₂	⋯	S _n		

x_i —第 i 部门产品的总产出;

y_i —第 i 部门产品的最终产出;

x_{ij} —第 j 部门生产过程中所消耗的第 i 部门产品的数量;

E_{ij} —第 j 个消除污染部门在消除污染过程中所消耗的第 i 部门产品的数量;

P_{ij} —第 j 部门生产过程中所产生的第 i 种污染物的数量;

F_{ij} —第 j 个消除污染部门本身所产生的第 i 种污染物的数量；

R_i —最终需求领域所产生的第 i 种污染物数量；

Q_i —第 i 种污染物的总量；

S_j —第 j 个消除污染部门消除污染物的总消除量；

d_j —第 j 个生产部门的固定资产折旧；

$\bar{d}_j$ —第 j 个消除污染部门的固定资产折旧；

v_j —第 j 个生产部门的劳动报酬；

$\bar{v}_j$ —第 j 个消除污染部门的劳动报酬；

m_j —第 j 个生产部门所创造的社会纯收入；

$\bar{m}_j$ —第 j 个消除污染部门所创造的社会纯收入。

从表 12-7 的水平方向来看，有两组平衡方程，一组是产品的生产与消耗的平衡方程；另一组是污染物的形成方程。即：

$$\sum_{j=1}^n x_{ij} + \sum_{j=1}^m E_{ij} + y_i = x_i \quad (i = 1, 2, \dots, n) \quad (9)$$

$$\text{及} \quad \sum_{j=1}^n P_{ij} + \sum_{j=1}^m F_{ij} + R_i = Q_i \quad (i = 1, 2, \dots, m) \quad (10)$$

(9)式说明，总产品 x_i 除去最终产品 y_i 以外，其余则用作产品生产的消耗和消除污染部门的消耗；(10)式说明，污染物 Q_i 来自三个方面，即生产领域，最终需求领域，以及消除污染部门本身。

我们仍以 a_{ij} 表示生产部门的直接消耗系数。令：

$$e_{ij} = \frac{E_{ij}}{S_j} \quad \begin{pmatrix} i = 1, 2, \dots, n \\ j = 1, 2, \dots, m \end{pmatrix} \quad (11)$$

$$p_{ij} = \frac{P_{ij}}{x_j} \quad \begin{pmatrix} i = 1, 2, \dots, m \\ j = 1, 2, \dots, n \end{pmatrix} \quad (12)$$

$$\text{及} \quad f_{ij} = \frac{F_{ij}}{S_j} \quad (i, j = 1, 2, \dots, m) \quad (13)$$

则， e_{ij} 表示消除污染部门消除一个单位的第 j 种污染物所消耗的第 i 部门产品的数量，它称为消除污染部门的直接消耗系数； p_{ij} 表示第 j 部门单位产品生产过程中的第 i 种污染物的数量，它称为生产部门污染物的产生系数； f_{ij} 表示第 j 个消除污染部门在消除一个单位污染物中所新生产的第 i 种污染物的数量，它称为污染部门污染物的产生系数。

我们引入以下系数矩阵：

生产部门直接消耗系数矩阵：

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ M & M & M & M \\ a_{n1} & a_{n2} & \cdots & a_{nn} \end{pmatrix}$$

消除污染部门直接消耗系数矩阵：

$$E = \begin{pmatrix} e_{11} & e_{12} & \cdots & e_{1m} \\ e_{21} & e_{22} & \cdots & e_{2m} \\ M & M & M & M \\ e_{n1} & e_{n2} & \cdots & e_{nm} \end{pmatrix}$$

生产部门污染物产生系数矩阵：

$$P = \begin{pmatrix} p_{11} & p_{12} & \cdots & p_{1n} \\ p_{21} & p_{22} & \cdots & p_{2n} \\ M & M & M & M \\ p_{m1} & p_{m2} & \cdots & p_{mn} \end{pmatrix}$$

消除污染部门污染物产生系数矩阵：

$$F = \begin{pmatrix} f_{11} & f_{12} & \cdots & f_{1m} \\ f_{21} & f_{22} & \cdots & f_{2m} \\ M & M & M & M \\ f_{m1} & f_{m2} & \cdots & f_{mm} \end{pmatrix}$$

以及 $X=[x_1, x_2, \cdots, x_n]^T$, $Y=[y_1, y_2, \cdots, y_n]^T$, $S=[s_1, s_2, \cdots, s_m]^T$, $R=[R_1, R_2, \cdots, R_m]^T$, $Q=[Q_1, Q_2, \cdots, Q_m]^T$

则方程组(9)式和(10)式就可以分别表示成矩阵形式：

$$AX+ES+Y=X \quad (14)$$

$$PX+FS+R=Q \quad (15)$$

进一步，我们以 $d_i (0 \leq d_i \leq 1)$ 表示第 i 种污染物的消除比例，则

$$S_i = a_i Q_i (i=1, 2, \cdots, m) \quad (16)$$

如果作对角矩阵：

$$\bar{a} = \begin{pmatrix} a_1 & 0 & \cdots & 0 \\ 0 & a_2 & \cdots & 0 \\ M & M & M & M \\ 0 & 0 & & a_m \end{pmatrix} \quad (17)$$

那么向量 S 和 Q 就有如下关系：

$$S = \bar{a}Q$$

将(17)式分别代入(14)式与(15)式，稍加整理便有：

$$\begin{bmatrix} I-A & -E\bar{a} \\ -P & I-F\bar{a} \end{bmatrix} \begin{bmatrix} X \\ Q \end{bmatrix} = \begin{bmatrix} Y \\ R \end{bmatrix} \quad (18)$$

求解(18)式可以得到：

$$\begin{bmatrix} X \\ Q \end{bmatrix} = \begin{bmatrix} I-A & -E\bar{a} \\ -P & I-F\bar{a} \end{bmatrix}^{-1} \begin{bmatrix} Y \\ R \end{bmatrix} \quad (19)$$

(19)式表明，如果已知最终产品 Y 和最终产品领域所产生的污染量 R ，就可以求出应生产的总产品量 X 和产生的污染物总量 Q 。由于(17)式表示污染物的消除总量，因而残存污染物为：

$$Q_{残} = Q - \bar{a}Q = (I - \bar{a})Q \quad (20)$$

以上内容是从表 12-7 的水平方向上研究得到的结论。如果从垂直方向上研究表 12-7，并以价值单位作为生产部门的计量单位，我们可以研究消除污染的费用及其对产品价格的影响。

1. 生产部门费用构成 对于生产部门，在考虑消除污染费用之前的平衡关系是：

$$\sum_{i=1}^n x_{ij} + d_j + v_j + m_j = x_j \quad (j=1, 2, \dots, n) \quad (21)$$

如果进行消除污染活动，势必要提高产品的价格，我们设 π_i ($i=1, 2, \dots, n$) 表示第 j 部门产品价格的提高率； ϕ_i ($i=1, 2, \dots, m$) 表示消除一个单位的第 i 种污染物的费用。这时，新的平衡关系式为：

$$\begin{aligned} & \sum_{i=1}^n (1 + \pi_i) x_{ij} + \sum_{i=1}^m \phi_i a_i P_{ij} + d_j + v_j + m_j \\ & = (1 + \pi_j) x_j \quad (j=1, 2, \dots, n) \end{aligned} \quad (22)$$

由两组平衡关系式(21)和(22)，我们得到：

$$\sum_{i=1}^n \pi_i x_{ij} + \sum_{i=1}^m \phi_i a_i P_{ij} = \pi_j x_j \quad (j=1, 2, \dots, n)$$

上式两端同除以 x_j 得：

$$\sum_{i=1}^n \pi_i a_{ij} + \sum_{i=1}^m \phi_i a_i P_{ij} = \pi_j \quad (j=1, 2, \dots, n)$$

用矩阵形式表示上式，有：

$$A^T \pi + P^T \hat{a} \Phi = \pi, \quad (23)$$

在(23)式中， $\pi = [\pi_1, \pi_2, \dots, \pi_n]^T$ ， $\Phi = [\phi_1, \phi_2, \dots, \phi_m]^T$ 。

2. 消除污染部门的费用第 j 个消除污染部门的费用总额为 $\phi_j S_j$ ，因此第 j 个消除污染部门的费用的平衡关系为：

$$\begin{aligned} & \sum_{i=1}^n (1 + \pi_i) E_{ij} + \sum_{i=1}^m \phi_i a_i F_{ij} + \bar{d}_j + \bar{v}_j + \bar{m}_j \\ & = \phi_j S_j \quad (j=1, 2, \dots, m) \end{aligned} \quad (24)$$

(24)式两端除以 S_j ，并合：

$$h_j = \sum_{i=1}^n e_{ij} + \frac{\bar{d}_j + \bar{v}_j + \bar{m}_j}{S_j} \quad (j=1, 2, \dots, m)$$

于是有：

$$\sum_{i=1}^n \pi_i e_{ij} + \sum_{i=1}^m \phi_i a_i f_{ij} + h_j = \phi_j \quad (j=1, 2, \dots, m)$$

上式用矩阵形式表示，则有：

$$E^T \pi + F^T \hat{a} \Phi + H = \Phi \quad (25)$$

(25)式中， $H = (h_1, h_2, \dots, h_m)^T$ 。

联合(23)式和(25)式，可以得到：

$$\begin{bmatrix} \pi \\ \Phi \end{bmatrix} = \begin{bmatrix} A^T & P^T \hat{a} \\ E^T & F^T \hat{a} \end{bmatrix} + \begin{bmatrix} 0 \\ H \end{bmatrix} \quad (26)$$

求解(26)式得到：

$$\begin{bmatrix} \pi \\ \Phi \end{bmatrix} = \begin{bmatrix} I - A^T & -P^T \hat{a} \\ -E^T & I - F^T \hat{a} \end{bmatrix} \quad (27)$$

利用(27)式，我们就可以计算出消除一个单位污染物的费用向量 Φ ，以及由于消除污染各生产部门价格的上涨率向量 π 。荷兰曾于1973年用类似的方法计算出消除污染对各部门产品价格的影响(见表12-8)。

表 12-8 消除污染对各部门产品价格的影响

部门 时期	农业	纺织业	煤矿	化石品	煤油	金属制品	建筑业
中期(π)	0.22	1.00	0.10	0.47	0.11	0.11	0.18
长期(π)	1.67	6.25	0.96	2.99	1.65	0.97	1.30

从表12-8可以看出，中期消除污染对各部门产品价格的影响的百分率比长期的小，这是因为中期各种污染物的消除比例较长期低的缘故。

