



格致方法·定量研究系列 吴晓刚 主编

自助法：一种统计推断的非参数估计法

[美] 克里斯托弗·Z.穆尼 (Christopher Z. Mooney) 著
罗伯特·D.杜瓦尔 (Robert D. Duval)
李兰 译 李忠路 校

- ★ 革新研究理念
- ★ 丰富研究工具
- ★ 最权威、最前沿的定量研究方法指南

格致出版社  上海人民出版社

56

自助法：一种统计 推断的非参数估计法

[美] 克里斯托弗·Z.穆尼(Christopher Z.Mooney) 著
罗伯特·D.杜瓦尔(Robert D.Duval)
李兰 译 李忠路 校

SAGE Publications, Inc.

格致出版社 上海人民出版社

图书在版编目(CIP)数据

自助法:一种统计推断的非参数估计法/(美)克
里斯托弗·Z.穆尼,(美)罗伯特·D.杜瓦尔著;李兰译.
—上海:格致出版社:上海人民出版社,2017.1
(格致方法·定量研究系列)
ISBN 978-7-5432-2713-2

I. ①自… II. ①克… ②罗… ③李… III. ①非参数
统计-研究 IV. ①0212.7
中国版本图书馆 CIP 数据核字(2017)第 017219 号

责任编辑 张苗凤

格致方法·定量研究系列
自助法:一种统计推断的非参数估计法

[美]克里斯托弗·Z.穆尼 罗伯特·D.杜瓦尔 著
李兰 译 李忠路 校

出 版 世纪出版股份有限公司 格致出版社 世纪出版集团 上海人民出版社 (200001 上海福建中路 193 号 www.ewen.co)		印 刷 浙江临安曙光印务有限公司	
 编辑部热线 021-63914988 市场部热线 021-63914081 www.hibooks.cn		开 本 920×1168 1/32	
发 行 上海世纪出版股份有限公司发行中心		印 张 4	
		字 数 68,000	
		版 次 2017 年 3 月第 1 版	
		印 次 2017 年 3 月第 1 次印刷	

出版说明

由香港科技大学社会科学部吴晓刚教授主编的“格致方法·定量研究系列”丛书,精选了世界著名的 SAGE 出版社定量社会科学研究丛书,翻译成中文,起初集结成八册,于 2011 年出版。这套丛书自出版以来,受到广大读者特别是年轻一代社会科学工作者的热烈欢迎。为了给广大读者提供更多的方便和选择,该丛书经过修订和校正,于 2012 年以单行本的形式再次出版发行,共 37 本。我们衷心感谢广大读者的支持和建议。

随着与 SAGE 出版社合作的进一步深化,我们又从丛书中精选了三十多个品种,译成中文,以飨读者。丛书新增品种涵盖了更多的定量研究方法。我们希望本丛书单行本的继续出版能为推动国内社会科学定量研究的教学和研究作出一点贡献。

总序

2003年,我赴港工作,在香港科技大学社会科学部教授研究生的两门核心定量方法课程。香港科技大学社会科学部自创建以来,非常重视社会科学研究方法论的训练。我开设的第一门课“社会科学里的统计学”(Statistics for Social Science)为所有研究型硕士生和博士生的必修课,而第二门课“社会科学中的定量分析”为博士生的必修课(事实上,大部分硕士生修完第一门课后都会继续选修第二门课)。我在讲授这两门课的时候,根据社会科学研究生的数理基础比较薄弱的特点,尽量避免复杂的数学公式推导,而用具体的例子,结合语言和图形,帮助学生理解统计的基本概念和模型。课程的重点放在如何应用定量分析模型研究社会实际问题上,即社会研究者主要为定量统计方法的“消费者”而非“生产者”。作为“消费者”,学完这些课程后,我们一方面能够读懂、欣赏和评价别人在同行评议的刊物上发表的定量研究的文章;另一方面,也能在自己的研究中运用这些成熟的方法论技术。

上述两门课的内容,尽管在线性回归模型的内容上有少

量重复,但各有侧重。“社会科学里的统计学”从介绍最基本的社会研究方法论和统计学原理开始,到多元线性回归模型结束,内容涵盖了描述性统计的基本方法、统计推论的原理、假设检验、列联表分析、方差和协方差分析、简单线性回归模型、多元线性回归模型,以及线性回归模型的假设和模型诊断。“社会科学中的定量分析”则介绍在经典线性回归模型的假设不成立的情况下的一些模型和方法,将重点放在因变量为定类数据的分析模型上,包括两分类的 logistic 回归模型、多分类 logistic 回归模型、定序 logistic 回归模型、条件 logistic 回归模型、多维列联表的对数线性和对数乘积模型、有关删节数据的模型、纵贯数据的分析模型,包括追踪研究和事件史的分析方法。这些模型在社会科学研究中有着更加广泛的应用。

修读过这些课程的香港科技大学的研究生,一直鼓励和支持我将两门课的讲稿结集出版,并帮助我将原来的英文课程讲稿译成了中文。但是,由于种种原因,这两本书拖了多年还没有完成。世界著名的出版社 SAGE 的“定量社会科学研究”丛书闻名遐迩,每本书都写得通俗易懂,与我的教学理念是相通的。当格致出版社向我提出从这套丛书中精选一批翻译,以飨中文读者时,我非常支持这个想法,因为这从某种程度上弥补了我的教科书未能出版的遗憾。

翻译是一件吃力不讨好的事。不但要有对中英文两种语言的精准把握能力,还要有对实质内容有较深的理解能力,而这套丛书涵盖的又恰恰是社会科学中技术性非常强的内容,只有语言能力是远远不能胜任的。在短短的一年时间里,我们组织了来自中国内地及香港、台湾地区的二十几位

研究生参与了这项工程,他们当时大部分是香港科技大学的硕士和博士研究生,受过严格的社会科学统计方法的训练,也有来自美国等地对定量研究感兴趣的博士研究生。他们是香港科技大学社会科学部博士研究生蒋勤、李骏、盛智明、叶华、张卓妮、郑冰岛,硕士研究生贺光烨、李兰、林毓玲、肖东亮、辛济云、於嘉、余珊珊,应用社会经济研究中心研究员李俊秀;香港大学教育学院博士研究生洪岩璧;北京大学社会学系博士研究生李丁、赵亮员;中国人民大学人口学系讲师巫锡炜;中国台湾“中央”研究院社会学所助理研究员林宗弘;南京师范大学心理学系副教授陈陈;美国北卡罗来纳大学教堂山分校社会学系博士候选人姜念涛;美国加州大学洛杉矶分校社会学系博士研究生宋曦;哈佛大学社会学系博士研究生郭茂灿和周韵。

参与这项工作的许多译者目前都已经毕业,大多成为中国内地以及香港、台湾等地区高校和研究机构定量社会科学方法教学和研究的骨干。不少译者反映,翻译工作本身也是他们学习相关定量方法的有效途径。鉴于此,当格致出版社和 SAGE 出版社决定在“格致方法·定量研究系列”丛书中推出另外一批新品种时,香港科技大学社会科学部的研究生仍然是主要力量。特别值得一提的是,香港科技大学应用社会经济研究中心与上海大学社会学院自 2012 年夏季开始,在上海(夏季)和广州南沙(冬季)联合举办《应用社会科学研究方法研修班》,至今已经成功举办三届。研修课程设计体现“化整为零、循序渐进、中文教学、学以致用”的方针,吸引了一大批有志于从事定量社会科学研究博士生和青年学者。他们中的不少人也参与了翻译和校对的工作。他们在

繁忙的学习和研究之余，历经近两年的时间，完成了三十多本新书的翻译任务，使得“格致方法·定量研究系列”丛书更加丰富和完善。他们是：东南大学社会学系副教授洪岩璧，香港科技大学社会科学部博士研究生贺光烨、李忠路、王佳、王彦蓉、许多多，硕士研究生范新光、缪佳、武玲蔚、臧晓露、曾东林，原硕士研究生李兰，密歇根大学社会学系博士研究生王骁，纽约大学社会学系博士研究生温芳琪，牛津大学社会学系研究生周穆之，上海大学社会学院博士研究生陈伟等。

陈伟、范新光、贺光烨、洪岩璧、李忠路、缪佳、王佳、武玲蔚、许多多、曾东林、周穆之，以及香港科技大学社会科学部硕士研究生陈佳莹，上海大学社会学院硕士研究生梁海祥还协助主编做了大量的审校工作。格致出版社编辑高璇不遗余力地推动本丛书的继续出版，并且在这个过程中表现出极大的耐心和高度的专业精神。对他们付出的劳动，我在此致以诚挚的谢意。当然，每本书因本身内容和译者的行文风格有所差异，校对未免挂一漏万，术语的标准译法方面还有很大的改进空间。我们欢迎广大读者提出建设性的批评和建议，以便再版时修订。

我们希望本丛书的持续出版，能为进一步提升国内社会科学定量教学和研究水平作出一点贡献。

吴晓刚

于香港九龙清水湾

序

长期以来,由于非参数估计不需要做正态分布的假设,因此非参数统计在社会科学研究中一直备受关注。简·狄更生·吉本斯(Jean Dickinson Gibbons)写的《非参数统计简介》(*Nonparametric Statistics: An Introduction*, 本丛书第 90 册)和《相关关系的非参数测量》(*Nonparametric Measures of Association*, 本丛书第 91 册)介绍了许多单变量和双变量的“分布任意”(distribution-free)的统计量。穆尼和杜瓦尔这两位教授执笔的本专著所介绍的推断方法与经典的参数估计方法不同。自助法利用计算机从原样本中“重新抽取”(resample)大量的新样本,通过这些新样本得到一个统计量抽样分布的估计。(根据作者介绍,我们可以利用蒙特卡洛法从一个样本量为 50 的原始样本中有放回地抽取 1 000 个样本量为 50 的随机样本,计算每一次的 $\hat{\beta}$ 值。这 1 000 个 $\hat{\beta}$ 的频率分布将组成抽样分布的

估计。)然后,我们再利用这个估计的抽样分布(而不是事先假设的分布)来做总体推断,例如推断 β 值是否不为 0。

因此,当统计量的潜在抽样分布不能假设为正态分布,且利用普通最小二乘法(ordinary least squares,简称 OLS)估计回归系数得到的残差有偏时,我们可以利用自助法来估计。当没有可用的分析方法来分析抽样分布(如两个样本中位数之差的估计)时,我们也可利用自助法来估计。在这些情况下,我们不能用传统方法来估计置信区间(和做显著性检验),而可能倾向于利用以下四种自助置信区间法(bootstrap confidence interval methods):正态近似法(normal approximation),百分位法(percentile),偏差校正百分位法(bias-corrected percentile),或百分位 t 法(percentile- t)。虽然每种方法都有各自的优缺点,这在本书中有详细的讨论,但穆尼和杜瓦尔稍稍倾向于百分位 t 法,至少当主要目标是假设检验的精确性时。而且,即使分析人员最终依赖于传统的推断方法,他们也可利用自助法来评估某些模型假设是否不成立。

作者运用许多真实数据来举例说明自助法。这些例子包括美国各州的石油生产、标准都市统计区(SMSA)的人均个人收入、美国人争取民主行动组织(Americans for Democratic Action,简称 ADA)对国会成员的排名,以及立法委员会成员和整个立法机关的偏好的中位值之差。最后,在附录中,作者总结了怎样利用不同的软件包来应用

这个计算机运算密集型方法。利用本书和恰当的计算机支持,分析人员应该能很容易地运用自助法做一些统计推断的探索。

迈克尔·S.刘易斯-贝克

前言

我们在 1991 年 7 月美国北卡罗来纳州达勒姆举行的第八届政治学方法论会议上讨论过本书的各章节。我们非常感谢那次参会人员的有效评论,尤其感谢约翰·费里曼(John Freeman)、菲利普·A.施罗德(Philip A. Schrodtt)、加里·金(Gary King)、道格·里弗斯(Doug Rivers)和梅尔·辛尼克(Mel Hinnich)。我们也非常感谢以下参与项目人员的辛勤工作:乔治·克劳斯(George Krause)、布拉德利·埃弗龙(Bradley Efron)、罗伯特·斯泰恩(Robert Stine)、基思·克雷比尔(Keth Krehbiel)、托马斯·J.迪西西欧(Thomas J. DiCiccio)、威廉·雅各比(William Jacoby)、迈克尔·刘易斯-贝克(Michael Lewis-Beck)和两位匿名评审。

目 录

序	1
前言	1
第 1 章 简介	1
第 1 节 传统参数统计推断	6
第 2 节 自助统计推断	13
第 3 节 自助回归模型	21
第 4 节 理论依据	27
第 5 节 刀切法	32
第 6 节 自助法的蒙特卡洛评估	38
第 2 章 利用自助法进行统计推断	43
第 1 节 偏差估计	45
第 2 节 自助置信区间	50
第 3 章 自助置信区间的应用	63
第 1 节 抽样分布未知的统计量的置信区间	65

第 2 节	当传统分布假设不成立时的推断	80
第 4 章	结论	87
第 1 节	未来的研究工作	89
第 2 节	自助法的局限性	92
第 3 节	结语	95
附录	利用统计软件应用自助法	97
	注释	102
	参考文献	104
	译名对照表	109

第 **1** 章

简 介

定量社会科学研究的最基本任务是，利用从总体中抽取的样本得到的估计值 $\hat{\theta}$ 对总体参数 θ 做一个基于概率的统计推断。自助法是一种做这样推断的计算密集型非参数技术。自助法与传统参数推断方法的区别在于，前者利用大量的重复计算来估计一个统计量的抽样分布形状，而后者是通过很强的分布假设和分析公式来估计。基于此，研究人员可以在找不到可行的分析方法或假设不成立的情况下做相应的统计推断。因此，自助法本身不是一个统计量。更确切地说，它是一种利用统计量对总体参数进行推断的方法。但是，它与 z 检验和 t 检验这样的传统参数法存在本质差别。在过去 70 年里，社会科学领域一直教授的是传统参数估计方法。

自助法依赖于样本和其来源总体的类似程度。这里的中心思想是，我们要得出有关总体参数的结论，有时严格地根据样本比根据对此总体做可能不现实的假设要更好。为了生成一个统计量完整的抽样分布的经验估计，自

自助法需要非常多次有放回地“重抽样”数据。在统计推断中二次抽样并不新颖(例如 Jones, 1956; McCarthy, 1969; Tukey, 1958),自助法新颖有趣的地方在于它通过使用丰富便宜的计算机资源,把这种抽样方法广泛地应用到许多统计量上。

为了理解自助法是什么以及它与传统参数统计推断的差异,我们首先必须弄清抽样分布这个概念。一个统计量 $\hat{\theta}$ 的抽样分布可理解为是根据一个样本计算的所有 $\hat{\theta}$ 可能值的相对频率,这个样本是从一个给定总体中抽取的,其样本量为 n (Mohr, 1990:13—28)。考虑到从总体中抽样是随机的, $\hat{\theta}$ 是一个随机变量,其抽样分布是 θ 的函数。下面我们举例来说明。从上社会学导论课的 300 位学生总体中抽取 10 个学生样本,计算这 10 个学生的 IQ 均值。这个样本均值的抽样分布将由从这个班级获得的每个可能 IQ 均值的概率组成。IQ 均值为 90 或 150 的概率低些,而均值为 120 的概率高些。而且,这个分布可能有些右偏,因为相对于 IQ 非常低的学生来说,被学校录取的 IQ 非常高的学生更多。我们可以从这个班级无限地有放回地抽取 10 个学生样本,每次计算抽取的 10 个学生的 IQ 均值,得到一个 IQ 均值的频率分布,这个频率分布就构成了一个抽样分布。图 1.1 呈现了一个可能的抽样分布。

直观地看,这个抽样分布的形状和位置好像受到这个班级的整体 IQ 均值(θ)、这些 IQ 的离差和样本大小的影