

Partial Identification of Probability Distributions



应用统计·数量经济精品系列

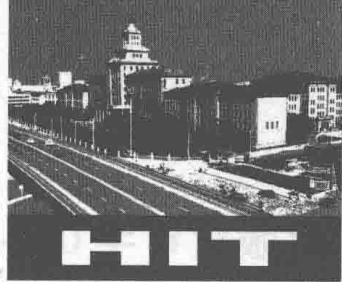
概率分布的部分识别

[美] 查尔斯·曼斯基 (Charles F. Manski) 著 王忠玉 译

非
外
借



哈尔滨工业大学出版社
HARBIN INSTITUTE OF TECHNOLOGY PRESS



应用统计·数量经济精品系列

Partial Identification of Probability Distributions

概率分布的部分识别

● [美] 查尔斯·曼斯基 (Charles F. Manski)

王忠玉 译

常州大学图书馆
藏书章



哈尔滨工业大学出版社
HARBIN INSTITUTE OF TECHNOLOGY PRESS

黑版贸审字 08 - 2017 - 098 号

内容简介

全书采用一种统一方式加以讨论,即首先对生成可用数据的抽样过程进行设定,并考察仅利用实证证据时,探讨了解认识总体参数的情况,然后研究倘若在施加各种各样的假设条件下,这些参数的集值识别域会如何缩小。所用的推断方法是传统的且完全非参数的方法。

本书适合于统计学、应用数学、数量经济学、经济学等专业的研究生和教师,以及相关专业对部分识别方法和应用感兴趣的研究人员、教师等参考使用。

图书在版编目(CIP)数据

概率分布的部分识别/(美)查尔斯·曼斯基(Charles F. Manski)著;王忠玉译.——哈尔滨:哈尔滨工业大学出版社,2018.7

ISBN 978 - 7 - 5603 - 7310 - 2

I. ①概… II. ①查…②王… III. ①概率分布—研究 IV. ①0211.1

中国版本图书馆 CIP 数据核字(2018)第 067342 号

Translation from English language edition:

Partial Identification of Probability Distributions by Charles F. Manski

Copyright © 2003 Springer New York

Springer New York is a part of Springer Science + Business Media

All Rights Reserved

策划编辑 刘培杰 张永芹

责任编辑 张永芹 杜莹雪

封面设计 孙茵艾

出版发行 哈尔滨工业大学出版社

社 址 哈尔滨市南岗区复华四道街 10 号 邮编 150006

传 真 0451 - 86414749

网 址 <http://hitpress.hit.edu.cn>

印 刷 哈尔滨市工大节能印刷厂

开 本 787mm × 1092mm 1/16 印张 13.5 字数 240 千字

版 次 2018 年 7 月第 1 版 2018 年 7 月第 1 次印刷

书 号 ISBN 978 - 7 - 5603 - 7310 - 2

定 价 68.00 元

(如因印装质量问题影响阅读,我社负责调换)

本书献给亚瑟·戈德伯格 (Arthur Goldberger),
他鼓励我不断飞翔。

I am pleased that Harbin Institute of Technology Press is publishing the authorized Chinese translation of my book Partial Identification of Probability Distributions. I am grateful to Professor Wang Zhongyu for his effort to perform the translation. I hope that the Chinese translation will enable Chinese – speaking students and researchers to become familiar with partial identification analysis and to apply it usefully to empirical research.

Charles F. Manski

2018.3.20

我非常高兴哈尔滨工业大学出版社出版了我的书《概率分布的部分识别》授权的中译本,感谢王忠玉教授为翻译这本书所付出的辛苦努力。我希望中译本能够促进中国学生和研究人员熟悉部分识别分析,并将其作为实证研究的有力工具。

查尔斯 F. 曼斯基

2018 年 3 月 20 日

◎
推
荐
序

自从拉格纳·弗里希(Ragnar Frisch)于20世纪30年代创立了计量经济学以来,经过计量经济学家们的不断努力,已经建立起新的方法论和计量经济理论,实现了从政治经济学向理性经济学质的转变,这不仅是在经济学研究的方法论上取得的重大突破,而且作为实证经济学的主流方法的计量经济学理论体系同样取得了巨大的成就。

特别是在最近30年,计量经济学领域出现一些新的方法或探索领域,比如协整、部分识别、单位根、格兰杰因果性、金融计量经济方法、面板数据的计量经济理论及方法、结构宏观计量经济方法、微观计量经济方法,等等。

对于识别理论,首先应该明确它的地位以及它与统计应用的关系。利用计量经济理论及方法进行实证研究问题时,一般步骤是:建立模型、识别模型、参数估计、显著性检验及应用。模型识别是连接模型设定与参数估计之间的桥梁,建立了联立方程计量经济学模型以后,要进行参数估计,必须先判别方程是否可识别。只有可识别的方程才能得出其结构式参数,可以认为,此时识别作为得以严谨的概念化的理论,已成为相对独立的部分,被从模型设定与估计中分离出来。从本质上看,由于识别涉及结构式方程参数的定值问题,是先于估计的逻辑问题,所以识别不是统计推断的问题,而是产生于模型建立与对变量概率分布解释之间的先验问题,从这一角度出发,有必要对识别进行单独研究。

实际上,对于计量经济模型来说,研究者做出的假设越强,所能得到的信息就会越多,因此,要获得较强的结构必然是以更强的假设为代价的。

部分识别分析方法致力于更灵活的识别概念,并为实证研究者提供可利用各种不同假设来控制 and 利用的感兴趣参数的信息范围。从某种意义上说,部分识别实证研究方法是当今理论经济学研究中前沿疆域之一。部分识别分析方法,将从各种不同实证模型所推断出的结论和以透明方式所做出的一系列假设联系起来,利用这种部分识别方法,实证研究者可以检查他们所做假设的信息内容,同时探索对所做推断的影响。

《概率分布的部分识别》是美国西北大学经济学教授曼斯基对“部分识别”深入探索系列成果的集成,是一本系统阐述计量经济学中部分识别分析方法的经典著作。本书中译本的出版是国内大学生和实证研究者、相关人员了解和学习“部分识别”这一新兴实证分析方法的一扇门,是进一步学习和掌握部分识别分析方法的基石。

科学的发展离不开现实的问题,科学创新来自继承和发展。作为新兴的实证研究分析的部分识别分析方法,势必会被越来越多的实证研究者和相关人员所掌握和运用,运用部分识别分析方法会使计量经济模型的假设更为灵活、应用更贴近现实、更有应用价值。

赵振全
吉林大学商学院教授

仅有样本数据不足以推断出有关总体的结论。研究人员在进行统计推断时,总是需要对总体和抽样过程做出一些假设。统计理论已经揭示了许多关于假设的强度如何影响点估计的精确度的问题,但是关于假设如何影响总体参数的可识别,这方面的内容还有许多事情要探索。事实上,将识别看成一个二元事件,也就是参数要么是可识别的,要么是不可识别的,这样做司空见惯,并且将点识别视为统计推断的先决条件。但是,利用数据和部分识别总体参数的假设,对于富有成效的推断来说具有很大的发展空间。

这本书解释了为什么要研究部分识别,并说明怎样利用部分识别进行推断。本书以严谨而全面的方式阐述了查尔斯·曼斯基对概率分布部分识别进行研究的主要内容。其中一个重点是,关于带有缺失结果或协变量数据的预测问题;另一个重点是关于有限混合的分解,这可应用于污染抽样和生态推断的分析;第三个重点是聚焦处理响应分析。

无论所研究的对象是什么,这本书的阐述都遵循一种共同的方式:作者首先对生成可用数据的抽样过程进行设定,同时考察仅利用实证证据时,探讨了解认识总体参数的情况,然后研究如果施加各种各样的假设,那么这些参数的(通常)集值识别域会如何缩小。贯穿全书的推断方法是有意采用传统的且完全非参数的方法。

在没有施加不可预测的假设条件下,传统的非参数分析会促使研究者从可利用的数据中学习。这种方法能在众多研究人员之间建立一个共识领域,这些研究者对关于什么样假设是适宜的的可能持有不同的信念。

◎ 前言

在早期,我对部分识别的研究工作是孤独的,亚瑟·戈德伯格 (Arthur Goldberger) 看到了它的潜力,并支持我继续这样做. 当我激动地向他展示一个新研究发现时,他给予我精神上特别的鼓舞,他说“现在你正在飞翔”. 多年以来,作为同事、评论家和朋友,我没有太多方式来表达这对我来说是多么的重要. 将这本书献给他就是我的一个鼓励.

随着人们对部分识别的兴趣,同时对部分识别的研究做出的贡献,我的研究工作少了些孤独感. 本书的 10 章内容中有 4 章是基于和我合作的共同作者而完成的,我很喜欢和合作者进行富有成效的合作. 第 3 章和第 4 章是基于和乔尔·霍罗威茨 (Joel Horowitz) 合作的几篇已经发表的论文而写成的;第 5 章是基于和菲利普·克罗斯 (Philip Cross) 合作论文而完成的;第 9 章是基于和约翰·佩珀 (John Pepper) 共同撰写的论文而成的. 针对这些特定的项目,我从与乔尔、菲利普和约翰一起研究工作中获益良多,还有对共同关注主题的讨论,亦是如此.

在 2002 年夏天,杰夫·多米尼茨 (Jeff Dominitz)、弗朗西斯卡·莫利纳里 (Francesca Molinari)、约翰·佩珀以及丹尼尔·沙尔夫斯泰因 (Daniel Scharfstein) 对我完成本书的草稿提出了周全

细致而又富有建设性的意见. 幸运的是,这四个人全都对这本书感兴趣,并且热心地帮助我改进本书内容的研究范围及阐述. 在 2002 年春季和秋季,我非常感谢在西北大学的研读博士学位的研究生参与我讲授的计量经济学课程. 在 2002 年春季课程中,我尝试了本书各种不同形式的各章早期版本,并在 2002 年秋季课程中完成了整本书的草稿. 我也感谢约尔格·斯托伊(Joerg Stoye),他仔细阅读并修改了本书手稿.

美国国家科学基金会对我的研究项目提供了持续的资助. 我这本书的前期准备研究工作还部分地得到了 SES-0001436 项目的资助.

伊利诺伊州,芝加哥

查尔斯 F. 曼斯基

2003 年 1 月

◎ 目 录

引言 部分识别与可信推断 //1

第1章 结果数据缺失 //6

1.1 问题解析 //6

1.2 均值 //8

1.3 遵从随机占优的参数 //10

1.4 多重抽样过程组合 //13

1.5 结果的区间测量 //17

补充1A 就业概率 //18

补充1B 盲人摸象 //20

注释 //22

第2章 工具变量 //25

2.1 分布假设与可信推断 //25

2.2 利用工具变量的某些假设 //26

2.3 结果数据随机缺失 //27

2.4 统计独立性 //28

2.5 均值独立性与均值单调性 //30

2.6 其他利用工具变量的某些假设 //33

补充2A 带有无回答权重的估计 //34

注释 //35

第3章 带有数据缺失的条件预测 //37

3.1 以协变量为条件的结果预测 //37

3.2 结果数据缺失 //38

3.3 联合结果数据缺失与协变量 //39

3.4 协变量缺失 //43

3.5 一般数据缺失模式 //45

3.6 条件分布的联合推断 //48

补充3A 失业率 //50

补充3B 带有数据缺失的参数预测 //51

注释 //53

第4章 污染结果 //55

4.1 数据误差的混合模型 //55

4.2 结果数据分布 //57

4.3 事件概率 //58

4.4 遵从随机占优的参数 //60

补充4A 通过补算的污染 //63

补充4B 识别和稳健推断 //65

注释 //66

第5章 回归、短回归及长回归 //68

5.1 生态回归 //68

5.2 问题解析 //69

5.3 长均值回归 //70

5.4 工具变量 //76

补充5A 结构预测 //78

注释 //79

第6章 基于响应的抽样 //81

6.1 逆向回归 //81

6.2 结果或协变量的辅助数据 //83

6.3 罕见病假设 //83

6.4 相对危险度和归因危险度的界 //85

6.5 从响应层抽样 //87

补充6A 吸烟与心脏病 //90

注释 //91

第7章 处理响应分析 //92

7.1 问题解析 //92

7.2 异质性总体中的处理选择 //95

7.3 选择问题与处理选择 //97

7.4 工具变量 //100

补充7A 识别与模糊不清 //102

补充7B 判刑和案犯 //104

补充7C 结果数据缺失与协变量 //105

补充7D 研究总体与处理总体 //108

注释 //109

第8章 单调处理响应 //111

8.1 形式约束 //111

8.2 单调性 //114

8.3 半单调性 //117

8.4 凹单调性 //121

补充8A 向下倾斜的需求 //125

补充8B 计量经济响应模型 //127

注释 //128

第9章 单调工具变量 //130

9.1 等式与不等式 //130

9.2 均值单调性 //132

9.3 均值单调性与均值处理响应 //134

9.4 单调工具变量专题变化形式 //137

补充9A 受教育回报 //137

注释 //141

第10章 混合问题 //142

10.1 组内处理变异 //142

10.2 已知处理份额 //145

10.3 仅利用实验的推断 //147

补充10A 没有协变量数据条件下的实验 //149

注释 //153

参考文献 //154

译后记 //166

编辑手记 //177

部分识别与可信推断

引

言

统计推断是指运用样本数据经由某种推理而得到关注总体(Population of Interest)的某些结论。可是,仅有数据并不能满足推理要求。人们在进行推断时,总是要对总体和抽样过程做出一些假设。统计理论通过说明怎样将数据与假设相结合来得出结论,以此阐明统计推断的逻辑。

实证研究者不仅应该研究推断的逻辑,而且还应该探讨其推断的可信性。尽管可信性是一件主观之事,但是我认为,存在一个广泛的意见一致原理,我们称之为可信性递减定律(the law of Decreasing Credibility)。

可信性递减定律:推断结论的可信性会随着所做假设效力的增强而递减。

这个原理蕴含着实证研究者面临着他们决定支持什么样假设的困境:倘若做出的假设越强,则得出的结论就会越具有说服效力,但其可信程度便会越低。统计理论虽不能解决此困境,但可以澄清其本质特性。

数据和假设的结合可用于对关注总体的参数进行点识别,而另一种方式则是将参数置于集值识别域,区分上述这两种方式是十分有益的。对于参数的一致点估计来说,点识别是十分基本的必要条件。对业已取得的点识别假设进一步加以强化,这可以提高参数的可达精确度。统计理论对诸如此类之事阐述得非常多。通过费希尔(Fisher)信息矩阵,经典的局部渐近有效性理论刻画了可达精确度是怎样随着所假定的总体内容的增多而增强的。非参数回归分析说明,估计的可达收敛速率

是怎样随着所假定的回归形状内容增多而增快的。这些业已取得的成果以及那些进展为实证研究者面对各种不同点的估计方法,为对可信性与精确度进行权衡提供了重要的指南。

对于不是点识别的参数(参看引言后面的历史评注)来说,统计理论对此极少谈及。将识别考虑成二值事件(参数要么是可识别的,要么是不可识别的),并且将点识别看成是有意义推断的先决条件。然而,有关利用数据与可部分识别总体参数的假设来进行推断方面取得了大量的丰硕成果。本书将解释为什么要进行这样的推断,并说明怎样进行推断。

本书起源与组织安排 这本书根植于我在 20 世纪 80 年代后期,对带有数据缺失的非参数分析所做的探索。运用回归估计的实证研究者通常会假定:缺失性在如下意义上是随机的,即结果的可观测性与结果之值是统计独立的。可是,这个点识别假设与其他的点识别假设,由于是不可能真实的而时常受到批评,所以我打算研究,如果关于缺失性过程什么都不知道,或者如果假设比所施加的作为广泛可信的那种假设充分地弱,则什么样的带有结果的部分可观测性的随机抽样可以揭示有关均值与分位数的回归呢? 研究结果表明,准确界(Sharp bounds)形式会随着关注回归与所做假设的变化而变化。这些界是可用非参数回归分析的标准方法很容易地估计出来。

对带有结果数据缺失的回归进行研究,激发了对更一般不完全数据问题的探索。一些样本实现值可能有不可观测的结果,有些样本实现值可能有不可观测的协变量,而另一些样本实现值可能全部缺失。有时候,可以利用结果的区间测量或者协变量,而不是点测量。带有结果与协变量不完全观测的随机抽样一般会产生回归的部分识别。其挑战是,当人们做出另一些假设时,如何刻画和估计由不完全数据过程所产生的识别域。

对带有结果数据缺失的回归进行研究,也很自然地导致了对处理响应的推断方面的考察。对处理响应进行分析必须研究下面的基本问题:反事实结果是不可观测的;因此,我对带有结果数据缺失的回归的部分识别研究方面取得了许多成果,都可以直接得以应用。然而,处理响应分析远远超越一般数据缺失问题。一个原因是,当结合了合适的假设之后,结果的实现观测值可以提供反事实事件的信息。另一个原因是,处理选择的实际问题激发了对处理响应做出

更多的探索,因此,要确定什么样的总体参数是人们感兴趣的。进而,我发现将处理响应推断作为一个专题确实有其自身价值,这样做是富有成效的。

另外一个研究主题是,对有限概率混合形式的成分进行推断。对有限概率混合形式进行分解的数学问题,从本质上看,会在许多截然不同的背景中产生,具体来说包括污染抽样、生态推断以及带有协变量数据缺失。有限概率混合形式的研究成果,不仅在所有这些专题方面有应用,甚至也在其他别的方面得到应用。

本书系统地阐述了我对概率分布的部分识别进行研究所取得的主要成果。第1章至第3章构成对涉及带有结果数据缺失或协变量数据情况进行预测的单元。第4章和第5章构成对有限混合形式执行分解的单元。第6章则是自成体系的,组成了基于响应的抽样分析单元。第7章至第10章构成处理响应分析单元。

不论所研究的专题内容如何特殊,阐述讲解都遵循一种共同的途径。首先,我设定抽样过程生成了可利用的数据,并探寻在没有限制总体分布的假设条件下,什么会成为推断总体参数呢?然后,就要探寻如果施加某些假设,这些参数的集值识别域会怎样(典型地)收缩呢?当然,存在数量极其众多的可能具有各自特定目的的假设。在本书里,我主要研究统计独立性与单调性假设。

贯穿于本书的推断方法是有意采用传统的标准且完全非参数的方法。研究抽样过程的传统方法可以部分识别总体参数,这种传统方法将可利用数据与充分强的假设相结合,从而获得点识别。这种假设往往没有很好的动机,而实证研究者时常反复考虑假设的有效性。传统的标准非参数分析会使研究者在施加了站不住脚的假设条件下,从可利用数据中知晓信息。非参数方法能在那些可能各自持有什么样假设是适宜的信念的研究者中间建立起共识,也会使可利用数据的局限性清晰可见。当可信的识别域被证明是令人不满意的时候,研究者应该大胆地面对如下事实:可利用数据确实不支持他们所希望达到的尽可能精炼的推断。

从总体上看,本书的分析是建立在最基本的初等概率论基础上。数量众多的有关识别,可以从全概率定律和贝叶斯定理的有见地的应用中知道,这点将变得十分明显。为了在不牺牲严谨条件下保持阐述简单起见,我们自始至终地