

统计预测

主
编
译

中国发明创造者基金会 中国预测研究会

统计预测

中国预测研究会



60193

C93
86-20

统 计 预 测

〔英〕Warren Gilhrist著

1976年第一版，1978年4月再版

王丽媿 译

中国发明创造者基金会
中国预测研究会

一九八五年二月

译者的话

由于译者对统计预测的业务不够熟悉，对翻译工作经验不足，译文中难免有缺点错误，欢迎读者批评指正。

原版本中每章的后面都附有参考书目，因这些文献在国内不易找到，故略去。前言在翻译时也略去。

由于条件所限，本书中的插图不象原版那样精致，所用符号也与原文不完全一致，希读者见谅。

2014/28

目 录

第一篇 序论	(1)
第1章 预测简介.....	(1)
1.1 人人预测	(1)
1.2 预测方法	(2)
1.3 预测的基础	(2)
第2章 模型	(6)
2.1 科学的预测	(6)
2.2 基本的模型	(9)
2.3 全局和局部有效模型	(11)
2.4 预测模型和公式	(12)
第3章 预测标准	(12)
3.1 标准的必要性	(12)
3.2 根据数据测试预测公式的标准的分类	(14)
3.3 统计特征	(16)
3.4 瞬态特征	(19)
3.5 稳态特征	(19)
3.6 根据模型测试预测公式的标准	(19)
3.7 模拟研究	(21)
第二篇 根据几类基本模型的预测	(23)
第4章 常数平均数模型	(23)
4.1 全局常数平均数模型	(23)
4.2 局部常数平均数模型	(28)
第5章 线趋势模型	(34)
5.1 全局线趋势模型	(34)
5.2 局部线趋势模型	(38)
5.3 多项式模型	(48)

第6章 回归模型	(48)
6.1 引言	(48)
6.2 配合回归模型	(50)
6.3 选择变量	(53)
6.4 计量经济模型	(56)
6.5 作为计算回归的自变量的时间	(57)
第7章 随机模型	(58)
7.1 引言	(58)
7.2 预测和条件期望	(64)
7.3 移动平均过程	(65)
7.4 自回归过程	(67)
7.5 自回归——移动平均模型	(69)
7.6 包含随机和确定成份的模型	(71)
第8章 季节模型	(74)
8.1 基本概念	(74)
8.2 季节指标法	(79)
8.3 傅里叶 (Fourier) 法	(89)
8.4 随机方法	(95)
8.5 方法的比较	(96)
第9章 增长曲线	(96)
9.1 增长曲线模型	(96)
9.2 配合增长曲线	(98)
9.3 预测增长曲线	(102)
第10章 概率模型	(103)
10.1 引言	(103)
10.2 预测概率	(104)
10.3 情况模型	(106)
10.4 预测标准	(109)
10.5 概率预测和决策	(111)
第11章 多变量模型	(112)
11.1 引言	(112)
11.2 指数平滑的推广	(112)
11.3 多变量随机模型	(113)

第12章 预测方法和模型 (114)

- 12.1 预测方法..... (114)
- 12.2 建立和识别模型..... (114)
- 12.3 模型检验..... (116)
- 12.4 判断的应用..... (120)

第三篇 预测过程 (122)

第13章 数据..... (122)

- 13.1 引言..... (122)
- 13.2 数据的来源..... (122)
- 13.3 数据的质量..... (123)
- 13.4 数据的校正..... (124)

第14章 自适应的方法和其它的推广 (127)

- 14.1 引言..... (127)
- 14.2 自适应的方法..... (127)
- 14.3 递归公式的推广..... (129)
- 14.4 误差校正方法的推广—Kalman模型..... (131)
- 14.5 线性预测公式..... (133)
- 14.6 应用混合的方法..... (135)

第15章 分析和方法的比较 (137)

- 15.1 引言..... (137)
- 15.2 预测值分析..... (138)
- 15.3 选择预测参数..... (149)
- 15.4 方法的比较..... (150)
- 15.5 预报区间..... (151)
- 15.6 灵敏度分析..... (152)

第16章 预测控制 (154)

- 16.1 预测中的质量控制..... (154)
- 16.2 偏离的追踪信号..... (155)
- 16.3 累积和方法..... (157)
- 16.4 自相关的追踪信号..... (159)

第17章 两步预测 (160)

- 17.1 引言..... (160)

17.2 预测误差的应用	(161)
17.3 外部预测资料的应用	(165)
17.4 组合预测值	(166)
17.5 代价和其它标准的应用	(168)
第18章 实践中的问题	(171)
18.1 引言	(171)
18.2 实际的预测标准——一个库存控制的实例	(171)
18.3 作为一个系统的组成部份的预测值——一个库存拥有量的实例	(174)
18.4 预测值和水晶球——一个调查表的实例	(175)
18.5 作为系统的组成部份的预测值——自动生效和自动失效的预测值	(180)
18.6 实际的预测——结束语	(181)
附录A 随机过程的某些专门术语和定义	(182)
附录B 条件期望给定最小均方误差预测值的证明	(185)
附录C 两路指数平滑	(186)
附录D 打折扣的最小二乘的技术特征	(187)

符 号

- x_t 在时间 t 的一个观测值
 $\hat{x}_t$ X_t 的预测值或估计值
 $\tilde{x}_{t+h}$ 在时间 t 做出的 x_{t+h} 的预测值, 也就是带有超前时间 h
 $\tilde{x}_t$ 超前一步预测值 $\tilde{x}_{t+1}$ 的缩写
 e_t 预测误差 $x_t - \hat{x}_t$ 。通常是 $x_t - \tilde{x}_{t-1}$
 e_{t+h} 超前时间 h 的预测误差 $x_{t+h} - \tilde{x}_{t+h}$
 μ_t 在时间 t 的基础平均数, 也就是 $E(x_t)$
 B_t 在时间 t 的基本趋势

第一篇 序论

第一章 预测简介

1.1 人人预测

每当我们约会时我们就作出一个关于我们将来守约的能力的预测。在“会见十分钟”的水平上我们甚至不去想预测的问题，但是对于比较长期的约会，我们力图思考可能妨碍约会的以预测形式出现的任何明显的事物。我们对于预测值思考得越多，就越是希望它将是正确的，因为我们不大可能忽略某些先前的约会或者这种状态的某些其它方面。在这个例子中，我们甚至可能试图利用当违约时我们感到意外的总数来估价预测值的置信度。这样一些不自觉的或半自觉的预测行动是我们日常生活的组成部份；它们也是那些涉及工业或商业问题的人们的生活的组成部份。一个好的经理与其说是能把过去错误的影响减到最低程度的人，宁可说是能够成功地驾驭未来的人。例如，考虑下列问题：

下月销路如何？

这个月应当生产多少？

应当保有多少库存？

应当购进多少材料？

应当在何时购进材料？

销售的目标是什么？

劳动力是否应当增加？

应当承担的代价是什么？

需要多少经营人员？

利润是多少？

为了得到这些问题的最好答卷，我们需要有可能了解将来的情况；为了能够给出这些问题的最好的实际答案，我们必需能够预测未来。我们的预测值可能错误，但是应当得到它们。过去，上述类型问题的大多数答案建立在不自觉或半自觉的预测值的基础上。在这些预测值中常常假定未来会和最近的过去恰巧一样。对预测的许多方面的日益增长的兴趣是建立在下述信心的基础上，即自觉地和仔细地思考肯定有助于改进我们的预测技能以及我们得到上述类型的问题的答案的能力。在某些情况下，一种完全科学的预测研究可能并不比老的“凭臆测”的方法得出更好的预测值，即使如此，这种研究可能得到对于他们的状态较好的理解从而改善他们对状态的控制。

统计预测方法有许多种，但它们中没有一种能够提供每个人的预测问题的全部答案。在本书中选择出来论述的技术是由于它们已经得到广泛运用。无论如何，本书的重点不是特别的技术而是一般的方法和对于预测的态度。因此希望读者最好能够运用统计预测方法于他需要的问题上。每个月都公布预测的新技术和运用，因此，如果简单地把本书当作预测技术的纲要，它将很快变成过时的东西。但是，如果读者把它当作一本关于统计预测的一般方法和态度的教

科书，那么他将能够把这个迅速扩展的领域中的新发展包括到他的知识的基础结构中去。

1.2 预测方法

在本节中我们将定义三种一般的预测方法，它们包含很不相同的途径。这并不意味着仅有这些方法是有用的类型或任何特定的方法应当自动地归入这些类中的一类，我们介绍这些方法是为了提供一般的背景情况。

1.2.1 直观法

这是预测的传统方法。它们基本上是在个体对于状态的感觉或专家舆论的基础上。人们常常由于情况的复杂和缺乏可靠的及有关的资料而被迫采用此方法。设法改善这种方法的一条明显的途径就是用小组或者较大数量的人群来代替单个的专家。例如，人们可能把销售人员的或者地区营业主任委员会的概括意见用于预测销售量。另一种方法，潜在顾客意图的抽样调查可能获得更可靠的资料。

1.2.2 因果关系法

这些方法建立在试图根据事物的原因的知识来预测结果的基础上。在许多预测状态中，原因是经济方面的，因此这个领域的大量工作也是经济性质的。人们研究原因和结果之间的经济关系，由于结果发生在原因之后，关于原因的直接知识导致可以预测结果。但是，象常发生的那样，当原因和结果之间的时间落后很小时我们可能还必须预测原因，因为在结果发生以后的某个时间之前，关于原因的资料可能是得不到的资料。由于原因能够比结果以更可靠的形式来预测，我们选择用于预测的原因。如果关系的结构已知，可以设计出很完善的方法。但是，关系的性质通常是含混不清的，因而我们被迫运用关系的一些完全经验式的形式。第六章提供相应的一些预测方法的论述。有相当多的书论述经济预测；参考书目列举了一些例子。

1.2.3 外推法

这些方法都是建立在把过去的有关数据显示的特征外推到将来的基础上。这些方法通常是数学性质或统计性质的，且为本书的大部份提供基础。在一些关于经济预测的书中这些方法有时被称为朴素的方法。但是，为了通晓任何状态，我们必须从朴素的开始，考察最简单的方法和手段。我们彻底理解了这些方法，然后我们可以前进到较高级的领域。

1.3 预测的基础

企图对预测未来进行科学研究可能似乎是无希望的行为，因为每一个问题明显地是独一无二的。但是，当我们更仔细地考察预测的一些特定的例子，就发现一些共同的原则。让我们考察一些例子。

例1.1

多数政治上的决策者都关心预测没有采取新行动时未来的情况。然后他们按照他们的看法来考察社会结构和进行活动以便采用根据他们的政治准则将会改进社会的政策。为了能够决定打算进行什么活动，他们必须有能力和识别将和他们的政策相对应的社会结构。执行失败的政策不是未知数。执行失败的政策可能是由于对结构的错误认识，或者可能是由于这个结构缺乏稳定性。为了正确地作出这样的预测值和确定政策，决策者必需能够识别合适的结构及对它适当时期内的稳定性有若干信心。

例1.2

H.T.Davis (1941)指出：如果太阳的质量和行星相对来说不是这样大，统计方法将必须

用来调查研究引力的物理定律和预言行星运动的未来。在现状下，太阳对行星的引力拉力是如此强有力，以致行星的轨道几乎就象其它行星不存在时一样。从这个轨道的许多观测值中能够推导出太阳和行星之间的关系的结构，因而推断引力的平方反比定律是配合于大多数被观测现象的模型。这种情况的巨大稳定性使得产生行星未来位置的很准确的预测值成为可能。假如太阳的质量较小，它的作用已不能压倒其它行星的影响，那么就会产生更复杂得多的轨道。这种复杂情况的后果将在于：虽然平方反比定律仍然还是基本的定律，但从观测数据中要找出它来将会很艰难。更进一步，预测行星的未来位置将会困难得多，因为已不会出现我们具有的轨道的简单性和稳定性。

例1.3

作为第3个例子我们将考察可能和一家厂商每周收到的定货相对应的一套数据所取的某些形式。表1.1给出适当的数据。假定表中(a)行给出数据的形式。在这里，定货恰巧来自一个需要为常数的顾客。数据显示一种很简单的结构：

表1.1

接受的定货(千单位)

	周											
	1	2	3	4	5	6	7	8	9	10	11	12
(a)	6	6	6	6	6	6	6	6	6	6	6	6
(b)	6	6	6	6	9	6	6					
(c)	6	6	6	6	9	6	6	6	6	9	6	6
(d)	7	6	6	4	7	6	8	6	5	6	7	5
(e)	8	8	9	8	12	12	15	14	14	16	18	17

需要不变和完全稳定的状态。从而我们在预测第13周有6000单位的定货时将感到完全有信心。但是，假定数据采取表1.1中(b)行所示形式。在这里，在第5周发生的9,000单位的需要导致对其余各周显示的结构稳定性有些怀疑，因此，我们还可以预测下周即第8周为6000单位，但是不再有这样的信心。假定我们等待几周以观察有什么情况发生和调查9,000单位的需要的原因。我们可能发现(b)行的数据是两个顾客定货的结果，一个就是上述的顾客，另一个顾客规则地每隔5周定货3000单位。现在切合实际的预测值将为：第13周和第14周是6000单位，第15周为9000单位。于是我们看到：通过考虑更精致的结构形式，有时可能妥善处理结构中出现的不稳定性。当获得如数据(c)中那样的较多的资料时，数据(b)的明显的不稳定性变为稳定结构。

实际上，如象(a)、(b)、(c)中那样的数据不常发生。我们最常见的数据是如表1.1的(d)行中所示的那些数据。这些数据不表示(a)的最简单的数学结构，但它们仍然表示如从图1.1所见的结构。这时给出一个数据的简单柱形图，它显示不同定货量发生的次数，揭示了一个很清晰的结构，但是，在这种情况下，它是一个统计结构，而不是一个数学结构。很明显，平均定货量接近6000单位，而且，根据数据的这个小样本的迹象，看来定货量是相当均匀地分布在这个数的周围，离平均数越远，出现的次数越低。

定货量 (千单位)	频率
4	1
5	2
6	5
7	3
8	1

期	定资本(千单位)
1	8
2	8
3	9
4	8
5	12
6	12
7	15
8	14
9	14
10	18
11	19
12	18

的数值在数学中是不稳定的！假设稳定的意思是：这种随机行为的结构如图1.1对12周的样本的说明那样。在预测第13周的定货时，有理由预测为6000单位，但是，现在我们既不象对数据（a）那样完全肯定我们的预测值，也不象对数据（b）那样带有模糊的怀疑。图1.1的资料，也就是统计结构，能够形成预测值的可能误差的很清晰的概念。如果我们用数据来提供误差概率的估计值而且仅考虑以千为单位，当我们的预测值为6000时预测是正确的概率约为 $\frac{5}{12}$ ，有1000单位的误差的概率约为 $\frac{5}{12}$ ，有2000单位的误差的概率约为 $\frac{1}{6}$ 。

例1.4

两 年 的 数 据

月												
年	1	2	3	4	5	6	7	8	9	10	11	12
1	33	34	30	36	26	22	14	25	19	20	38	37
2	42	49	42	33	37	34	23	40	23	31	49	40

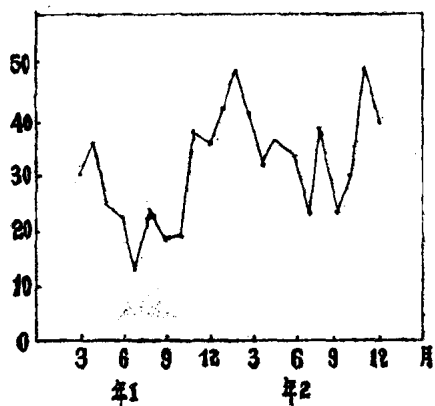


图1.3 表1.2的数据的散点图

表1.2给出两年的数据及图1.3给出这套数据的散点。一眼看去，这些数据好象是无定形的。对这种状态的研究，提出一种设想：即它可能由一个向上的趋势，某种季节变动及一个机会变化（或如通常所称的随机变化）所组成。表1.3显示把数据分成这三种成份的分类，同时带有随机变化的次数表，它给出如图1.1的同类的资料。从而数据具有同时由数学和统计成分组成的结构。从这个例子可以清楚地看到为一套数据寻找基础结构可能是一个困难的任务。按照所显示的形式获得的表1.2的数据无疑是不明显的。

(a) 表1.3

表1.2的数据分类

月	趋 势		季 节	随机变化		总 数	
	年 1	年 2	加量	年 1	年 2	年 1	年 2
1	21	33	10	2	- 1	33	42
2	22	34	12	0	3	34	49
3	23	35	10	- 3	- 3	30	42
4	24	36	8	-1*(4)	-11	31*(36)	33
5	25	37	0	1	0	26	37
6	26	38	- 6	2	2	22	34
7	27	39	-16	3	0	14	23
8	28	40	-10	7	10	25	40
9	29	41	- 9	- 1	- 9	19	23
10	30	42	- 7	- 3	- 4	20	31
11	31	43	2	5	- 1	38	44
12	32	44	6	- 1	-10	37	40

(b)

随 机 变 化 次 数 表

数值	小于-7	-7-5	-4-2	-1-1	2-4	5-7	8-10
次数	3	0	4	9*(8)	5*(6)	2	1

〔注 与表1.2的数字对照原文年1、4月的总数错了，因此有关数字也错了，改正的数字在括号内〕

那也是很清楚的，即给出足够的努力和机敏，对任一套数据，人们能够创造出某种这样的结构。不过如果为了预测的目的，我们打算假定结构的稳定性，很清楚，所用的结构应当

和状态中的某种实际的结构相符。用于预测的数学和统计方程式，对于将用它们来预测的状态来说，应当是有意义的。因此，用于预测的结构的研究，不仅是数学或统计的任务，而且需要广阔的关于要调查的具体的预测问题的性质和背景情况的知识。

上述的例子中我们已经看到，预测的基本要求是存在一个稳定的结构。这种结构可能是纯粹的数学形式，或者纯粹的统计形式，或者，更常见的是二者的混合形式。迄今为止，我们都以单纯描述的方式来对待这些结构，今后，我们在描述结构时应当准确和严格。预测状态的数学和统计模型提供了这样一种严格的描述。对这样的一些模型的研究的介绍将是下一章的主要内容。

第二章 模型

2.1 科学的预测

在第一章中已经看到预测是工商业活动固有的组成部份，然而直到最近还主要是依靠直观的方法来作出预测。当前经营方法革命的一个必要的组成部分就是力图用科学的方法去作出经营决策。现在让我们研究用科学方法去预测问题的若干步骤。

2.1.1 第一步——收集数据

在我们能够采用任何科学的预测方法于一种状态之前，我们自然应当详尽地考察状态，尤其要搜集和该状态有关的尽可能多的数据。这样的数据主要有两类：第一类，它是公司内部的数据，例如过去的销售量、产量、库存等的记录；第二类，它是公司外部的数据，例如政府和贸易协会统计经常提供背景情况，和背景情况相对照可以解释内部资料的意义。有时，这些外部资料在产生预测值方面可能有直接价值。这方面的一个例子将是整个工业的一些预测值为预测那种工业内部的一个公司的销售额提供了基础资料。

2.1.2 第二步——缩减数据

这里目的是从第一步所获得的尽可能广泛的大量资料中选出最有关的资料，因而它将减到最低限度。例如，在销售额预测情况中公司过去的销售额的纪录大概将提供最低限度的资料。为了决定哪些项目应当包括在这种基本的最低限度的资料中，常常运用三条标准。最低限度的资料应当符合下述要求：

适当：资料是可能得到的最直接适当的资料吗？

可靠：数据是怎样获得的？它的可靠程度怎样？

时新：这是得到的最新资料吗？新资料将能迅速获得吗？

当这种最低限度的资料已经过透彻的研究而且理解它的特性后我们可以开始返回来运用其它资料。例如，在销售额预测中，我们可能利用经济资料或者由销售人员作出的关于状态的估计。但是，暂时我们仅考虑运用这种最低限度的资料。虽然我们已有运用于选择我们的资料的若干标准，但是我们通常不可能以完全客观的方式做到这点。必须在广阔和深刻理解实际情况的基础上做出选择。对这类资料的较近的科学研究将指出我们的选择是好的或者是坏的。通过从资料得到好的预测值的成就，将显示出资料选择的好坏。很遗憾，当我们可能发现我们的资料不是最好的时候已经太晚了，因此，在这个过程中，我们的直观选择是必需的和重要的步骤。

2.1.3 第三步——建立模型

为了识别它的结构和产生这种结构的准确特性曲线，现在我们来研究我们的最低限度的资料。在大多数场合，这种特性曲线将由结构的数学描述或统计描述组成；这种描述就是我们通常所谓的建立模型。为了使这个概念更具体，让我们再考察表1.1给出的数据。在前章中我们曾看见用纯粹语言来描述数据。现在我们将用数学术语来描述它们。让 x_t 代表第 t 周的定货， $t=1,2,3\cdots$ 对于数据（a）， x_t 总是6,000单位，所以我们写：

$$x_t = 6,000 \quad (t = 1, 2, 3, \dots)$$

对于数据（c）， x_t 在第5周，第10周，第15周……变为9,000，所以我们写：

$$x_t = 6,000 \quad (t = 1, 2, 3, 4, 5, \Delta 6, 7, 8, 9, 11)$$

（注： Δ 处原文如此，5应去掉）

$$x_t = 9,000 \quad (t = 5, 10, 15, 20, \dots)$$

对于数据（d），我们发现，定货与6,000的差别，不是一个确定的数值；而是一个由

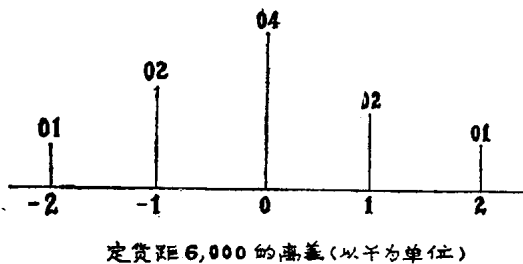


图2.1 概率分布

统计结构决定的数值。让我们用符号 ε_t 代表在第 t 周发生的距6,000的机会或随机离差，这个量 ε_t 可能以一定的概率取值……，

-2, -1, 0, 1, 2, …… (千单位)。图1.1（注：原文如此，应为图2.1）提供的数据给出抽取一个这些随机离差的样本的结果。假定图2.1提供的数值给出 ε_t 的实际的基础概率。因而状态的这种结构可以被想象为一顶

含有按图2.1给定的比例的标明-2,000, -1,000, 0, 1,000, 2,000的标签的宽帽子，这些比例事实上是选择相应标签的张数的概率。在 t 周末，随机选出一张标签，它的数值我们叫做 ε_t 。这个数值被加到6,000上以得到 x_t ，因而这个模型是：

$$x_t = 6,000 + \varepsilon_t$$

这里，用更准确的术语来说， ε_t 是具有图2.1给出的概率分布的随机变量。对于表1.1中（e）行的数据，一个趋势被加到结构上。这个趋势，如图1.2所示，在数学上被表达为：

在第 t 周的趋势值 = 1,000 t

因此定货通过方程式被作成模型：

$$x_t = 6,000 + 1,000t + \varepsilon_t$$

在我们的例子中，我们拿数据作为始点然后建立模型。帽中的标签的实例使得我们可以假定：把一个模型作为始点并从它建立具有这种结构的数据总是可能的。模型的纯粹数学部分几乎不产生问题，但由 ε_t 产生的数值可以变得很复杂。一顶字面上的帽子将处理图2.1的分布。对于更复杂的分布，我们需要更精巧的设备，最有用的是一台计算机。为什么有能力制造具有任何给定模型的特性的数据是有用的，有两个理由。第一，它给我们从给定结构的数据去获得预期特性的经验；第二，它使我们有可能考察预测的不同方法的性质。从不同种类模型的特性获得的经验，对于决定何种模型可考虑用于我们自己所有的一套具体的数据，有很大帮助。

如果我们用图2.1的概率分布去产生一组 ε_t 以便用于趋势模型，将得到和表1.1中我们原有的数据不相等的一套数据，但是它将有相同的结构。因而我们看到：模型既不等于一套数据，也完全不能说明为什么数据表现为这种或那种琐细的形式。模型仅给出状态的数

学的或统计的公式，而一个公式能够产生类似于观测数据的特性的一些数。

2.1.4 第四步——把模型外推

对我们的目的来说，模型和数据之间的重要差别在于：资料的最新部分结束了数据的收集，而模型将产生我们想要代入的任何 t 的数值。因而在表达式

$$x_t = 6,000 + 1,000t + \varepsilon_t$$

中我们可以代入任何 t 值而并不恰好是原有数据显示的 $t=1$ 到 12 。对 t 的不同值求数学部分 $6,000 + 1,000t$ 是一个简单的代换问题。随机部份 ε_t 的求值，将依赖于我们正设法要做的事。如果目的是产生一套具有给定特性的数据，应当从适当的概率分布中找出 ε_t 的数值。如果目的是预测，一种合理的方法就是用 ε_t 的推理的平均值来代替它的未来值，这个平均值通常是零。因此 x_{14} 的一个预测值，用符号 $\hat{x}_{14}$ 表示，将是：

$$\hat{x}_{14} = 6,000 + 1,000 \times 14 + 0 = 2,0000$$

因为在图2.1的分布中，这个 ε_t 的推理平均数或期望是零。我们假定这个结构是稳定的，因而把我们的模型，我们的结构外推到未来。

在图2.2中阐明了我们已论述过的基本步骤。有许多和建立模型有关的问题，特别是关于选择模型形式和用于模型的数值问题，在以前的一般论述中已被忽略。例如，我们如何知道上述的一个趋势模型应当和一个6,000的开始值及每周增加1,000单位的常数比率一起使用。在实际状态中我们由完全无知开始，也就是用模型

$$x_t = ?$$

判断这个问号应当是什么被称为识别问题，模型的选择问题。在上述数据中我们已经假定解决了这个问题，而且我们有一个形式

$$x_t = \alpha + \beta t + \varepsilon_t$$

的模型，这里的 α 和 β 是常数， β 是直线的斜率及 α 是与 x 轴的截距。不过我们还不知道常数 α 和斜率 β 所取的数值。这里面对的问题被称为估计或模型配合问题，用于模型的数值选择问题。当我们已经解决上述模型的这个问题时，然后我们可以取

$$x_t = 6,000 + 1,000t + \varepsilon_t$$

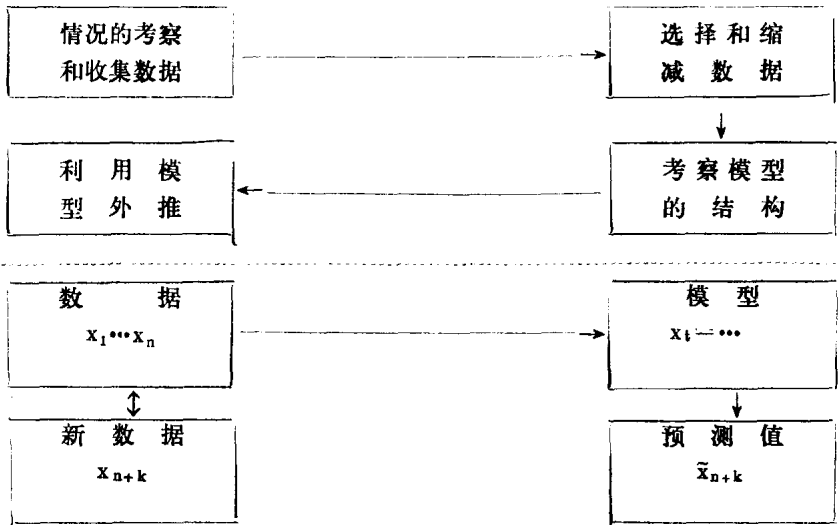


图2.2 统计预测的四个基本步骤

注意：在我们开始预测之前，不是一次就能解决所有的问题。我们希望，在任何阶段，我们的模型将是对迄今为止的观测数据给出一个好的描述的模型。然而随着时间的前进，将得到新数据而且结构可能改变，或者增加的有用资料的量可能揭露那些以前不明显的特性。这种可能性有两个后果。第一个是：判断和估计的过程需要以某种形式继续进行，只要预测在继续。第二个是：即使至今为止我们有一个对于数据来说很好的模型，它已经给出极好的预测值，这仍旧不能保证这个模型在将来会继续给出好的预测值。考虑的时期越长，结构保持稳定的机会越少。因此，我们外推我们的模型到将来的期间越长，对我们的预测值的可靠性的怀疑越大。我们已经正确地识别一个稳定结构这个基本假设的可能失败，意味着预测者可能未曾真正投出他的水晶球。（译者注：这里水晶球意指可以清楚地看到未来的实际情况的透明球）

2.2 基本的模型

在我们以前的论述中可以看出模型的选择是预测的一个基本步骤。许多从事实际工作的人对这点的一种自然的反应就是去探索那个模型，因为他们的状态或许是独一无二的，他们将必须创造自己的模型以和他们的状态相配合。因为可能感到这是他们力所不及的，他们可能丧失统计预测的希望。幸而实际上已经发现各式各样的数据能够根据一些相对来说数量较少的基本模型来适当地做成模型。表2.1给出若干基本模型的一张一览表。本书的第二篇研究这些模型中的每一种模型在预测中的应用。通常这些模型提供实际预测所需要的一切。但是，有时会需要把它们结合、修正和以某种形式改制以符合正在处理的问题的要求。做这件事的一些方法和若干例子在本书的第三篇给出。

让我们用表2.1的模型对若干术语作一个初步介绍。表中如（a）、（b）、（c）之类的模型是纯粹的数学形式。例如，从纯粹线性趋势中，如果在时间0, 1, 2, 3, 4, 我们得到数据2, 4, 6, 8, 10, 那么这些观测值中的任意两个数，都可以用来证明 $\alpha = 2$ 和 $\beta = 2$ ；因此下一个观测值的准确的预测值将是12 ($= 2 + 2 \times 5$)。这样的模型我们将称为确定性模型。如象我们在例1.4中曾发现的那样，这类模型难得描述归因于机会（或机遇）的现实世界，而那种变化是由大量数据来显示的。例如表2.1中（g）表示的随机数列常用作这方面的模型。因此，象我们已看到的那样，我们可能有由下式

$$x_t = \text{线性趋势} + \text{随机数列}$$

或者，由数学式

$$x_t = \alpha + \beta t + \varepsilon_t$$

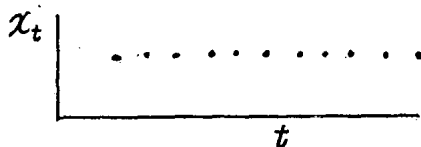
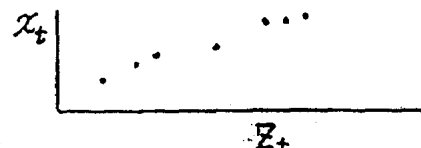
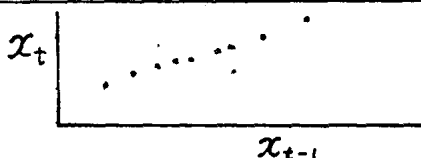
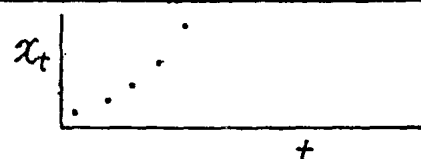
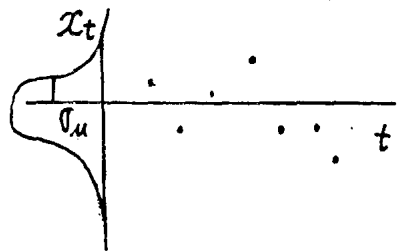
做成模型的数据。这里 ε_t 用于代表散布在零的上下的独立随机变量数列。因为这个确定的部份是这类模型的不变的基础，我们通常还是把这类模型当作本质上是确定性的模型。很遗憾，从这样一类模型得出的数据看来如同图1.2所示，因此，恰好拿出两个观测值以便找出 α 和 β 将是不可能的。我们将必须发展这些方法以便找出 α 和 β 的数值，估计 α 和 β 的数值，只要我们的模型显示和数据配合得好，那么就可以有理由外推以得到一个预测值。如象 α 和 β 这样的参数估计问题，从而模型和数据的配合问题将是以后各章大量论述的基本内容。

现在让我们来考察表2.1中（d）所描述的模型。按现实情况来说，被这种方程式描述的情况是确定性的。如果我们把它变得更实际而且写为

$$x_t = \alpha + \beta x_{t-1} + \varepsilon_t$$

以便考虑到现实世界的机会变化，我们获得一种很不同的状态。如果 x_t 受以前的 x 的影

表 2.1 几种常用模型

名 称	散 点	形 式
(a) 常数平 均 数		对于全部 t $x_t = u$
(b) 线 性 趋 势		对于全部 t $x_t = \alpha + \beta t$
(c) 线 性 回 归		$x_t = \alpha + \beta z_t$ z_t 是另一变量
(d) 自 回 归		$x_t = \alpha + \beta x_{t-1}$ x_{t-1} 是 x 以前的数值
(e) 周期 性 的		x_t 每隔 T 个观测值 重复图形对于全部 t , $x_{t+nT} = x_t \quad (n=1, 2, 3)$
(f) 指数 增 长		对于全部 t $x_t = \alpha e^{\beta t}$
(g) 随机 正 态 数 列		x_t 是具有平均数 μ 和 方差 σ^2 的正态分布对于 全部 t , x_t 的值彼此独立