

第八章 时间序列

世界是在发展变化的，我们所关心的一些指标(经济的等等)随着时间变化而变化，这种按发生时间先后变化的指标称为时间序列

(*Time Series*)。现在要问这些变化是否有规律（如是否有周期，是否有趋势）？能否进行预测？预测精度如何？研究随时间发生变化的指标（可以是一维，也可以是多个指标，即多维的）的变化规律就称为时间序列分析，

例1 观察Jasbeer, Rajan&Co. 七年半按季度的销售情况。

年.季度	时间	(y)	(T)	(y-T)
1.1	1	166	253.65	-87.65
1.2	2	52	263.67	-211.67
1.3	3	140	273.68	-133.68
1.4	4	733	283.70	449.30
2.1	5	224	293.71	-69.71
2.2	6	114	303.73	-189.73
2.3	7	181	313.74	-132.74
2.4	8	753	323.76	429.24
3.1	9	269	333.77	-64.77
3.2	10	214	343.79	-129.79
3.3	11	210	353.80	-143.80
3.4	12	860	363.82	496.18
4.1	13	345	373.83	-28.83
4.2	14	203	383.84	-180.84
4.3	15	233	393.86	-160.86

4.4	16	922	403.87	518.13
5.1	17	324	413.89	-89.89
5.2	18	224	423.90	-199.90
5.3	19	284	433.92	-149.92
5.4	20	822	443.93	378.07
6.1	21	352	453.95	-101.95
6.2	22	280	463.96	-183.96
6.3	23	295	473.98	-178.98
6.4	24	930	483.99	446.01
7.1	25	345	494.01	-149.01
7.2	26	320	504.02	-184.02
7.3	27	390	514.04	-124.04
7.4	28	978	524.05	453.95
8.1	29	483	534.06	-51.06
8.2	30	320	544.08	-224.08
合计	465	11966		

我们不禁要问为什么会产生这样的时间序列？
哪些因素的共同作用产生了上述数据。我们分二部分介绍所用的模型，第一部分模型是描述性的。第二部分是用数理统计知识建模的，(仅作简单介绍)

第一部分模型是描述性的

§ 1 因子与模型

时间序列反映某一指标随时间的推移而呈现的变动。影响这种变动的因素很多，不仅有自然的，还有经济的和社会的。在诸多的因素中，有些对指标的波动起着长期的或决定性的作用，有些则对指标的波动起着短期的和不确定性的作用，从而使指标的变化呈现波动性，周期性和不规则性。因此时间序列的值是诸多因素共同作用的结果。作为基本的分析，最简单的模型是把时间序列的值看成为四种基本因素---趋势(trend)，季节因子(seasonal factor)，

循环波动（**cyclical factor**）和随机误差（**random factor**）---的和或积。下面我们来一一讨论。

（一）因子

如上所说，我们主要考虑的如下的四个因子（factor），即

1. 趋势:表达指标数据运动的方向，
2. 季节因子：在一个完整周期内发生的正常波动（若数据按季，则4个数据发生一次波动），
3. 循环因子：长期（一般为若干年）的正常波动，
4. 随机因子：许多其他因素共同作用对时间序列产生的影响。

(二) 模型

常用的模型有

1. 加法模型(the additive model)

以 y 表示实际数据，则模型为

$$y = T + S + C + R \quad (1)$$

对许多模型，一般没有足够的数据来识别循环周期，故(1)常常简化为 $y = T + S + R \quad (2)$

这儿的基本假定为 $R=0$ 或误差序列的平均值 $\bar{R}=0$

当趋势变化在每年或每周有基本相同的大小时，加法模型较为适用。

2. 乘法模型(multiplicative model) 此模型把时间序列的值表达为四个因子的乘积：

$$y = T \times S \times C \times R \quad (2)$$

在这个模型中 y 和 T 为实际量值，而 S, C, R 为比值，且假定 $\bar{R} = 1$ 。有时把这个模型表达为

$$y = T \times S \times C + R \quad (3)$$

这儿 y, T, R 为实际量，而 S, C 为比值，且 R 的平均值 $\bar{R}$ 假定为0。

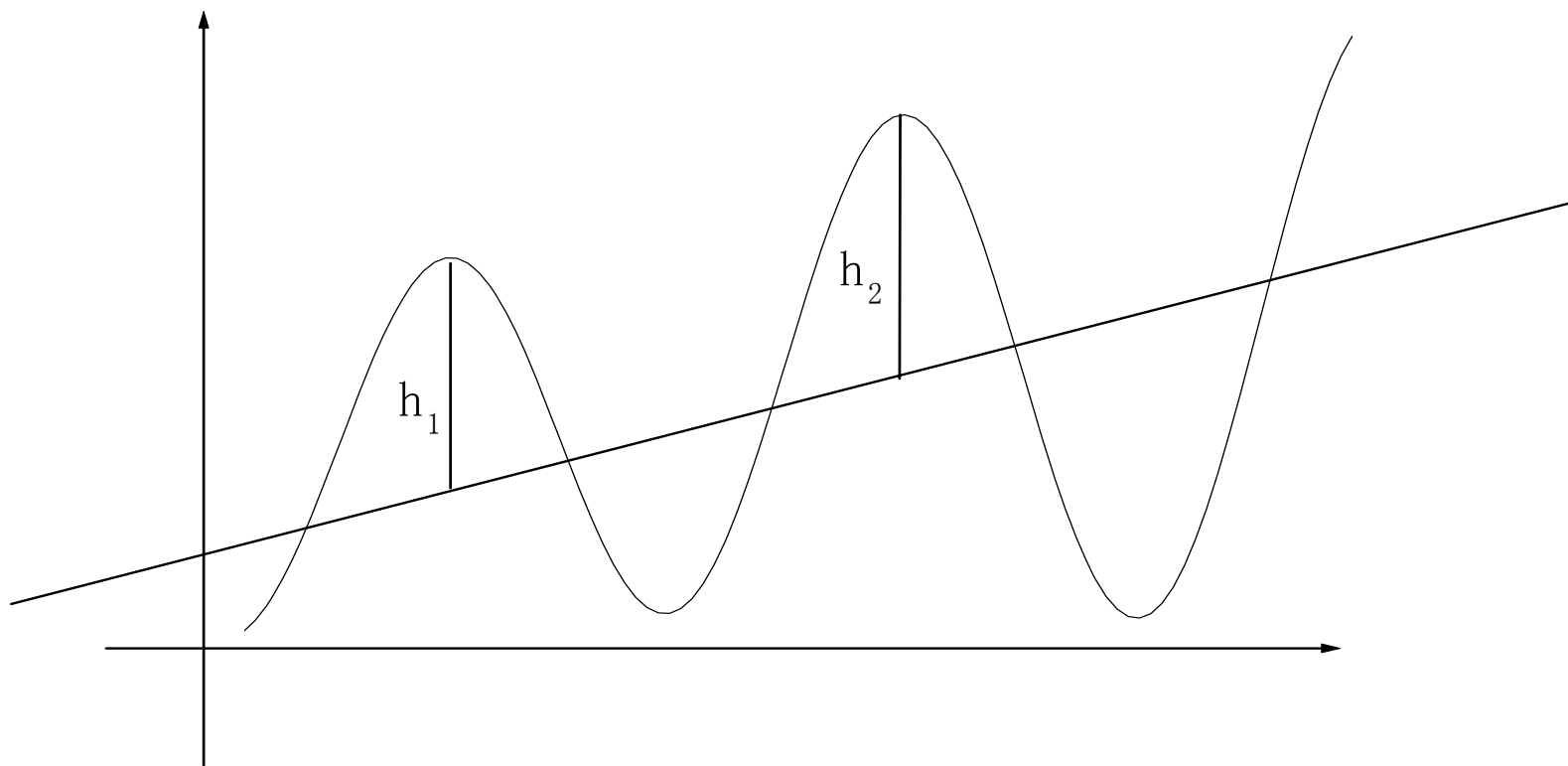
如同加法模型，通常由于数据少而不能识别循环因子，故和通常简化为

$$y = T \times S \times R \quad (2')$$

和

$$y = T \times S + R \quad (3')$$

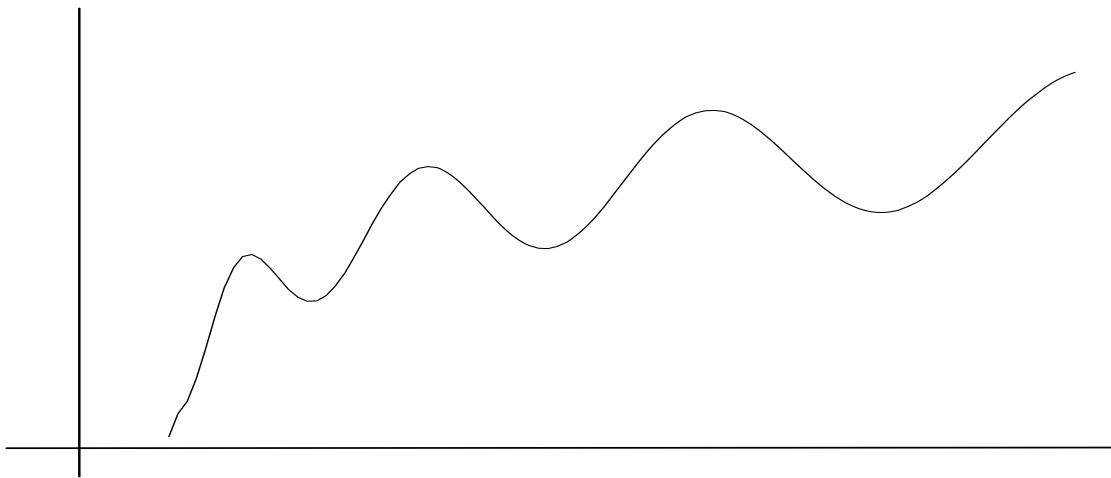
(2)和(3)适用于在每年（每周）相同时期内关于趋势的变化有与趋势相同的比例或百分比。



当趋势是一条水平线时，趋势差 h_1 不变（即图中
 $h_1 = h_2$ ）

§ 2 趋势

为观察时间序列有无趋势（即随着时间的延伸时间序列的运动方向），首先以时间为横坐标把时间序列的值标在图上（我们称为散点图）。然后用眼大体观察一下时间序列的走向。如果走向是向上的，则称时间序列有向上的趋势；如果走向是向下的，则称时间序列有向下的趋势。当然也有曲折前进的。



我们用最小二乘法讨论直线趋势，用滑动平均法讨论曲线趋势。

1. 直线趋势

设时间序列数据为 (i, y_i) , $i=1, \dots, n$, 其中 i 表示第 i 季度（年，日等），如果时间序列有直线趋势，设直线方程为

$$y = a + bi$$

我们的任务是用数据求出系数 a 和 b 。原则是最小二乘估计(**LS**估计)。由前一章最小二乘估计原理线性回归得和的估计为

$$\hat{b} = \frac{n \sum i y_i - \sum i \sum y_i}{n \sum i^2 - (\sum i)^2} \quad \text{这儿}$$

$$\hat{a} = \bar{y} - \hat{b} \bar{i} \quad \bar{i} = \frac{1}{n} \sum_{i=1}^n i = \frac{n+1}{2}, \quad \bar{y} = \frac{1}{n} \sum_{i=1}^n y_i$$

例2. 数据见例1, 先作散点图(见图3)。由带统计功能的计算器或在计算机用SPSS软件包可算得

$$\hat{b} = 10.015, \quad a = 243.639$$

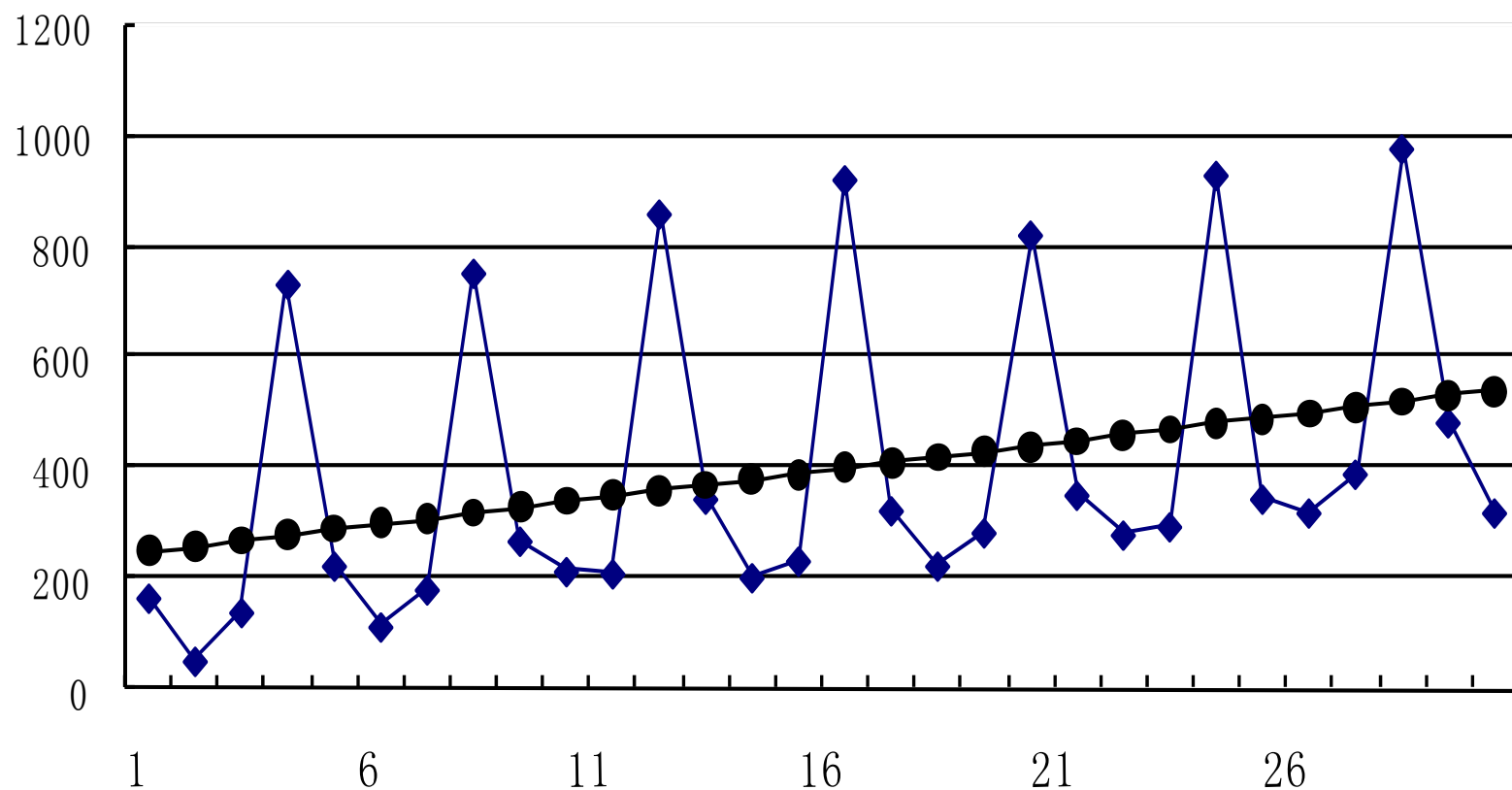


图 3

趋势直线（见图3）为

$$y = 243.639 + 10.015 i$$

由(4)可找出趋势在不同时间点的值。如 $i = 1$ 时趋势值为

$$y = 243.639 + 10.015 \times 1 = 253.654$$

注意：上式得出的是趋势值，并不是实际值。正如前面提到的，实际值的估计要涉及到用什么模型。（软件包**Excel**也可做上述工作）

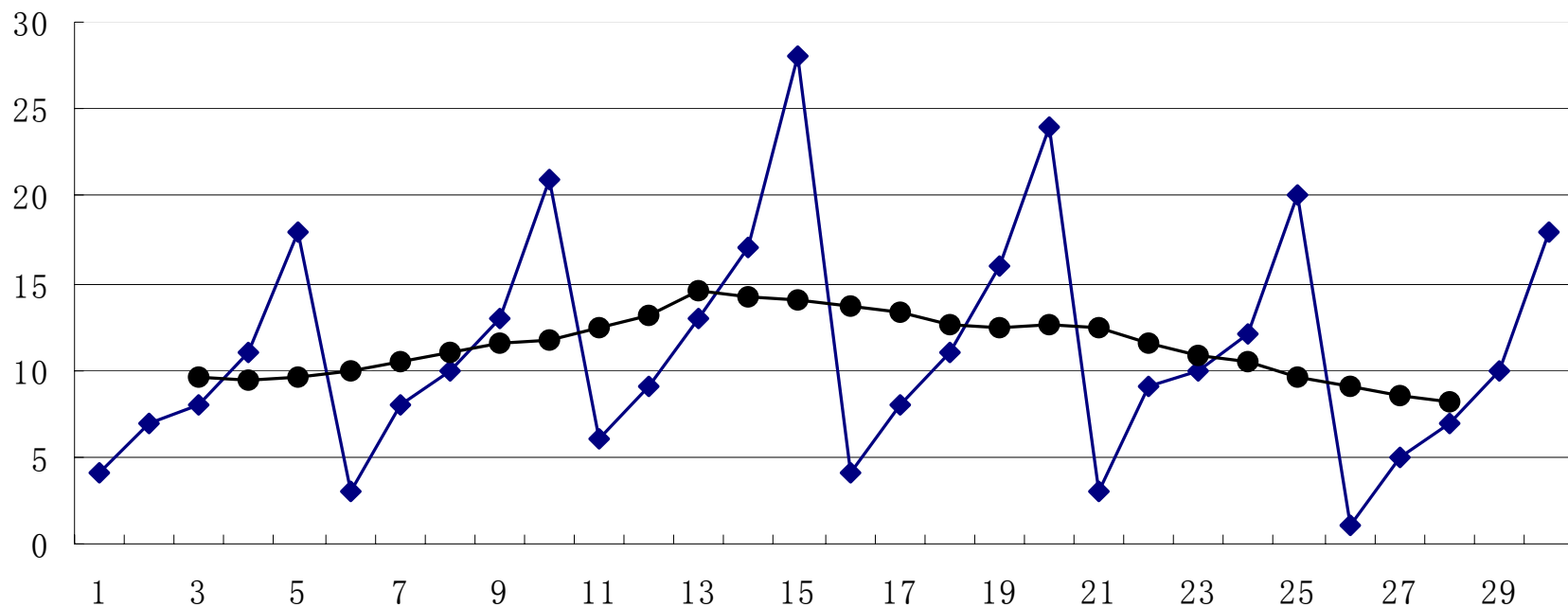
（二）曲线趋势

有时候时间序列的趋势不能简单地用直线趋势来描述，或者希望能直观地观察时间序列运动的方向和幅度。这时可用**滑动平均法**（moving average）。

其原理是用相对小的一段数据，找出平均值，以光滑原始序列的波动，然后滑动到另一段上，每小段的长短与我们观察的数据类型有关。如果是季节数据则把 4 个数据作为一小段来作平均；如果是月数据则把 12 个数据作为一小段来作平均；如果是日数据则把 5 个或 7 个数据作为一小段来作平均。

例3. 关于Jasbeer, Rajan. Co. 员工每周病事假记录。为预测每日病事假我们用滑动平均法求趋势曲线。由上所讨论可知，这儿每段合适的大小为。前个平均为，把对准第三行；从第二个数据开始到第六个数据的平均值为人，把对准第四行...。（见表7.1）。把它们画在原先图上，得趋势曲线（见图）。

当每小段数据为偶数时，要作两次滑动平均。首先把 d 点的平均值对准这 d 个点的中间 $(d+1)/2$ ，然后再对相邻平均值作平均，把这一次的平均值对准 $1+d/2$ 。把这些值连接起来，首项对准 $1+d/2$ 第个数据，构成一条趋势曲线。



图：病假人数散点图和平滑趋势图

例4：对例1给出的时间序列作滑动平均(数据见下表)

周	星期	病事假人数	$\sum 5s$	平均
1	1	4		
	2	7		
	3	8	48	9.6
	4	11	47	9.4
	5	18	48	9.6
2	1	3	50	10.0
	2	8	52	10.4
	3	10	55	11.0
	4	13	58	11.6
	5	21	59	11.8
3	1	6	62	12.4
	2	9	66	13.2
	3	13	73	14.6
	4	17	71	14.2
	5	28	70	14.0

4	1	4	68	13.6
	2	8	67	13.4
	3	11	63	12.6
	4	16	62	12.4
	5	24	63	12.6
5	1	3	62	12.4
	2	9	58	11.6
	3	10	54	10.8
	4	12	52	10.4
	5	20	48	9.6
6	1	1	45	9.0
	2	5	43	8.6
	3	7	41	8.2
	4	10		
	5	18		

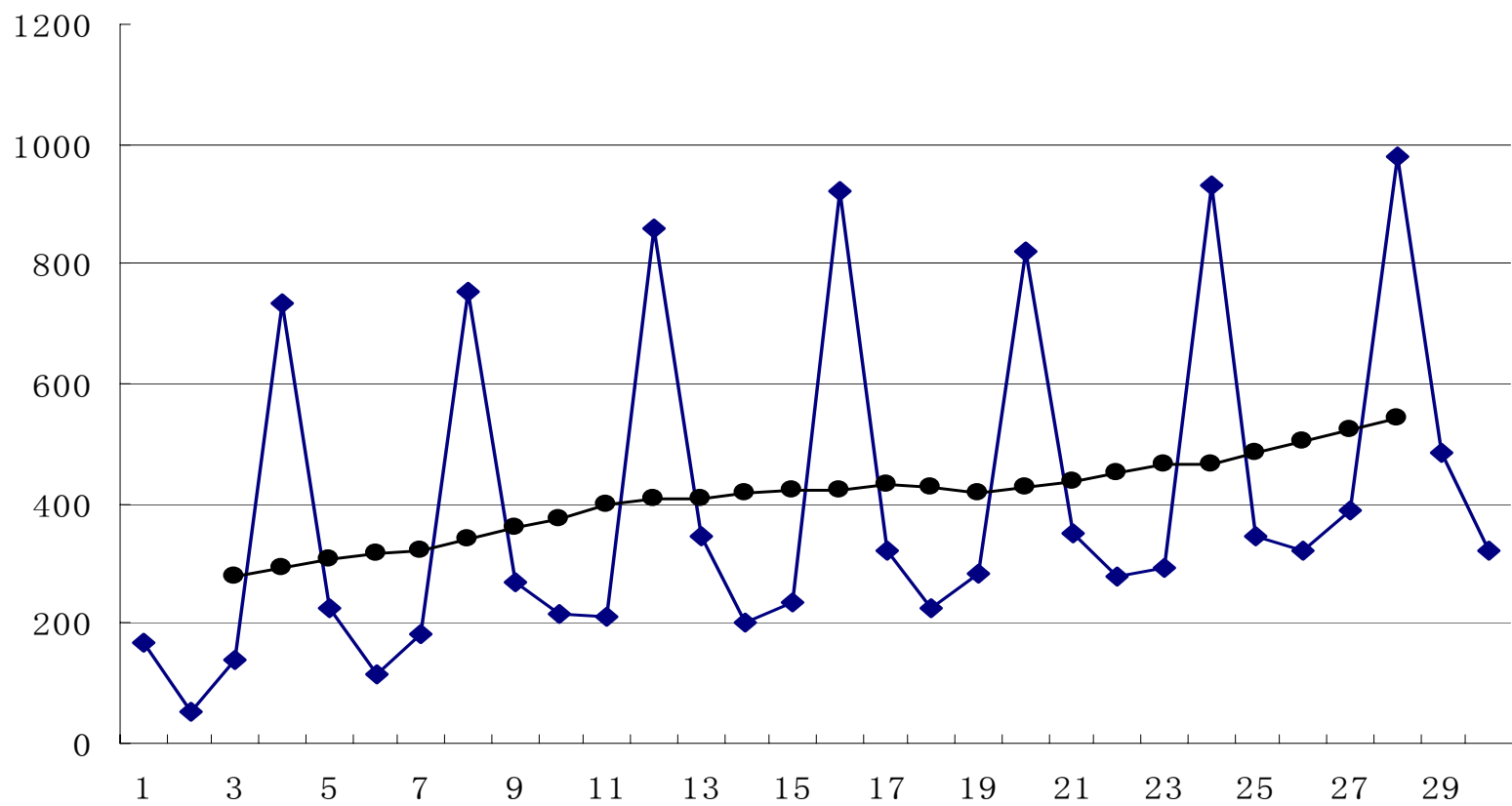
年.季度	销售量	$\sum 4s$	$\sum 8s$	滑动平均
1.1	166			
1.2	52			
		1091		
1.3	140		2240	280.00
		1149		
1.4	733		2360	295.00
		1211		
2.1	224		2463	307.88
		1252		

2.2	114		2524	315.500
		1272		
2.3	181		2589	323.625
		1317		
2.4	753		2734	341.750
		1417		
3.1	269		2863	357.875
		1446		
3.2	214		2999	374.875
		1553		
3.3	210		3182	397.750
		1629		

3.4	860		3247	405.875
		1618		
4.1	345		3259	407.375
		1641		
4.2	203		3344	418.000
		1703		
4.3	233		3385	423.125
		1682		
4.4	922		3385	423.125
		1703		
5.1	324		3457	432.125
		1754		

5.2	224		3408	426.000
		1654		
5.3	284		3336	417.000
		1682		
5.4	822		3420	427.000
		1738		
6.1	352		3487	435.875
		1749		
6.2	280		3606	450.750
		1857		
6.3	295		3707	463.375
		1850		

6.4	930		3740	467.500
		1890		
7.1	345		3875	484.375
		1985		
7.2	320		4018	502.250
		2033		
7.3	390		4204	525.500
		2171		
7.4	978		4342	542.750
		2171		
8.1	483			
8.2	320			



销售散点图和平滑趋势图

§ 3. 季节因素 (The seasonal factors)

时间序列的值与季节有关，这是大家很熟悉的。如产品销售等的季节效果，在政府统计数据中也有“季节调整”。如：“上月失业人数增加，但考虑到季节因素，整个趋势是下降的”等等。

分析时间序列季节因素目的之一是何时实际数据与平均值（趋势）不同，找出变化幅度，以便计划生产，储存和政策的制订。

1. 加法模型中季节因子的确定

考虑简化的时间序列加法模型：

$$y = T + S + R$$

其中趋势T可由前面得出（最小二乘所得直线，或滑动平均所得曲线），再设 $R=0$ （因为 $\bar{R}$ 应为0）则

$$y = T + S \Rightarrow S = y - T$$

即加法模型中的季节因子可以由上式得到。

例5. 考虑例1中时间序列的季节因子。由公式 $S=y-T$ 可得到每一季节的值。把同一季度的季节因子作一平均，得到一个季度的平均值，如下表所示。

表3 加法模型下的季节因子

年	第一季度	第二季度	第三季度	第四季度
1	-87.654	-211.668	-133.683	449.302
2	-69.712	-189.727	-132.742	429.243
3	-64.771	-129.786	-143.801	496.185
4	-28.830	-180.845	-160.859	518.126
5	-89.889	-199.903	-149.918	378.067
6	-101.947	-183.962	-178.977	446.009
7	-149.006	-184.021	-124.036	453.950
8	-51.065	-224.080		
总计	-642.875	-1503.992	-1024.015	3170.882
平均	-80.359	-187.999	-146.288	452.983

由上表得季节因子分别为

-80.359 -187.999 -146.288 455.983

由于季节因子是作为一年中指标上下浮动的平均效果，因此它们的和相加应该为零。但上述四个季节因子不满足这一条件：

$$-80.359 + (-187.999) + (-146.288) + 455.983 = 38.337$$

因此需要对季节因子作修正，其中修正因子为

$$\ell = \frac{\text{季节因子总和}}{\text{周期中点数}} \quad \left(\text{这儿 } \ell = \frac{38.337}{4} = 9.58425 \right)$$

在每个季节因子上减去 ℓ 得

-89.943 -197.583 -155.872 443.399

这就是修正后的季节因子。

2. 乘法模型中季节因子的确定

在乘法模型(2')中

$$y = T \times S \times R$$

T 用 § 2 中方法给出（直线或曲线）先设 R 为 1，则

$$y = T \times S \Rightarrow S = \frac{y}{T}$$

对每季度用此公式得每个季度的季节因子。如对例 1 中的数据用线性趋势和乘法模型(2),得到每个季度的季节因子，然后对同一季度的季节因子再作算术平均得出四个季节因子，（见下表）

年.季 度	销售 额 y	线性趋势 T	y/T
1.1	166	253.65	0.65
1.2	52	263.67	0.20
1.3	140	273.68	0.51
1.4	733	283.70	2.58
2.1	224	293.71	0.76
2.2	114	303.73	0.38
2.3	181	313.74	0.58
2.4	753	323.76	2.33
3.1	269	333.77	0.81
3.2	214	343.79	0.62
3.3	210	353.80	0.59
3.4	860	363.82	2.36
4.1	345	373.83	0.92
4.2	203	383.84	0.53
4.3`	233	393.86	0.59

4.4	922	403.87	2.28
5.1	324	413.89	0.78
5.2	224	423.90	0.53
5.3	284	433.92	0.65
5.4	822	443.93	1.85
6.1	352	453.95	0.78
6.2	280	463.96	0.60
6.3	295	473.98	0.62
6.4	930	483.99	1.92
7.1	345	494.01	0.70
7.2	320	504.02	0.63
7.3	390	514.04	0.76
7.4	978	524.05	1.87
8.1	483	534.06	0.90
8.2	320	544.08	0.59

表5 乘法模型下的季节因子

年	第一季度	第二季度	第三季度	第四季度
1	0.654	0.197	0.512	2.584
2	0.763	0.175	0.577	2.326
3	0.806	0.622	0.594	2.364
4	0.923	0.529	0.592	2.282
5	0.783	0.528	0.655	1.852
6	0.775	0.603	0.622	1.922
7	0.698	0.635	0.759	1.866
8	0.904	0.588		
总计	6.307	4.079	4.309	15.196
平均	0.788	0.510	0.616	2.171

由此得四个季节因子为

78.8% 51.0% 61.6% 217.1%

由于是乘法模型，它的季节因子用算术平均求得，因此它们百分数的和相加应该为400，这儿

$$78.8 + 51.0 + 61.6 + 217.1 = 408.5$$

它们的和不为400，因此要对它们作修正。修正因子为：

$$\ell = \frac{400}{\text{因子总和}} = \frac{400}{408.5} = 0.979192$$

然后每个季节因子乘以修正因子得修正后的季节因子：

$$77.16\% \quad 49.94\% \quad 60.32\% \quad 212.58\%$$

对乘法模型而言也可用几何平均来求季节因子，本例中4个季节因子分别

$$78.4\% \quad 48.2\% \quad 61.2\% \quad 215.4\%$$

再作修正。修正因子为

$$(78.4 + 48.2 + 61.2 + 215.4 - 400) / 4 = 0.008 \quad ,$$

固修正的季节因子为:

$$77.6\% \quad 47.4\% \quad 60.4\% \quad 214.7\%$$

用几何平均得到的季节因子与用算术平均得到的季节因子相比，相差不多。

3. 时间序列的季节调整

用长期数据计算出季节因子，然后从数据中剔除季节影响，以显示在没有季节变化条件下将来的趋势会如何运动。这在许多出版物中被称为"季节调整":

由上知，模型为

$$y - S = T + C + R$$

注意，模型好坏取决于季节因子的精确识别。

§ 4 循环因子

一般而言循环因子很难识别，因为由循环因子的定义知道循环因子的识别关键是要有大量的历史数据，通常我们得到的数据相对而言比较短，因此很难识别出循环因子。在社会生活中已知的一些循环有：

（1）**Kuznet**长波：各国的国民生产总值和人口迁移大体有**20**年的循环。

（2）商业循环：在贸易循环等流通领域，通常有**2-3**年的循环，它们是长返的但不是周期的。

§ 5 随机因素和残差

(The random factor and residual or factor)

在不计循环的加法模型中随机因子（也称为残差）可由下式算出：

$$R = y - T - \bar{S}$$

其中 $\bar{S}$ 为修正后的加法模型季节因子。在乘法模型中，随机因子为

$$R = \frac{y}{T \times \bar{S}}$$

残差为 $y - T\bar{S}$ ，其中 $\bar{S}$ 为修正后的乘法模型季节因子。

下表列出的是例1中在线性趋势下用加法模型和乘法模型得到的残差。

表6 Jasbeer,Rajan & Co.的残差因子

年. 季 度	加法模型 残差	平方和	乘法模 型残差	平方和
1.1	2.2898	5.2434	-29.81	888.7458
1.2	-14.0852	198.394	-79.64	6342.051
1.3	22.1890	492.350	-24.97	623.7087
1.4	5.9033	34.8488	129.96	16888.49
2.1	20.2311	409.298	-2.736	7.4845
2.2	7.8560	61.7171	-37.64	1416.498
2.3	23.1302	535.008	-8.12	65.9550
2.4	-14.1554	200.377	64.805	4199.634
3.1	25.1724	633.648	11.340	128.6019
3.2	67.7973	4596.47	42.364	1794.726
3.3	12.0715	145.721	-3.268	10.6823
3.4	52.7858	2786.34	86.653	7508.821
4.1	61.1136	3734.88	56.416	3182.803
4.2	16.7386	280.179	11.365	129.1588
4.3	-4.9872	24.8724	-4.415	19.4965

4.4	74.7271	5584.1	63.502	4032.5
5.1	0.0549	0.0030	4.4924	20.182
5.2	-2.3202	5.3832	12.365	152.90
5.3	5.9540	35.451	22.437	503.44
5.4	-65.3316	4268.2	-121.6	14798.
6.1	-12.0038	144.09	1.5685	2.4601
6.2	13.6211	185.53	48.366	2339.3
6.3	-23.1047	533.82	9.2903	86.310
6.4	2.6096	6.8102	-98.80	9761.4
7.1	-59.0626	3488.4	-36.36	1321.7
7.2	13.5624	183.94	68.367	4674.0
7.3	31.8366	1013.6	80.143	6422.9
7.4	10.5509	111.32	-136.0	18483.
8.1	38.8787	1511.6	70.721	5001.4
8.2	-26.4964	702.06	48.367	2339.4
平方和		31914	249.16	113146
平均值		1063.8		3771.5

一个现实的问题是我们作时间序列分析中是用加法模型还是乘法模型？一个常用的方法是看残差平方和或残差平方和的平均值，在加法模型中残差即为随机因子，在乘法模型中残差等于 $y - TS$ 残差的平方和。残差平方和或残差平方和的平均值小的那种模型较好。这儿是加法模型较好。

§ 6 预报 (Predictions)

由上述模型，如果我们能识别出趋势和季节因子。以及能够拓展趋势到将来(比如线形趋势)，则我们就能作预报(是否准确由趋势及季节因子及误差所决定)。如果趋势是用滑动平均来做的，则预报要困难些。一般作短期预报，因为长期的往往不准。

例6. 试在线性趋势下分别用加法模型和乘法模型预报 *Jasbeer Rajan & Co.* 第八年第三和第四季度的销售额。

解. 在线性趋势方程

$$y = 243.639 + 10.015 i$$

中, 分别取 $i = 31$ 和 $i = 32$, 得线性趋势值为

$$T_{31} = 554.092, \quad T_{32} = 564.106$$

在加法模型中, 第三和第四季度的季节因子分别为

$$S_3 = -155.872 \quad S_4 = 443.399$$

用预测公式 $\hat{y} = T + S$

得到第八年第三和第四季度预测销售额分别为

$$\hat{A}_{31} = 398.222, \quad \hat{A}_{32} = 1007.508$$

在乘法模型中，第三和第四季度的季节因子分布为

$$S_3 = 60.32\%, S_4 = 212.58\%$$

用预测公式 $\hat{A} = T \times S$

得到第八年第三和第四季度预测销售额分别为

$$\hat{A}_{31} = 334.228, \hat{A}_{32} = 1199.177$$

§ 7指数加权滑动平均

用滑动平均作趋势的一个缺点是无法用它作预报。若要作短期预报,(如商品储存控制, 可以用"指数加权滑动平均"(EWMA)技术。该技术考虑了过去已发生过的数据的走向, 对历史值和当前值作加权平均来预报下一时刻的值。加权的原则是现在比过去大, 且以指数形式下降-----降低过老数据在预报中的作用。具体为

$$\hat{y}_{t+1} = \alpha y_t + (1 - \alpha) \hat{y}_t \quad (1)$$

这儿 α 是平滑常数, 且 $0 < \alpha < 1$. 这种预报称为指数加权滑动平均 预报。

预报的好坏与 α 值的选取有较大关系, α 取小值, 导致过去数据遗忘得较慢。曲线趋势比较平缓。如果 α 大, 过去数据遗忘得快, 曲线趋势变化就较快,

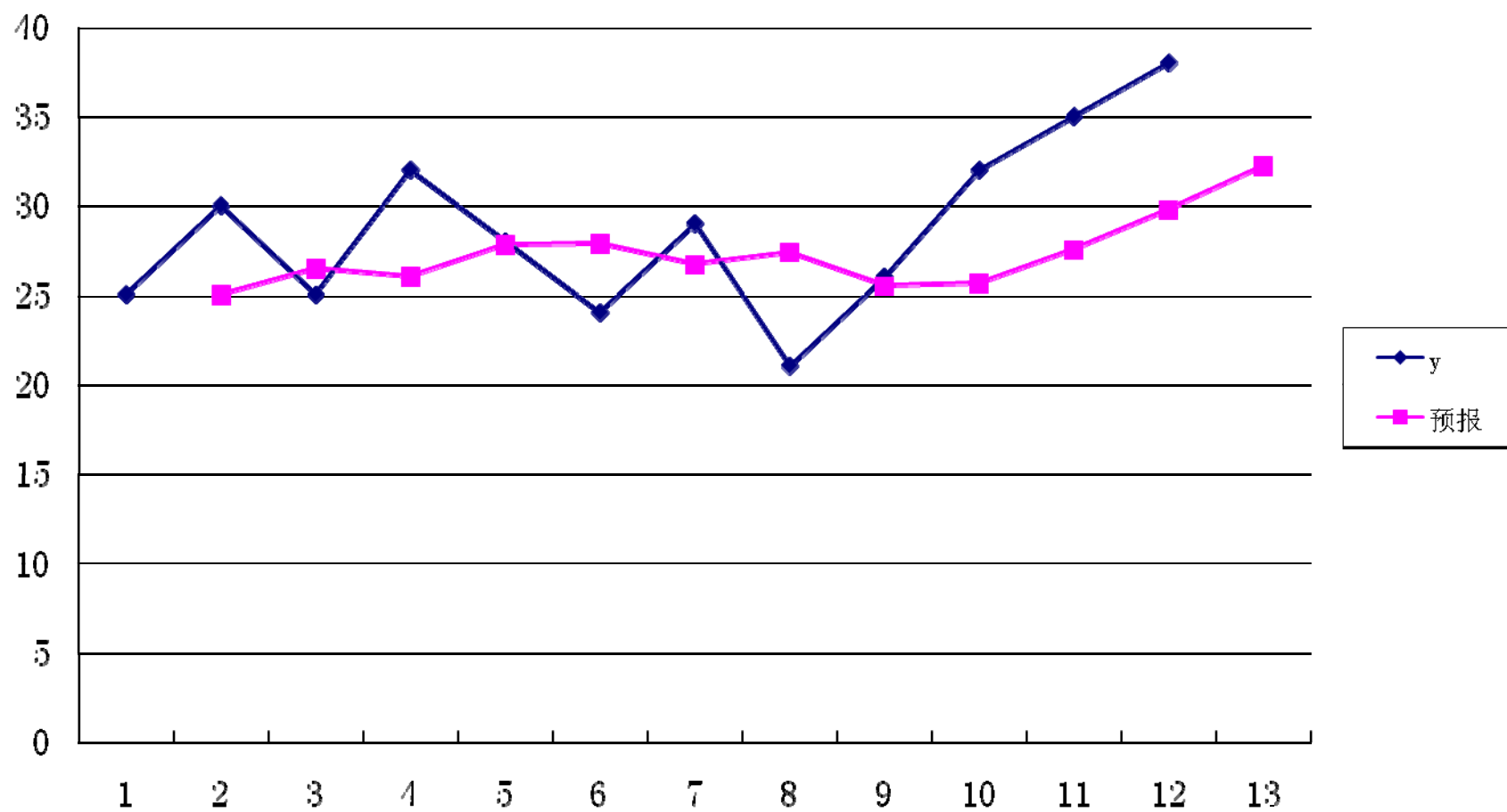
例7.某小商店用指数加权滑动平均(EWMA)模型预报下个月吸引新顾客。取 $\alpha=0.3$ 。数据如下：

月	顾客数 (Y_t)	$0.3 Y_t$	$0.7 \hat{y}_t$	$\hat{y}_{t+1}$
1	25	7.5	17.5	
2	30	9	17.5	25.00
3	25	7.5	18.55	26.500
4	32	9.6	18.235	26.050
5	28	8.4	19.485	27.835
6	24	7.2	19.519	27.885
7	29	8.7	18.703	26.719
8	21	6.3	19.182	27.403
9	26	7.8	17.838	25.482
10	32	9.6	17.946	25.638
11	35	10.5	19.282	27.546
12	38	11.4	20.848	29.782
1				32.248

由指数加权滑动平均公式可得，从二月份到到下一年一月的预报值（见上表）如取 $\alpha = 0.65$ 则有如下的预报（下表和图）：

月	顾客数(Y_t)	$0.65 Y_t$	$0.35 \hat{y}_t$	$\hat{y}_{t+1}$
1	25	16.25	8.75	
2	30	19.5	8.75	25.000
3	25	16.25	9.888	28.250
4	32	20.8	9.148	26.138
5	28	18.2	10.482	29.948`
6	24	15.6	10.039	28.682
7	29	18.85	8.974	25.639
8	21	13.65	9.738	27.824
9	26	16.9	8.186	23.388
10	32	20.8	8.780	25.086
11	35	22.75	10.353	29.580
12	38	24.7	11.586	33.103
1				36.286

例6 $\alpha=0.3$



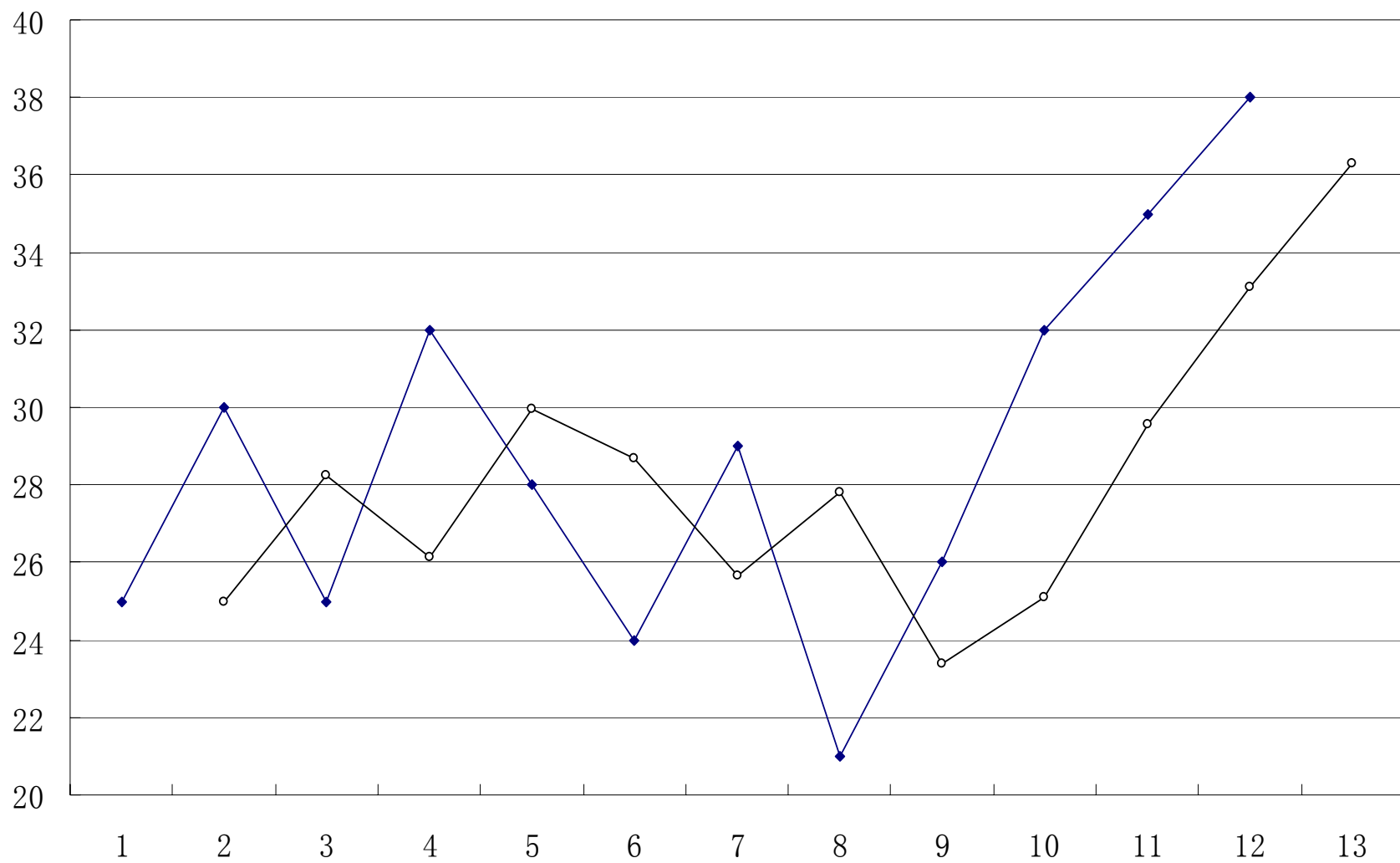


图 7 $\alpha=0.65$

在实际应用中用哪个 α 好？我们的原则是预报误差平方和越小越好。具体实施过程中可以试验不同的 α ，以决定最好的 α 值。如上面小商店预报中，通过计算得 $\alpha = 0.3$ 时，预报误差平方和为**283.81**；当 $\alpha = 0.65$ 时预报误差平方和为**261.49**。因此当 $\alpha = 0.65$ 时预报较好。

第二部分建模

一、统计时间序列模型简介

简要地介绍一下如何用统计方法来建模。以下
设 $\{\varepsilon_t\}$, $t = \dots, -1, 0, 1, \dots$

为有相同分布和相互独立的误差序列, $\{y_t\}$ 为时间序列。

(1) 滑动平均 (MA) 模型 (Moving Average Model)

如果存在 q 个常数 $\theta_1, \dots, \theta_q$, 使得时间序列

$$y_t = \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q} \quad (1)$$

则称该时间序列是 q 阶滑动平均, 记为 $MA(q)$ 。

(2).自回归 (AR) 模型 (Autoregressive Model)

若存在 p 个常数 $\alpha_1, \dots, \alpha_p$, 使得时间序列,

$$y_t = \alpha_1 y_{t-1} + \alpha_2 y_{t-2} + \dots + \alpha_p y_{t-p} + \varepsilon_t \quad (2)$$

则称该时间序列为 p 阶自回归模型, 记为 $AR(p)$

(3) ARMA (自回归滑动平均)模型

若存在 $p+q$ 个常数 $\alpha_1, \dots, \alpha_p; \theta_1, \dots, \theta_q$, 使得

$$y_t = \alpha_1 y_{t-1} + \alpha_2 y_{t-2} + \dots + \alpha_p y_{t-p} + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q} \quad (3)$$

则称 $\{y_t\}$ 为 (p, q) 阶混合自回归滑动平均模型, 记为 $ARMA(p, q)$ 。

滑动平均模型指出数据 y_t 由时刻 t 及 t 前 q 个时刻的误差构成。而自回归模型指出当前数据由前 p 个时刻的数据及当前误差组成，如 $p = 1$

$$y_t = \alpha_1 y_{t-1} + \varepsilon_t$$

我们可以把 y_t 看成当前时刻的股价，则上式表示当前股价与前一时刻有关。自回归滑动平均则是这两种模型的混合，在一定的条件下，它们可以转换。

商业上用的最多的是 $AR(p)$ 模型，其中 p 称为模型的阶，它可以直接由经验得到，也可以通过
对数据的分析得到 p 的一个估计。

二、随机时间序列的特性分析

1、时序特性的研究工具

(1) 自相关

构成时间序列的每个序列值 $X_t, X_{t-1}, X_{t-2}, \dots, X_{t-k}$ 之间的简单相关关系称为自相关。自相关程度由自相关系数 γ_k 度量, 表示时间序列中相隔 k 期的观测值之间的相关程度。

$$\gamma_k = \frac{\sum_{t=1}^{n-k} (X_t - \bar{X})(X_{t+k} - \bar{X})}{\sum_{t=1}^n (X_t - \bar{X})^2}$$

注1: n 是样本量, k 为滞后期, $\bar{X}$ 代表样本数据的算术平均值

注2: 自相关系数 γ_k 的取值范围是 $[-1, 1]$ 且 $|\gamma_k|$ 越接近1, 自相关程度越高

(2) 偏自相关

偏自相关是指对于时间序列 X_t ，在给定 $X_{t-1}, X_{t-2}, \dots, X_{t-k+1}$ 的条件下， X_t 与 X_{t-k} 之间的条件相关关系。其相关程度用偏自相关系数 φ_{kk} 度量，有 $-1 \leq \varphi_{kk} \leq 1$

$$\varphi_{kk} = \begin{cases} \gamma_1 & k=1 \\ \frac{\gamma_k - \sum_{j=1}^{k-1} \varphi_{k-1,j} \cdot \gamma_{k-j}}{1 - \sum_{j=1}^{k-1} \varphi_{k-1,j} \cdot \gamma_j} & k=2, 3, \dots \end{cases}$$

其中 γ_k 是滞后 k 期的自相关系数，

$$\varphi_{kj} = \varphi_{k-1,j} - \varphi_{kk} \varphi_{k-1,k-j}, j=1, 2, \dots, k-1$$

例如：平稳一阶自回归AR(1)过程

$$x_t = \phi_1 x_{t-1} + u_t, \quad |\phi_1| < 1$$

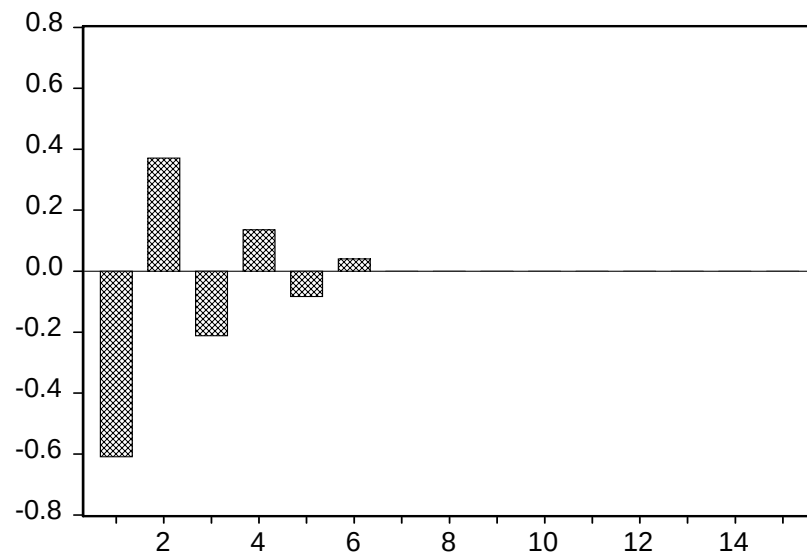
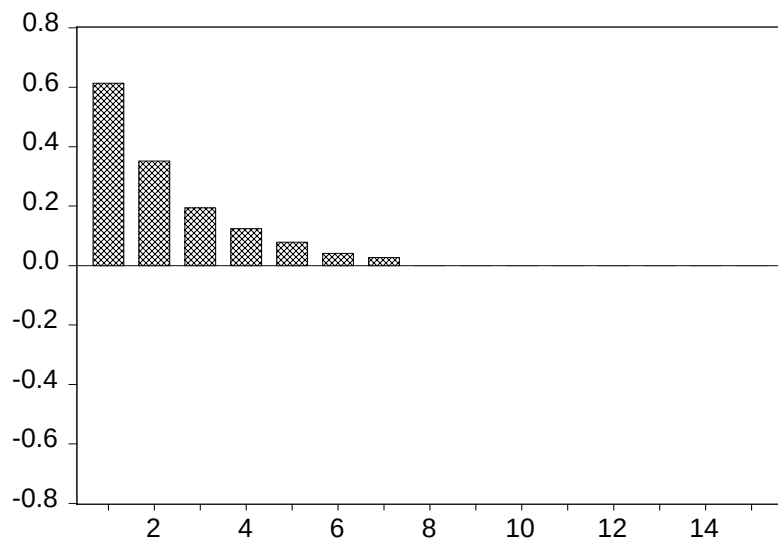
的自相关函数

$$\rho_k = \phi_1^k, \quad (k \geq 0)$$

$$\begin{aligned} x_t x_{t-k} &= \phi_1 x_{t-1} x_{t-k} + u_t x_{t-k} \\ &\vdots \\ \gamma_k &= \phi_1 \gamma_{k-1} \end{aligned} \Rightarrow \rho_k = \phi_1 \rho_{k-1} = \dots = \phi_1^k \rho_0 = \phi_1^k$$

两边取期望

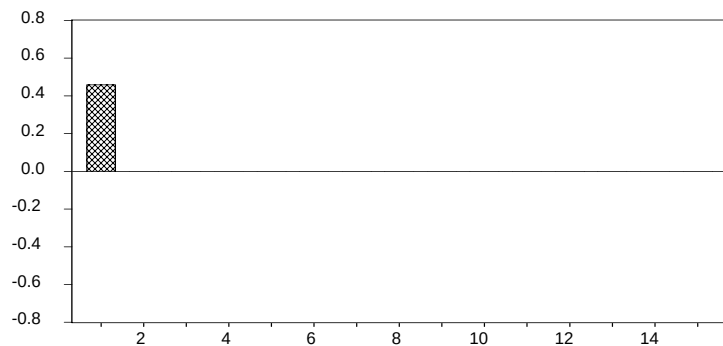
对于平稳序列有 $|\phi_1| < 1$ 。所以当 ϕ_1 为正时，自相关函数按指数衰减至零，当 ϕ_1 为负时，自相关函数正负交错地指数衰减至零。



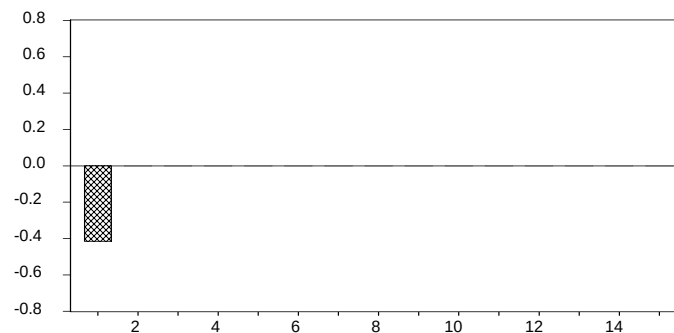
一阶移动平均**MA(1)** 过程的自相关函数。

$$X_t = u_t + \theta_1 u_{t-1}$$

$$\rho_k = \frac{\gamma_k}{\gamma_0} = \begin{cases} \frac{\theta_1}{1 + \theta_1^2} & k=1 \\ 0 & K > 1 \end{cases}$$



$$\theta_1 > 0$$

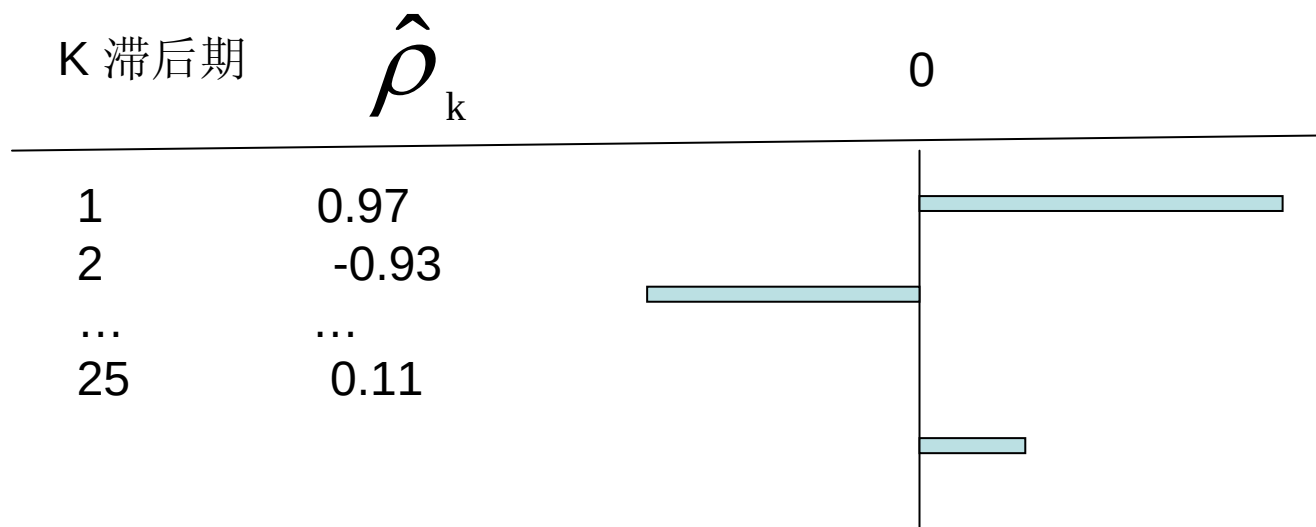


$$\theta_1 < 0$$

可见**MA(1)** 过程的自相关函数具有截尾特征。

➤ 相关图

相关图是对自相关函数的估计。由于**MA**过程和**ARMA**过程中的**MA**分量的自相关函数具有截尾特性，所以通过相关图可以估计**MA**过程的阶数 q 。相关图是识别**MA**过程阶数和**ARMA**过程中**MA**分量阶数的一个重要方法。实际应用中相关图一般取 $k = 15$ 就足够了。



三、模型的识别与建立

在需要对一个时间序列运用**B-J**方法建模时，应运用序列的自相关与偏自相关对序列适合的模型类型进行识别，确定适宜的阶数 d, D, p, q 以及 P, Q （消除季节趋势性后的平稳序列）

1、自相关函数与偏自相关函数

(1) MA (q) 的自相关与偏自相关函数

自协方差函数

$$\gamma_k = \begin{cases} (1 + \theta_1^2 + \cdots + \theta_q^2) \sigma^2, & k = 0 \\ (-\theta_k + \theta_1 \theta_{k+1} + \cdots + \theta_{q-k} \theta_q) \sigma^2, & 1 \leq k \leq q \\ 0, & k > q \end{cases}$$

$Du_t = \sigma^2$ 是白噪声序列的方差

样本自相关函数

$$\rho_k = \frac{\gamma_k}{\gamma_0} = \begin{cases} 1, & k = 0 \\ \frac{-\theta_k + \theta_1 \theta_{k+1} + \cdots + \theta_{q-k} \theta_q}{1 + \theta_1^2 + \cdots + \theta_q^2}, & 1 \leq k \leq q \\ 0, & k > q \end{cases}$$

MA (q) 序列的自相关函数 ρ_k 在 $k > q$ 以后全都是0，这种性质称为自相关函数的 q 步截尾性；偏自相关函数随着滞后期 k 的增加，呈现指数或者正弦波衰减，趋向于0，这种特性称为偏自相关函数的拖尾性

(2) AR (p) 序列的自相关与偏自相关函数

偏自相关函数

$$\varphi_{kk} = \begin{cases} \varphi_k, & 1 \leq k \leq p \\ 0, & k > p \end{cases}$$

是 p 步截尾的；

自协方差函数 γ_k 满足 $\varphi(B)\gamma_k = 0$

自相关函数 ρ_k 满足 $\varphi(B)\rho_k = 0$

它们呈指数或者正弦波衰减，具有拖尾性

(3) ARMA (p, q) 序列的自相关与偏自相关函数均是拖尾的

2、模型的识别

自相关函数与偏自相关函数是识别ARMA模型的最主要工具，B-J方法主要利用相关分析法确定模型的阶数。

若样本自协方差函数 γ_k 在 q 步截尾，则判断 X_t 是MA(q)序列

若样本偏自相关函数 φ_{kk} 在 p 步截尾，则可判断 X_t 是AR(p)序列

若 γ_k ， φ_{kk} 都不截尾，而仅是依负指数衰减，这时可初步认为 X_t 是ARMA序列，它的阶要由从低阶到高阶逐步增加，再通过检验来确定。

但实际数据处理中，得到的样本自协方差函数和样本偏自相关函数只是 γ_k 和 φ_{kk} 的估计，要使它们在某一步之后全部为0几乎是不可能的，而只能是在某步之后围绕零值上下波动，故对于 γ_k 和 φ_{kk} 的截尾性只能借助于统计手段进行检验和判定。

(1) γ_k 的截尾性判断

对于每一个 q ，计算 $\gamma_{q+1}, \dots, \gamma_{q+M}$ (M 一般取 $\sqrt{N}$ 左右)，考察其中满足

$$|\gamma_k| \leq \frac{1}{\sqrt{N}} \sqrt{\gamma_0^2 + 2 \sum_{l=1}^q \gamma_l^2} \quad \text{或} \quad |\gamma_k| \leq \frac{2}{\sqrt{N}} \sqrt{\gamma_0^2 + 2 \sum_{l=1}^q \gamma_l^2}$$

的个数是否为 M 的68.3%或95.5%。

如果当 $1 \leq k \leq q_0$ 时， γ_k 明显地异于0，而 $\gamma_{q_0+1}, \dots, \gamma_{q_0+M}$ 近似为0，且满足上述不等式的个数达到了相应的比例，则可近似地认为 γ_k 在 q_0 步截尾

(2) φ_{kk} 的截尾性判断

作如下假设检验: $M = \sqrt{N}$

$$H_0 : \varphi_{p+k, p+k} = 0, k = 1, \dots, M$$

$$H_1 : \text{存在某个 } k, \text{ 使 } \varphi_{kk} \neq 0, \text{ 且 } p < k < M + p$$

$$\text{统计量 } \chi^2 = N \sum_{k=p+1}^{p+M} \varphi_{kk}^2 \sim \chi_M^2$$

$\chi_M^2(\alpha)$ 表示自由度为 M 的 χ^2 分布的上侧 α 分位数点
对于给定的显著性水平 $\alpha > 0$, 若 $\chi^2 > \chi_M^2(\alpha)$, 则认为
样本不是来自 AR (p) 模型; $\chi^2 < \chi_M^2(\alpha)$, 可认为
样本来自 AR (p) 模型。

注: 实际中, 此判断方法比较粗糙, 还不能定阶, 目前流行的方法是 H. Akaike 信息定阶准则 (AIC)

(3) AIC准则确定模型的阶数

AIC定阶准则： S 是模型的未知参数的总数

$\hat{\sigma}^2$ 是用某种方法得到的方差 σ^2 的估计

N 为样本大小，则定义**AIC**准则函数

$$AIC(S) = \ln \hat{\sigma}^2 + \frac{2S}{N}$$

用**AIC**准则定阶是指在 p, q 的一定变化范围内，寻求使得 $AIC(S)$ 最小的点 $(\hat{p}, \hat{q})$ 作为 (p, q) 的估计。

AR (p) 模型：

$$AIC = \ln \hat{\sigma}^2 + \frac{2p}{N}$$

ARMA (p, q) 模型：

$$AIC = \ln \hat{\sigma}^2 + \frac{2(p+q)}{N}$$

3、参数估计

在阶数给定的情形下模型参数的估计有三种基本方法：矩估计法、逆函数估计法和最小二乘估计法，这里仅介绍矩估计法

(1) AR (p) 模型

$$\begin{pmatrix} \hat{\varphi}_1 \\ \hat{\varphi}_2 \\ \vdots \\ \hat{\varphi}_p \end{pmatrix} = \begin{pmatrix} 1 & \hat{\rho}_1 & \cdots & \hat{\rho}_{p-1} \\ \hat{\rho}_1 & 1 & \cdots & \hat{\rho}_{p-2} \\ \vdots & \vdots & & \vdots \\ \hat{\rho}_{p-1} & \hat{\rho}_{p-2} & \cdots & 1 \end{pmatrix}^{-1} \begin{pmatrix} \hat{\rho}_1 \\ \hat{\rho}_2 \\ \vdots \\ \hat{\rho}_p \end{pmatrix}$$

白噪声序列 $\mathbf{u}_t$ 的方差的矩估计为

$$\hat{\sigma}^2 = \gamma_0 - \sum_{j=1}^p \hat{\varphi}_j \hat{\gamma}_j$$

(2) MA (q) 模型

$$\begin{cases} (1 + \hat{\theta}_1^2 + \cdots + \hat{\theta}_q^2) \hat{\sigma}^2 = \hat{\gamma}_0 \\ (-\hat{\theta}_k + \hat{\theta}_1 \hat{\theta}_{k+1} + \cdots + \hat{\theta}_{q-k} \hat{\theta}_q) \hat{\sigma}^2 = \hat{\gamma}_k, k = 1, \cdots, q \end{cases}$$

(3) ARMA(p, q) 模型的参数矩估计分三步:

i) 求 $\varphi_1, \varphi_2, \cdots, \varphi_p$ 的估计

$$\begin{pmatrix} \hat{\varphi}_1 \\ \hat{\varphi}_2 \\ \vdots \\ \hat{\varphi}_p \end{pmatrix} = \begin{pmatrix} \hat{\gamma}_q & \hat{\gamma}_{q-1} & \cdots & \hat{\gamma}_{q-p+1} \\ \hat{\gamma}_{q+1} & \hat{\gamma}_q & \cdots & \hat{\gamma}_{q-p+2} \\ \vdots & \vdots & & \vdots \\ \hat{\gamma}_{q+p-1} & \hat{\gamma}_{q+p-2} & \cdots & \hat{\gamma}_q \end{pmatrix}^{-1} \begin{pmatrix} \hat{\gamma}_{q+1} \\ \hat{\gamma}_{q+2} \\ \vdots \\ \hat{\gamma}_{q+p} \end{pmatrix}$$

ii) 令 $Y_t = X_t - \hat{\phi}_1 X_{t-1} - \cdots - \hat{\phi}_p X_{t-p}$, 则 Y_t

的自协方差函数的矩估计为

$$\hat{\gamma}_k^{(Y)} = \sum_{i=0}^p \sum_{j=0}^p \hat{\phi}_i \hat{\phi}_j \hat{\gamma}_{k+j-i}, \quad \hat{\phi}_0 = -1$$

iii) 把 Y_t 近似看作MA (q) 序列, 利用 (2)

对MA (q) 序列的参数估计方法即可

4、模型检验

对于给定的样本数据 $X_1, \dots, X_N$ ，我们通过相关分析法和AIC准则确定了模型的类型和阶数，用矩估计法确定了模型中的参数，从而建立了一个ARMA模型，来拟合真正的随机序列。但这种拟合的优劣程度如何，主要应通过实际应用效果来检验，也可通过数学方法来检验。

下面介绍模型拟合的残量自相关检验，即白噪声检验：

对于ARMA模型，应逐步由ARMA（1，1），ARMA（2，1），ARMA（1，2），ARMA（2，2），...依次求出参数估计，对AR（ p ）和MA（ q ）模型，先由 γ_k 和 φ_{kk} 的截尾性初步定阶，再求参数估计。

一般地，对ARMA(p, q) 模型

$$u_t = X_t - \sum_{i=1}^p \hat{\varphi}_i X_{t-i} + \sum_{j=1}^q \hat{\theta}_j u_{t-j}$$

取初值 $u_0, u_{-1}, \dots, u_{1-q}$ 和 $X_0, X_{-1}, \dots, X_{1-p}$ (可取它们等于0, 因为它们均值为0), 可递推得到残量估计 $\hat{u}_1, \hat{u}_2, \dots, \hat{u}_N$
现作假设检验:

$H_0: \hat{u}_1, \hat{u}_2, \dots, \hat{u}_N$ 是来自白噪声的样本

$$\text{令 } \hat{\gamma}_j^{(u)} = \frac{1}{N} \sum_{t=1}^{N-j} \hat{u}_{t+j} \hat{u}_t \quad j = 0, 1, \dots, k$$

$$\hat{\rho}_j^{(u)} = \frac{\hat{\gamma}_j^{(u)}}{\hat{\gamma}_0^{(u)}} \quad j = 1, \dots, k$$

$$Q_k = \sum_{j=1}^k \left(\sqrt{N} \hat{\rho}_j^{(u)} \right)^2 = N \sum_{j=1}^k \left(\hat{\rho}_j^{(u)} \right)^2$$

其中 k 取 $\frac{N}{10}$ 左右。则当 H_0 成立时, Q_k 服从自由度为 k 的 χ^2 分布。

对给定的显著性水平 α , 若 $Q_k > \chi_k^2(\alpha)$, 则拒绝 H_0 , 即模型与原随机序列之间拟合得不好, 需重新考虑建模; 若 $Q_k < \chi_k^2(\alpha)$, 则认为模型与原随机序列之间拟合得较好, 模型检验被通过。

四、模型的预测

若模型经检验是合适的，也符合实际意义，可用作短期预测.

B-J方法采用L步预测，即根据已知 n 个时刻的序列观测值 $X_1, X_2, \dots, X_n$ ，对未来的 $n+L$ 个时刻的序列值做出估计，线性最小方差预测是常用的一种方法. 其主要思想是使预测误差的方差达到最小. 若 $\hat{Z}_n(L)$ 表示用模型做的L步平稳线性最小方差预测，那么，预测误差

$$e_n(L) = X_{n+L} - \hat{Z}_n(L)$$

并使 $E[e_n(L)]^2 = E[X_{n+L} - \hat{Z}_n(L)]^2$ 达到最小.

1、AR (p) 序列预测

模型 (1) :

$$X_t = \varphi_1 X_{t-1} + \varphi_2 X_{t-2} + \cdots + \varphi_p X_{t-p} + u_t$$

的L步预测值为

$$\hat{Z}_n(L) = \varphi_1 \hat{Z}_n(L-1) + \varphi_2 \hat{Z}_n(L-2) + \cdots + \varphi_p \hat{Z}_n(L-p)$$

其中 $\hat{Z}_n(-j) = X_{n-j}$ ($j \geq 0$)

2、MA (q) 的预测

对模型 (3) :

$$X_t = u_t - \theta_1 u_{t-1} - \theta_2 u_{t-2} - \cdots - \theta_q u_{t-q}$$

当 $L > q$ 时, 由于 $X_{n+L} = u_{n+L} - \theta_1 u_{n+L-1} - \theta_2 u_{n+L-2} - \cdots - \theta_q u_{n+L-q}$

可见所有白噪声的时刻都大于 n , 故与历史取值无关,

从而 $\hat{Z}_n(L) = 0$;

当 $L \leq q$ 时, 各步预测值可写成矩阵形式:

$$\begin{pmatrix} \hat{Z}_{n+1}(1) \\ \hat{Z}_{n+1}(2) \\ \vdots \\ \hat{Z}_{n+1}(q) \end{pmatrix} = \begin{pmatrix} \theta_1 & 1 & 0 & \cdots & 0 \\ \theta_2 & 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \theta_{q-1} & 0 & 0 & \cdots & 1 \\ \theta_q & 0 & 0 & \cdots & 0 \end{pmatrix} \begin{pmatrix} \hat{Z}_n(1) \\ \hat{Z}_n(2) \\ \vdots \\ \hat{Z}_n(q) \end{pmatrix} - \begin{pmatrix} \theta_1 \\ \theta_2 \\ \vdots \\ \theta_q \end{pmatrix} X_{n+1}$$

递推时，初值 $\hat{Z}_0(1), \hat{Z}_0(2), \dots, \hat{Z}_0(L)$ 均取为0。

第三部分 时间序列平稳性

时间序列的平稳性

单位根检验

协整分析与误差修正模型

§ 1 时间序列的平稳性及其检验

- (一) 问题的引出
- (二) 时间序列数据的平稳性
- (三) 时间序列模型分类
- (四) 平稳性的图示判断

（一）问题的引出：

- 经典回归分析暗含着一个重要假设：**数据是平稳的。**
- **数据非平稳，大样本下的统计推断基础——“一致性”要求——被破坏。**

满足统计推断中大样本下的“一致性”特性：

$$P\lim_{n \rightarrow \infty}(\hat{\beta}) = \beta$$

▲ **如果X是非平稳数据**（如表现出向上的趋势），则上式不成立，回归估计量不满足“一致性”，基于大样本的统计推断也就遇到麻烦。

◆数据非平稳，往往导致出现“虚假回归”问题

表现在:两个本来没有任何因果关系的变量，
却有很高的相关性（有较高的 R^2 ）:

时间序列分析模型方法就是在这样的情况下，以通过揭示时间序列自身的变化规律为主线而发展起来的全新的计量经济学方法论。

(二) 时间序列数据的平稳性

假定某个时间序列是由某一随机过程生成的，即假定时间序列 $\{X_t\}$ ($t=1, 2, \dots$) 的每一个数值都是从一个概率分布中随机得到，如果满足下列条件：

- 1) 均值 $E(X_t)=\mu$ 是与时间 t 无关的常数；
- 2) 方差 $\text{Var}(X_t)=\sigma^2$ 是与时间 t 无关的常数；
- 3) 协方差 $\text{Cov}(X_t, X_{t+k})=\gamma_k$ 是只与时期间隔 k 有关，与时间 t 无关的常数；

则称该随机时间序列是平稳的，而该随机过程是一平稳随机过程。

下面介绍两种基本的随机过程：

例1. 一个最简单的随机时间序列是一具有零均值同方差的独立分布序列：

$$X_t = \mu_t, \quad \mu_t \sim N(0, \sigma^2)$$

该序列常被称为是一个白噪声（**white noise**）。

由于 X_t 具有相同的均值与方差，且协方差为零，由定义，一个白噪声序列是平稳的。或者说白噪声是平稳的随机过程

例2. 另一个简单的随机时间列序被称为随机游走（**random walk**），该序列由如下随机过程生成：

$$X_t = X_{t-1} + \mu_t$$

这里， μ_t 是一个白噪声。

那么，随机游走是不是平稳的？

容易知道该序列有相同的均值： $E(X_t)=E(X_{t-1})$

为了检验该序列是否具有相同的方差，可假设 X_t 的初值为 X_0 ，则易知

$$X_1=X_0+\mu_1$$

$$X_2=X_1+\mu_2=X_0+\mu_1+\mu_2$$

... ..

$$X_t=X_0+\mu_1+\mu_2+\dots+\mu_t$$

由于 X_0 为常数， μ_t 是一个白噪声，因此 $\text{Var}(X_t)=t\sigma^2$

即 X_t 的方差与时间 t 有关而非常数，它是一非平稳序列。或者说随机游走过程是非平稳的随机过程。

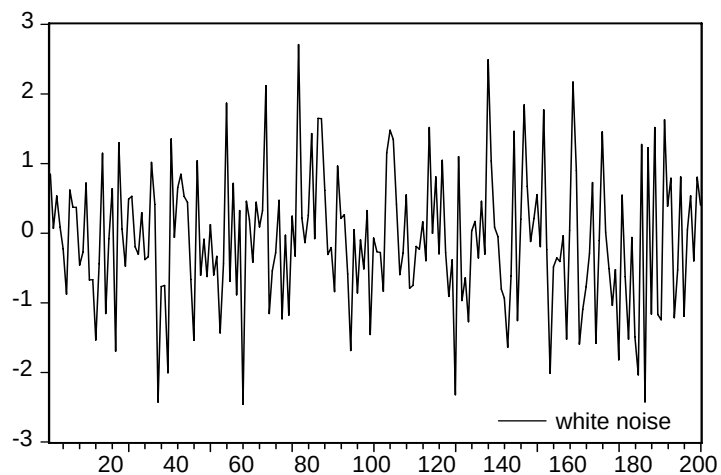


图1：由白噪声过程产生的时间序列

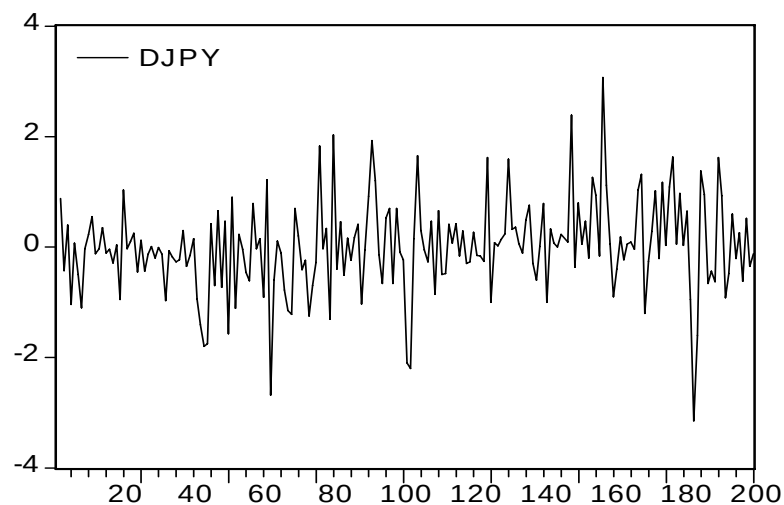


图2：日元对美元汇率的收益率序列

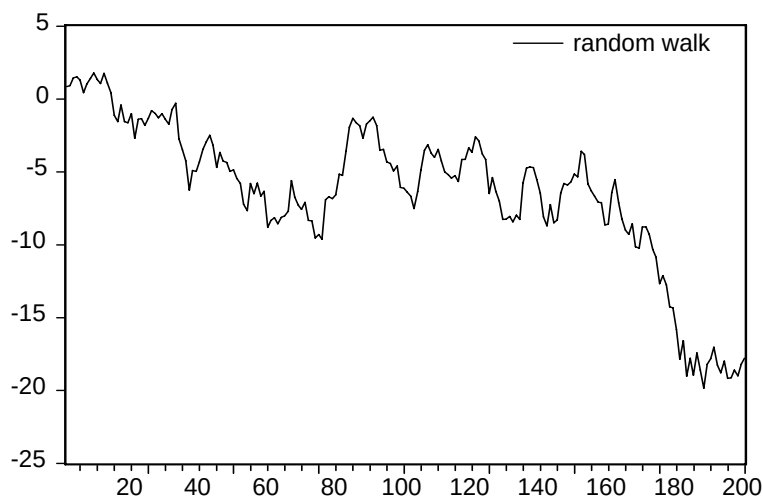


图3：由随机游走过程产生时间序列

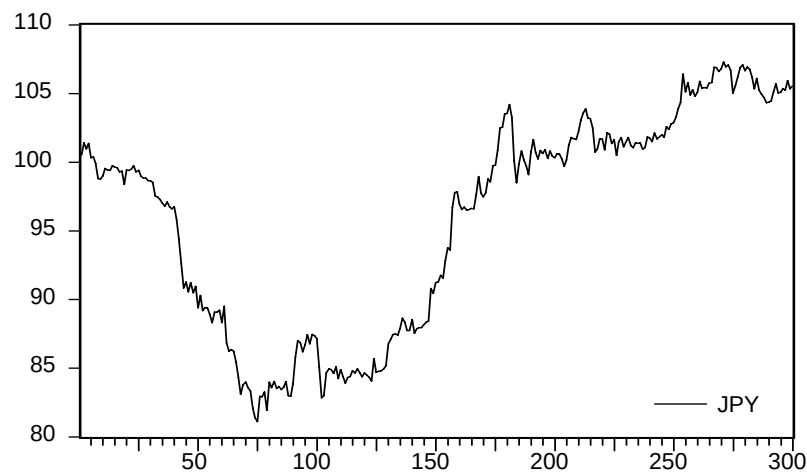


图4：日元对美元汇率（300天，1995年）

- 对随机游走，对 X 取一阶差分：

$$\Delta X_t = X_t - X_{t-1} = \mu_t$$

由于 μ_t 是一个白噪声，则序列 $\{\Delta X_t\}$ 是平稳的。

一般情况下，如果一个时间序列是非平稳的，它常常可通过取差分的方法而形成平稳序列。

（三）时间序列模型的分类

(1) 自回归模型

如果一个线性过程可表达为

$$X_t = \phi_1 X_{t-1} + \phi_2 X_{t-2} + \dots + \phi_p X_{t-p} + u_t,$$

其中 $\phi_i, i = 1, \dots, p$ 是自回归参数， u_t 是白噪声过程，则称 x_t 为 p 阶自回归过程，用AR(p)表示。 X_t 是由它的 p 个滞后变量的加权和以及 u_t 相加而成。

$$X_t - \phi_1 X_{t-1} - \phi_2 X_{t-2} - \dots - \phi_p X_{t-p} = u_t,$$

若用滞后算子表示

$$(1 - \phi_1 L - \phi_2 L^2 - \dots - \phi_p L^p) X_t = \Phi(L) X_t = u_t$$

其中 $\Phi(L) = 1 - \phi_1 L - \phi_2 L^2 - \dots - \phi_p L^p$
称为特征多项式或自回归算子。

与自回归模型常联系在一起的是平稳性问题。

对于自回归过程 **AR(p)**，如果其特征方程

$$\begin{aligned} \Phi(z) &= 1 - \phi_1 z - \phi_2 z^2 - \dots - \phi_p z^p \\ &= (1 - G_1 z)(1 - G_2 z) \dots (1 - G_p z) = 0 \end{aligned}$$

的所有根的绝对值都大于1，则 **AR(p)** 是一个平稳的随机过程。

AR(p) 过程中最常用的是 **AR(1)**、**AR(2)** 过程，

$$X_t = \phi_1 X_{t-1} + u_t$$

保持其平稳性的条件是特征方程

$(1 - \phi_1 L) = 0$ 根的绝对值必须大于1，满足

$$|1/\phi_1| > 1$$

也就是 $|\phi_1| < 1$

下面分析**AR(2)**过程

$x_t = \phi_1 x_{t-1} + \phi_2 x_{t-2} + u_t$ 具有平稳性的条件。

对于**AR(2)** 过程，特征方程式是

$$1 - \phi_1 L - \phi_2 L^2 = 0$$

上式的两个根是：

$$L_1, L_2 = \frac{\phi_1 \pm \sqrt{\phi_1^2 + 4\phi_2}}{-2\phi_2}$$

设 $\lambda_1 = 1/L_1, \lambda_2 = 1/L_2$

$$\lambda_1, \lambda_2 = \frac{-2\phi_2}{\phi_1 \pm \sqrt{\phi_1^2 + 4\phi_2}} = \frac{\phi_1 \mp \sqrt{\phi_1^2 + 4\phi_2}}{2}$$

则 $x_t = \phi_1 x_{t-1} + \phi_2 x_{t-2} + u_t$,

改写为 $(1 - \lambda_1 L)(1 - \lambda_2 L)x_t = u_t$ 。

AR(2)模型具有平稳性的条件是 $|L1| > 1, |L2| > 1$ （在单位圆外）

或 $|\lambda_1| < 1, |\lambda_2| < 1$

(1) 移动平均模型

如果一个线性随机过程可用下式表达

$$\begin{aligned} X_t &= u_t + \theta_1 u_{t-1} + \theta_2 u_{t-2} + \dots + \theta_q u_{t-q} \\ &= (1 + \theta_1 L + \theta_2 L^2 + \dots + \theta_q L^q) u_t = \Theta(L) u_t \end{aligned}$$

其中 $\theta_1, \theta_2, \dots, \theta_q$ 是回归参数， u_t 为白噪声过程，则上式称为 q 阶移动平均过程，记为 $MA(q)$ 。之所以称“移动平均”，是因为 x_t 是由 $q+1$ 个 u_t 和 u_t 滞后项的加权和构造而成。“移动”指 t 的变化，“平均”指加权和。

注：由定义知任何一个 q 阶移动平均过程都是由 $q+1$ 个白噪声变量的加权和组成，所以任何一个移动平均过程都是平稳的。

与移动平均过程相联系的一个重要概念是可逆性。
移动平均过程具有可逆性的条件是特征方程：

$\Theta(z) = (1 + \theta_1 z + \theta_2 z^2 + \dots + \theta_q z^q) = 0$ 的全部根的绝对值必须大于1。

(3) 自回归移动平均模型

由自回归和移动平均两部分共同构成的随机过程称为自回归移动平均过程，记为**ARMA** (p, q)，其中 p, q 分别表示自回归和移动平均部分的最大阶数。

ARMA (p, q) 的一般表达式是：

$$x_t = \phi_1 x_{t-1} + \phi_2 x_{t-2} + \dots + \phi_p x_{t-p} + u_t + \theta_1 u_{t-1} + \theta_2 u_{t-2} + \dots + \theta_q u_{t-q}$$

即

$$(1 - \phi_1 L - \phi_2 L^2 - \dots - \phi_p L^p) x_t = (1 + \theta_1 L + \theta_2 L^2 + \dots + \theta_q L^q) u_t$$

或

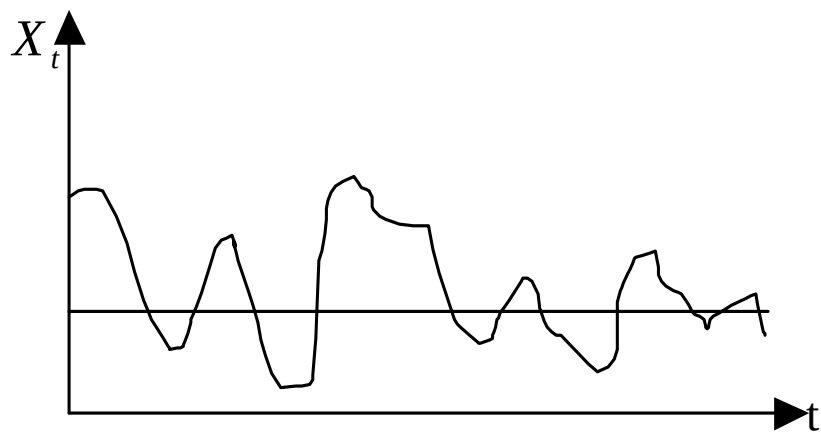
$$\Phi(L) x_t = \Theta(L) u_t$$

其中 $\Phi(L)$ 和 $\Theta(L)$ 分别表示 L 的 p, q 阶特征多项式。

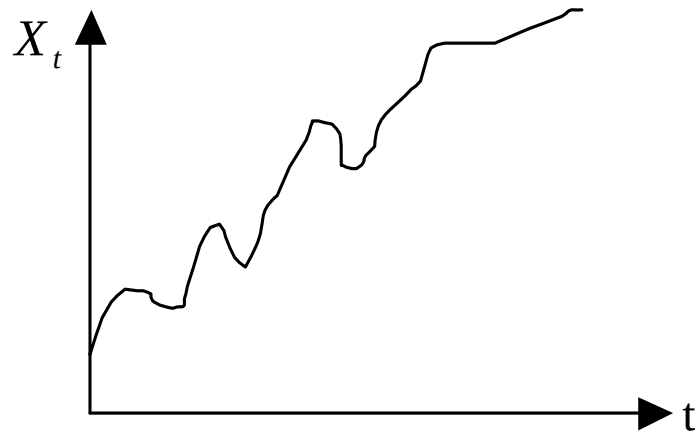
ARMA (p, q) 过程的平稳性只依赖于其自回归部分，即 $\Phi(L) = 0$ 的全部根取值在单位圆之外（绝对值大于1）其可逆性则只依赖于移动平均部分，即 $\Theta(L) = 0$ 的根取值应在单位圆之外。

（四）平稳性检验的图示判断

- 给出一个随机时间序列，首先可通过该序列的**时间路径图**来粗略地判断它是否是平稳的；
- 一个**平稳的时间序列**在图形上往往表现出一种围绕其均值不断波动的过程；
- 而**非平稳序列**则往往表现出在不同的时间段具有不同的均值（如持续上升或持续下降）。



(a)



(b)

图 9.1 平稳时间序列与非平稳时间序列图

图**a** 表示平稳时间序列；图**b** 表示非平稳时间序列

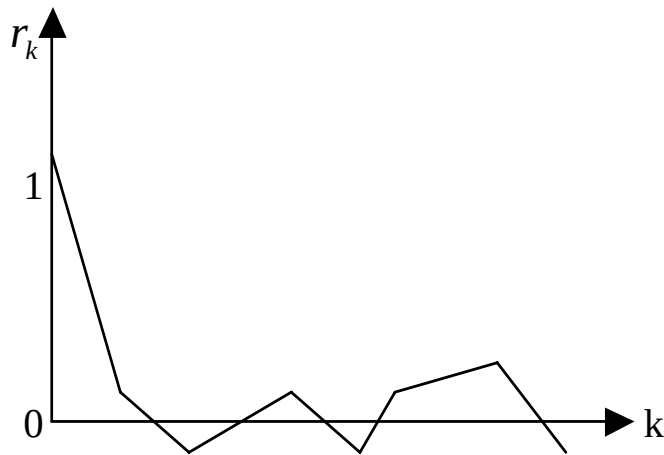
- 进一步的判断:

检验样本自相关函数及其图形

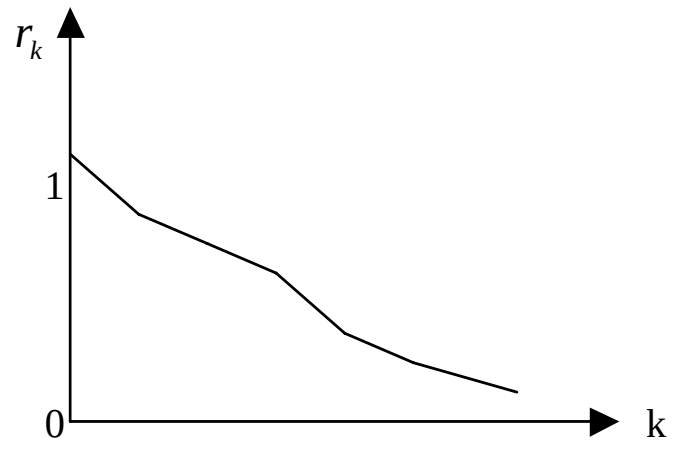
随机时间序列的自相关函数（**autocorrelation function, ACF**） $\rho_k = \gamma_k / \gamma_0$

自相关函数是关于滞后期 k 的递减函数。

易知，随着 k 的增加，样本自相关函数下降且趋于零。
但从下降速度来看，**平稳序列要比非平稳序列快得多。**



(a)



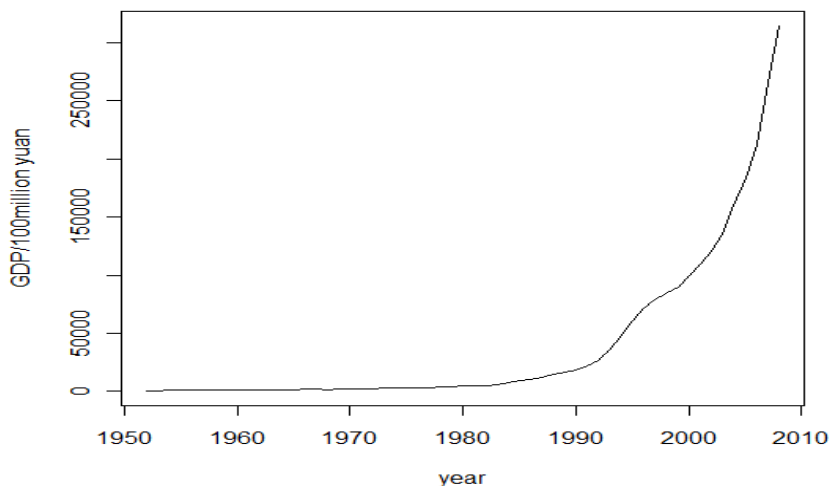
(b)

图 9.1.2 平稳时间序列与非平稳时间序列样本相关图

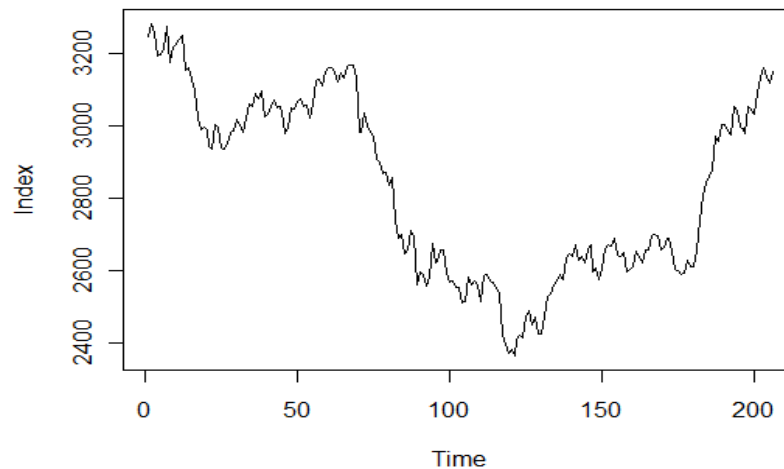
伪回归（虚假回归）

- 回归分析：
- 一个重要的前提假设：平稳性
- 但是，实际上大部分的宏观经济时间序列和金融时间序列都是非平稳的。

GDP of China



Shanghai Composite Index 2010



伪回归（虚假回归）

结果

案例

以1990年至2008年 **美国** 城镇居民家庭人均可支配收入和 **中国** 人均消费性支出为例：

```
> data <- read.table("data.txt")
> Usincome <- data[,1]
> Chinacoms <- data[,2]
> reg <- lm(Chinacoms ~ USincome)
> summary(reg)
> library(zoo)
> library(lmtest)
> dwtest(reg)
```

DW统计量很小，
存在严重的自相关

Call:

lm(formula = Chinacoms ~ USincome)

Residuals:

Min	1Q	Median	3Q	Max
-640.11	-350.44	-55.96	346.49	1139.42

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	-6.886e+03	4.896e+02	-14.06	8.56e-11 ***
USincome	4.788e-01	1.898e-02	25.23	6.54e-15 ***

较大的t
值

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 477.7 on 17 degrees of freedom
Multiple R-squared: 0.974, Adjusted R-squared: 0.9724
F-statistic: 636.3 on 1 and 17 DF, p-value: 6.544e-15

显著的R²

Durbin-Watson test

data: reg

DW = 0.4992, p-value = 4.485e-06

alternative hypothesis: true autocorrelation is greater than 0

线性回归模型

lm()函数:

用法:

```
> fitted.model <- lm(formula, data, subset, weights...)
```

参数:

formula response ~ terms

data 数据框。

subset 取出数据集的一个子集。

weights 权重向量，用于加权最小二乘估计。

例如:

```
> reg <- lm(y ~ x1 + x2, data=output)
```

提取回归模型中所需信息的函数:

coef(reg) resid(reg) plot(reg) summary(reg)

伪回归（虚假回归）

含义：指两个没有因果关系的时间序列之间，基于一些其他的外在因素，推断出因果关系。例如：事件C导致事件A和事件B，如果在A和B之间进行回归分析，则容易推断出A和B之间存在因果关系的错误结论。

特征：

- 1、对参数的检验（t检验）和对回归方程的检验（F检验）容易得到显著的结果，接近于1的 R^2 。
- 2、残差存在严重的正自相关。

结果：

许多非平稳经济变量之间显著的相关性可能并不存在，是虚假的。

§ 2 单位根检验

1. 单整

随机游走序列

$$X_t = X_{t-1} + \mu_t$$

经差分后等价地变形为

$$\Delta X_t = \mu_t$$

由于 μ_t 是一个白噪声，因此差分后的序列 $\{\Delta X_t\}$ 是平稳的。

如果一个时间序列经过一次差分变成平稳的，就称原序列是一阶单整（integrated of 1）序列，记为I(1)。

一般地，如果一个时间序列经过 d 次差分后变成平稳序列，则称原序列是 d 阶单整（integrated of d ）序列，记为 $I(d)$ 。

显然， $I(0)$ 代表一平稳时间序列。

现实经济生活中：

- 1) 只有少数经济指标的时间序列表现为平稳的，如利率等；
- 2) 大多数指标的时间序列是非平稳的，如一些价格指数常常是2阶单整的，以不变价格表示的消费额、收入等常表现为1阶单整。

大多数非平稳的时间序列一般可通过一次或多次差分的形式变为平稳的。

但也有一些时间序列，无论经过多少次差分，都不能变为平稳的。这种序列被称为非单整的（non-integrated）。

2. 确定性趋势和随机性趋势

一些非平稳的经济时间序列往往表现出共同的变化趋势，而这些序列间本身不一定有直接的关联关系，这时对这些数据进行回归，尽管有较高的 R^2 ，但其结果是没有任何实际意义的。这种现象我们称之为**虚假回归或伪回归**（**spurious regression**）。

如：用中国的劳动力时间序列数据与美国**GDP**时间序列作回归，会得到较高的 R^2 ，但不能认为两者有直接的关联关系，而只不过它们有共同的趋势罢了，这种回归结果我们认为是虚假的。

为了避免这种虚假回归的产生，通常的做法是引入作为趋势变量的时间，这样包含有时间趋势变量的回归，可以消除这种趋势性的影响。

然而这种做法，只有当趋势性变量是确定性的（**deterministic**）而非随机性的（**stochastic**），才会是有效的。

换言之，如果一个包含有某种确定性趋势的非平稳时间序列，可以通过引入表示这一确定性趋势的趋势变量，而将确定性趋势分离出来。

那么，什么是确定性趋势？什么是随机性趋势呢？

考虑如下的含有一阶自回归的随机过程：

$$X_t = \alpha + \beta t + \rho X_{t-1} + \mu_t \quad (*)$$

其中： μ_t 是一白噪声， t 为一时间趋势。

1) 如果 $\rho=1$ ， $\beta=0$ ，则 $(*)$ 式成为一带位移的随机游走过程：

$$X_t = \alpha + X_{t-1} + \mu_t \quad (**)$$

根据 α 的正负， X_t 表现出明显的上升或下降趋势，这种趋势称为随机性趋势（stochastic trend）

2) 如果 $\rho=0$ ， $\beta \neq 0$ ，则 $(*)$ 式成为一带时间趋势的随机变化过程：

$$X_t = \alpha + \beta t + \mu_t \quad (***)$$

根据 β 的正负， X_t 表现出明显的上升或下降趋势。这种趋势称为确定性趋势（deterministic trend）。

3) 如果 $\rho=1$, $\beta \neq 0$, 则 X_t 包含有确定性与随机性两种趋势。

判断一个非平稳的时间序列, 它的趋势是随机性的还是确定性的, 可通过**ADF**检验。

(1) 如果检验结果表明所给时间序列有单位根, 且时间变量前的参数显著为零, 则该序列显示出随机性趋势;

(2) 如果没有单位根, 且时间变量前的参数显著地异于零, 则该序列显示出确定性趋势。

3. 单位根检验 (Unit Root Test)

所建的模型里怎么看出有时间趋势？需要检验。

单位根检验是统计检验中普遍应用的一种检验方法。

1、DF检验

我们已知道，随机游走序列

$$X_t = X_{t-1} + \mu_t$$

是非平稳的，其中 μ_t 是白噪声。

而该序列可看成是随机模型

$$X_t = \rho X_{t-1} + \mu_t$$

中参数 $\rho=1$ 时的情形。

也就是说，我们对式

$$X_t = \rho X_{t-1} + \mu_t \quad (*)$$

做回归，如果确实发现 $\rho=1$ ，就说随机变量 X_t 有一个单位根，是非平稳的。

- (*) 式可变形形式成差分形式：

$$\begin{aligned} \Delta X_t &= (\rho - 1)X_{t-1} + \mu_t \\ &= \delta X_{t-1} + \mu_t \end{aligned} \quad (**)$$

检验 (*) 式是否存在单位根 $\rho=1$ ，也可通过
(**) 式判断是否有 $\delta = 0$ 。

一般地：

- 检验一个时间序列 X_t 的平稳性，可通过检验带有截距项的一阶自回归模型

$$X_t = \alpha + \rho X_{t-1} + \mu_t \quad (*)$$

中的参数 ρ 是否小于1。

或者：检验其等价变形式

$$\Delta X_t = \alpha + \delta X_{t-1} + \mu_t \quad (**)$$

中的参数 δ 是否小于0 。

- 因此，针对式 $\Delta X_t = \alpha + \delta X_{t-1} + \mu_t$
我们关心的检验为：**零假设 $H_0: \delta = 0$** 。
备择假设 $H_1: \delta < 0$

上述检验可通过**OLS**法下的**t**检验完成。

然而，在零假设（序列非平稳）下，即使在大样本下**t**统计量也是有偏误的（向下偏倚），通常的**t**检验无法使用。

Dicky和**Fuller**于**1976**年提出了这一情形下**t**统计量服从的分布（这时的**t**统计量称为 **τ 统计量**），即**DF**分布（见表**1.3**）。

由于**t**统计量的向下偏倚性，它呈现围绕小于零值的偏态分布。

表 9.1.3 DF 分布临界值表

显著性水平	样 本 容 量					t分布临界值 (n=∞)
	25	50	100	500	∞	
0.01	-3.75	-3.58	-3.51	-3.44	-3.43	-2.33
0.05	-3.00	-2.93	-2.89	-2.87	-2.86	-1.65
0.10	-2.63	-2.60	-2.58	-2.57	-2.57	-1.28

- 因此，可通过OLS法估计

$$\Delta X_t = \alpha + \delta X_{t-1} + \mu_t$$

并计算t统计量的值，与DF分布表中给定显著性水平下的临界值比较：

如果：**t < 临界值(或 t 的绝对值大于临界值的绝对值)**，则拒绝零假设**H₀: δ = 0**，
认为时间序列不存在单位根，是平稳的。

2、ADF检验

在上述使用

$$\Delta X_t = \alpha + \delta X_{t-1} + \mu_t$$

对时间序列进行平稳性检验中，实际上假定了时间序列是由具有白噪声随机误差项的一阶自回归过程**AR(1)**生成的。

但在实际检验中，时间序列可能由更高阶的自回归过程生成的，或者随机误差项并非白噪声，这样用**OLS**法进行估计均会表现出随机误差项出现自相关（**autocorrelation**），导致**DF**检验无效。

另外，如果时间序列包含有明显的随时间变化的某种趋势（如上升或下降），则也容易导致上述检验中的自相关随机误差项问题。

为了保证**DF**检验中随机误差项的白噪声特性，**Dicky**和**Fuller**对**DF**检验进行了扩充，形成了**ADF**（**Augment Dickey-Fuller**）检验。

ADF检验是通过下面三个模型完成的：

模型 1:
$$\Delta X_t = \delta X_{t-1} + \sum_{i=1}^m \beta_i \Delta X_{t-i} + \varepsilon_t \quad (*)$$

模型 2:
$$\Delta X_t = \alpha + \delta X_{t-1} + \sum_{i=1}^m \beta_i \Delta X_{t-i} + \varepsilon_t \quad (**)$$

模型 3:
$$\Delta X_t = \alpha + \beta t + \delta X_{t-1} + \sum_{i=1}^m \beta_i \Delta X_{t-i} + \varepsilon_t \quad (***)$$

- **模型3** 中的**t**是时间变量，代表了时间序列随时间变化的某种趋势（如果有的话）。
- 检验的假设都是：针对检验 **H0: $\delta=0$** ，即存在一单位根。模型**1**与另两模型的差别在于是否包含有常数项和趋势项。

- 实际检验时从模型**3**开始，然后模型**2**、模型**1**。

何时检验拒绝零假设，即原序列不存在单位根，为平稳序列，何时检验停止。否则，就要继续检验，直到检验完模型**1**为止。

检验原理与**DF**检验相同，只是对模型**1**、**2**、**3**进行检验时，有各自相应的临界值。

表**1.4**给出了三个模型所使用的**ADF**分布临界值表。

表 9.1.4 不同模型使用的 ADF 分布临界值表

模型	统计量	样本容量	0.01	0.025	0.05	0.10
1	τ_{δ}	25	-2.66	-2.26	-1.95	-1.60
		50	-2.62	-2.25	-1.95	-1.61
		100	-2.60	-2.24	-1.95	-1.61
		250	-2.58	-2.23	-1.95	-1.61
		500	-2.58	-2.23	-1.95	-1.61
		>500	-2.58	-2.23	-1.95	-1.61
2	τ_{δ}	25	-3.75	-3.33	-3.00	-2.62
		50	-3.58	-3.22	-2.93	-2.60
		100	-3.51	-3.17	-2.89	-2.58
		250	-3.46	-3.14	-2.88	-2.57
		500	-3.44	-3.13	-2.87	-2.57
		>500	-3.43	-3.12	-2.86	-2.57
	τ_{α}	25	3.41	2.97	2.61	2.20
		50	3.28	2.89	2.56	2.18
		100	3.22	2.86	2.54	2.17
		250	3.19	2.84	2.53	2.16
		500	3.18	2.83	2.52	2.16
		>500	3.18	2.83	2.52	2.16
3	τ_{δ}	25	-4.38	-3.95	-3.60	-3.24
		50	-4.15	-3.80	-3.50	-3.18
		100	-4.04	-3.73	-3.45	-3.15
		250	-3.99	-3.69	-3.43	-3.13
		500	-3.98	-3.68	-3.42	-3.13
		>500	-3.96	-3.66	-3.41	-3.12
	τ_{α}	25	4.05	3.59	3.20	2.77
		50	3.87	3.47	3.14	2.75
		100	3.78	3.42	3.11	2.73
		250	3.74	3.39	3.09	2.73
		500	3.72	3.38	3.08	2.72
		>500	3.71	3.38	3.08	2.72
	τ_{β}	25	3.74	3.25	2.85	2.39
		50	3.60	3.18	2.81	2.38
		100	3.53	3.14	2.79	2.38
		250	3.49	3.12	2.79	2.38
		500	3.48	3.11	2.78	2.38
		>500	3.46	3.11	2.78	2.38

一个简单的检验过程：

同时估计出上述三个模型的适当形式，然后通过**ADF**临界值表检验零假设 $H_0: \delta=0$ 。

这里所谓**模型适当的形式**就是在每个模型中选取适当的滞后差分项，以使模型的残差项是一个白噪声（主要保证不存在自相关）。

- 1) 只要其中有一个模型的检验结果拒绝了零假设，就可以认为时间序列是平稳的；
- 2) 当三个模型的检验结果都不能拒绝零假设时，则认为时间序列是非平稳的。

生成时间序列

stat包中的ts函数：

用法：

```
ts(data = NA, start = 1, end = numeric(0), frequency = 1, deltat = 1, ts.eps = getOption("ts.eps"),  
class = , names = )
```

参数：

data 一个数值向量。

start 时间序列的起始时刻。可以是一个整数，也可以是两个整数组成的向量。

end 时间序列的最后时刻。

frequency 频率，一个单位时间内的观测次数。

names 字符型向量。

时间序列举例

年度数据:

```
>ts(1:30,start=1980,end=2009,  
frequency=1)
```

Time Series:

Start = 1980

End = 2009

Frequency = 1

```
[1] 1 2 3 4 5 6 7 8 9 10 11 12 13 14 15  
[16] 16 17 18 19 20 21 22 23 24 25 26 27 28 29 30
```

季度数据:

```
>ts(1:30,start=c(2003,3),end=  
c(2010,4),frequency=4)
```

	Qtr1	Qtr2	Qtr3	Qtr4
2003			1	2
2004	3	4	5	6
2005	7	8	9	10
2006	11	12	13	14
2007	15	16	17	18
2008	19	20	21	22
2009	23	24	25	26
2010	27	28	29	30

时间序列举例

月度数据:

```
> ts(1:30, start=c(2008,5),  
frequency=12)
```

	Jan	Feb	Mar	Apr	May	Jun	Jul	Aug	Sep	Oct	Nov	Dec	
2008						1	2	3	4	5	6	7	8
2009	9	10	11	12	13	14	15	16	17	18	19	20	
2010	21	22	23	24	25	26	27	28	29	30			

每日数据:

```
> a <-  
ts(1:30, start=c(12,1), frequency=7)  
> print(a, calendar=TRUE)
```

	p1	p2	p3	p4	p5	p6	p7
12	1	2	3	4	5	6	7
13	8	9	10	11	12	13	14
14	15	16	17	18	19	20	21
15	22	23	24	25	26	27	28
16	29	30					

单位根检验

目的：检验时间序列是否存在单位根，即检验时间序列是否平稳。

原假设：序列存在一个单位根。

检验过程：原假设 $\Leftrightarrow \pi = 0$

$$\Delta y_t = \beta_1 + \beta_2 t + \pi y_{t-1} + \sum_{j=1}^k \gamma_j \Delta y_{t-j} + u_t$$

单位根检验

urca包中的ur.df()函数:

用法:

```
ur.df(y, type = c("none", "drift", "trend"), lags = 1, selectlags = c("Fixed", "AIC", "BIC"))
```

参数:

y 被检验的时间序列

type 检验类型: “none”, “drift” 或者 “trend”.

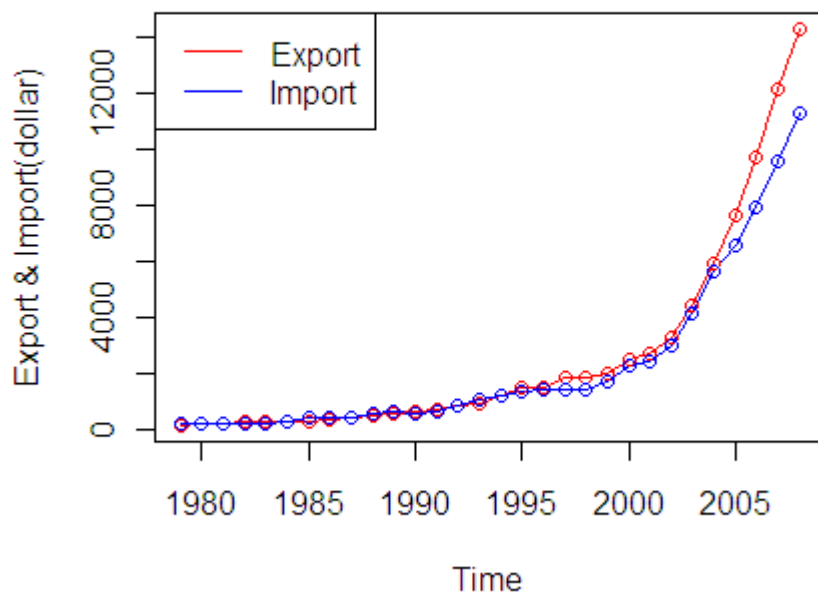
lags 内生变量的滞后阶数

selectlags 滞后阶数确定方法: the Akaike “AIC” 或者 the Bayes “BIC” 信息准则, 默认值是fixed, 由lags确定滞后阶数。

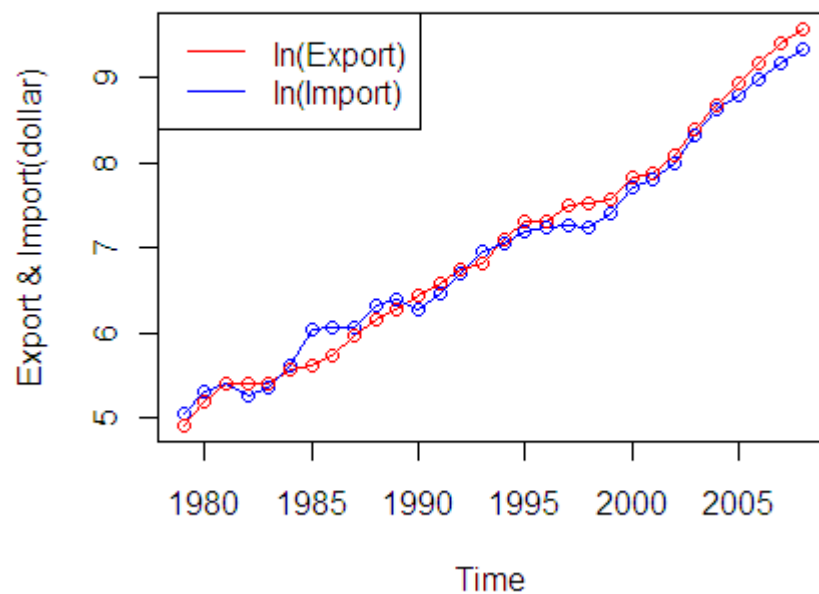
$$\Delta y_t = \beta_1 + \beta_2 t + \pi y_{t-1} + \sum_{j=1}^k \gamma_j \Delta y_{t-j} + u_t$$

案例分析：中国进出口贸易之间关系

Export and Import of China



Export and Import of China



数据来源：国家统计局数据库 (<http://www.stats.gov.cn/tjsj/>)

单位根检验

```
> library(urca)
> urt.ex <-
ur.df(lnex,type='trend',selectlags='AIC')
> summary(urt.ex)

> urt.im <-
ur.df(lnim,type='trend',selectlags='AIC')
> summary(urt.im)
```

Value of test-statistic is: -2.2266 6.0966 3.3771

Critical values for test statistics:

	1pct	5pct	10pct
tau3	-4.15	-3.50	-3.18
phi2	7.02	5.13	4.31
phi3	9.31	6.73	5.61

```
#####
# Augmented Dickey-Fuller Test Unit Root Test #
#####
```

Test regression trend

Call:

lm(formula = z.diff ~ z.lag.1 + 1 + tt + z.diff.lag)

Residuals:

Min	1Q	Median	3Q	Max
-0.17362	-0.05933	0.02236	0.05216	0.12439

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	0.77701	0.53253	1.459	0.158
z.lag.1	-0.15325	0.11600	-1.321	0.199
tt	0.02699	0.01723	1.566	0.130
z.diff.lag	0.20945	0.20765	1.009	0.323

Residual standard error: 0.08564 on 24 degrees of freedom

Multiple R-squared: 0.2354, Adjusted R-squared: 0.1398

F-statistic: 2.462 on 3 and 24 DF, p-value: 0.08697

Value of test-statistic is: -1.3211 7.0028 3.0289

Critical values for test statistics:

	1pct	5pct	10pct
tau3	-4.15	-3.50	-3.18
phi2	7.02	5.13	4.31
phi3	9.31	6.73	5.61

单位根检验

```
dlnex <- diff(lnex)
```

对 *dlnex* 单位根检验结果:

Value of test-statistic is: **-3.2348**
5.2379

Critical values for test statistics:

	1pct	5pct	10pct
tau2	-3.58	-2.93	-2.60
phil	7.06	4.86	3.94

```
dlnim <- diff(lnim)
```

对 *dlnim* 单位根检验结果:

Value of test-statistic is: **-4.8372**
11.7049

Critical values for test statistics:

	1pct	5pct	10pct
tau2	-3.58	-2.93	-2.60
phil	7.06	4.86	3.94

结论: 中国进口和出口的对数时间序列不平稳, 但一阶差分后平稳, 说明是一阶单整序列, 即 $\ln ex, \ln im \sim I(1)$, 满足协整检验条件。

§ 3 协整和误差修正模型

一、协整

(一) 问题的提出

- 经典回归模型是建立在稳定数据变量基础上的，对于非稳定变量，不能使用经典回归模型，否则会出现虚假回归等诸多问题。
- 由于许多经济变量是非稳定的，这就给经典的回归分析方法带来了很大限制。
- 但是，如果变量之间有着长期的稳定关系，即它们之间是协整的 (Cointegration)，则是可以使用经典回归模型方法建立回归模型的。
- 例如，中国居民人均消费水平与人均GDP变量的例子中：
因果关系回归模型要比ARMA模型有更好的预测功能，其原因在于，从经济理论上说，人均GDP决定着居民人均消费水平，而且它们之间有着长期的稳定关系，即它们之间是协整的。

(二) 协整

如果序列 $\{X_{1t}, X_{2t}, \dots, X_{kt}\}$ 都是 **d阶单整**，存在向量 $\alpha = (\alpha_1, \alpha_2, \dots, \alpha_k)$ ，使得 **$Z_t = \alpha X^T \sim I(d-b)$**

其中， **$b > 0$** ， **$X = (X_{1t}, X_{2t}, \dots, X_{kt})^T$** ，则认为序列 $\{X_{1t}, X_{2t}, \dots, X_{kt}\}$ 是 **(d,b)阶协整**，记为 **$X_t \sim CI(d, b)$** ， α 为协整向量（**Cointegrated Vector**）。

两变量协整概念：

若存在 x_t, y_t 都是一阶单整 **$I(1)$** ，如果存在 β 使 **$y_t - \beta x_t = \varepsilon_t \sim I(0)$** 则称 x_t, y_t 为协整。

由此可见：如果两个变量都是单整变量，只有当它们的单整阶数相同时，才可能协整；如果它们的单整阶数不相同，就不可能协整。

从协整的定义可以看出：

(**d,d**) 阶协整是一类非常重要的协整关系，**它的经济意义在于：**两个变量，虽然它们具有各自的长期波动规律，但是如果它们是 (d, d) 阶协整的，则它们之间存在着一个长期稳定的比例关系。

例如：消费**C**和国内生产总值**GDP**，它们各自都是**2**阶单整，**并且它们是(2, 2)阶协整**，说明它们之间存在着一个长期稳定的比例关系，从计量经济学模型的意义讲，建立如下居民人均消费函数模型

$$C_t = \alpha_0 + \alpha_1 GDP_t + \mu_t$$

变量选择是合理的，随机误差项一定是“白噪声”（即均值为**0**，方差不变的稳定随机序列），模型参数有合理的经济解释。**从这里，我们已经初步认识到：**检验变量之间的协整关系，在建立计量经济学模型中是非常重要的。

(三) 协整检验

1、两变量的Engle-Granger检验

为了检验两变量 Y_t, X_t 是否为协整，Engle和Granger于1987年提出两步检验法，也称为EG检验。

第一步，用OLS方法估计方程 $Y_t = \alpha_0 + \alpha_1 X_t + \mu_t$

并计算非均衡误差，得到：

$$\hat{Y}_t = \hat{\alpha}_0 + \hat{\alpha}_1 X_t$$

$$\hat{e}_t = Y_t - \hat{Y}_t$$

称为协整回归(cointegrating)或静态回归(static regression)。

第二步，检验 $\hat{e}_t$ 的单整性。如果 $\hat{e}_t$ 为稳定序列，则认为变量 Y_t, X_t 为(1,1)阶协整；如果 $\hat{e}_t$ 为1阶单整，则认为变量 Y_t, X_t 为(2,1)阶协整；...

$\hat{e}_t$ 的单整性的检验方法仍然是**DF**检验或者**ADF**检验。

由于协整回归中已含有截距项，则检验模型中无需再用截距项。如使用模型1

$$\Delta e_t = \delta e_{t-1} + \sum_{i=1}^p \theta_i \Delta e_{t-i} + \varepsilon_t$$

进行检验时，拒绝零假设 **$H_0: \delta=0$** ，意味着误差项 **e_t** 是平稳序列，从而说明**X与Y**间是协整的。

需要注意的是，这里的**DF**或**ADF**检验是针对协整回归计算出的误差项 $\hat{e}_t$ 而非真正的非均衡误差 μ_t 进行的。而**OLS**法采用了残差最小平方和原理，因此估计量 **δ** 是向下偏倚的，这样将导致拒绝零假设的机会比实际情形大。

于是对 e_t 平稳性检验的**DF**与**ADF**临界值应该比正常的**DF**与**ADF**临界值还要小。

- **MacKinnon(1991)**通过模拟试验给出了协整检验的临界值，表**3.1**是双变量情形下不同样本容量的临界值。

表 9.3.1 双变量协整 ADF 检验临界值

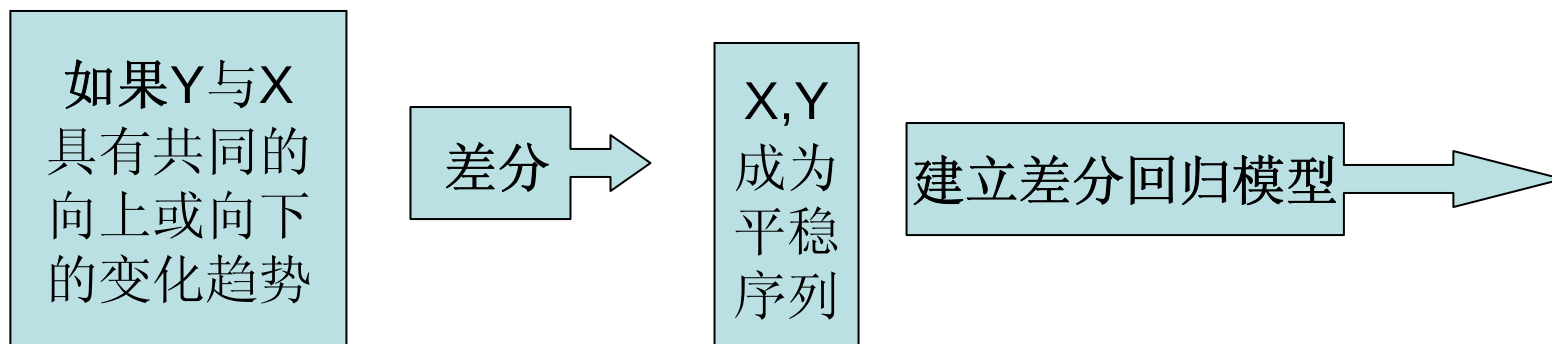
样本容量	显 著 性 水 平		
	0.01	0.05	0.10
25	-4.37	-3.59	-3.22
50	-4.12	-3.46	-3.13
100	-4.01	-3.39	-3.09
∞	-3.90	-3.33	-3.05

二、误差修正模型

对于非稳定时间序列，可通过差分的方法将其化为稳定序列，然后才可建立经典的回归分析模型。

如：建立人均消费水平（Y）与人均可支配收入（X）之间的回归模型：

$$Y_t = \alpha_0 + \alpha_1 X_t + \mu_t$$



$$\Delta Y_t = \alpha_1 \Delta X_t + v_t$$

式中， $v_t = \mu_t - \mu_{t-1}$

然而，这种做法会引起两个问题：

(1)如果X与Y间存在着长期稳定的均衡关系

$$Y_t = \alpha_0 + \alpha_1 X_t + \mu_t$$

且误差项 μ_t 不存在序列相关，则差分式

$$\Delta Y_t = \alpha_1 \Delta X_t + v_t$$

中的 v_t 是一个一阶移动平均时间序列，因而是序列相关的；

(2)如果采用差分形式进行估计，则关于变量水平值的重要信息将被忽略，这时模型只表达了X与Y间的短期关系，而没有揭示它们间的长期关系。

因为，从长期均衡的观点看，Y在第t期的变化不仅取决于X本身的变化，还取决于X与Y在t-1期末的状态，尤其是X与Y在t-1期的不平衡程度。

另外，使用差分变量也往往会得出不能令人满意回归方程。

例如，使用 $\Delta Y_t = \alpha_1 \Delta X_t + v_t$ 回归时，很少出现截距项显著为零的情况，即我们常常会得到如下形式的方程：

$$\Delta Y_t = \hat{\alpha}_0 + \hat{\alpha}_1 \Delta X_t + v_t \quad \hat{\alpha}_0 \neq 0 \quad (*)$$

在 X 保持不变时，如果模型存在静态均衡（**static equilibrium**）， Y 也会保持它的长期均衡值不变。

但如果使用（*）式，即使 X 保持不变， Y 也会处于长期上升或下降的过程中，这意味着 X 与 Y 间不存在静态均衡。

这与大多数具有静态均衡的经济理论假说不相符。

可见，简单差分不一定能解决非平稳时间序列所遇到的全部问题，因此，误差修正模型便应运而生。

误差修正模型（Error Correction Model，简记为ECM）是一种具有特定形式的计量经济学模型，它的主要形式是由**Davidson、Hendry、Srba和Yeo**于**1978**年提出的，称为**DHSY模型**。

为了便于理解，我们通过一个具体的模型来介绍它的结构。

假设两变量X与Y的长期均衡关系为：

$$Y_t = \alpha_0 + \alpha_1 X_t + \mu_t$$

由于现实经济中X与Y很少处在均衡点上，因此实际观测到的只是X与Y间的短期的或非均衡的关系，假设具有如下(1, 1)阶分布滞后形式

$$Y_t = \beta_0 + \beta_1 X_t + \beta_2 X_{t-1} + \mu Y_{t-1} + \varepsilon_t$$

该模型显示出第t期的Y值，不仅与X的变化有关，而且与t-1期X与Y的状态值有关。

由于变量可能是非平稳的，因此不能直接运用**OLS**法。
对上述**分布滞后模型**适当变形得

$$\begin{aligned}\Delta Y_t &= \beta_0 + \beta_1 \Delta X_t + (\beta_1 + \beta_2) X_{t-1} - (1 - \mu) Y_{t-1} + \varepsilon_t \\ &= \beta_1 \Delta X_t - (1 - \mu) \left(Y_{t-1} - \frac{\beta_0}{1 - \mu} - \frac{\beta_1 + \beta_2}{1 - \mu} X_{t-1} \right) + \varepsilon_t\end{aligned}$$

或

$$\Delta Y_t = \beta_1 \Delta X_t - \lambda (Y_{t-1} - \alpha_0 - \alpha_1 X_{t-1}) + \varepsilon_t \quad (**)$$

$$\lambda = 1 - \mu \quad \alpha_0 = \beta_0 / (1 - \mu) \quad \alpha_1 = (\beta_1 + \beta_2) / (1 - \mu)$$

式中如果将（**）中的参数，与 $Y_t = \alpha_0 + \alpha_1 X_t + \mu_t$ 中的相应参数视为相等，则（**）式中括号内的项就是t-1期的非均衡误差项。

（）式表明：**Y的变化决定于X的变化以及前一时期的非均衡程度。同时，（**）式也弥补了简单差分模型 $\Delta Y_t = \alpha_1 \Delta X_t + v_t$ 的不足，因为该式含有用X、Y水平值表示的前期非均衡程度。因此，Y的值已对前期的非均衡程度作出了修正。

$$\Delta Y_t = \beta_1 \Delta X_t - \lambda(Y_{t-1} - \alpha_0 - \alpha_1 X_{t-1}) + \varepsilon_t \quad (**)$$

称为一阶误差修正模型(**first-order error correction model**)。

(**) 式可以写成:

$$\Delta Y_t = \beta_1 \Delta X_t - \lambda ecm + \varepsilon_t \quad (***)$$

其中: **ecm**表示误差修正项。

由分布滞后模型

$$Y_t = \beta_0 + \beta_1 X_t + \beta_2 X_{t-1} + \mu Y_{t-1} + \varepsilon_t$$

知, 一般情况下 $|\mu| < 1$, 由关系式 $\lambda = 1 - \mu$ 得 $0 < \lambda < 1$ 。可以据此分析 **ecm** 的修正作用:

(1) 若 (t-1) 时刻 Y 大于其长期均衡解 $\alpha_0 + \alpha_1 X$, **ecm** 为正, 则 $(-\lambda ecm)$ 为负, 使得 ΔY_t 减少;

(2) 若 (t-1) 时刻 Y 小于其长期均衡解 $\alpha_0 + \alpha_1 X$, **ecm** 为负, 则 $(-\lambda ecm)$ 为正, 使得 ΔY_t 增大。

需要注意的是：在实际分析中，变量常以对数的形式出现。

其主要原因在于变量对数的差分近似地等于该变量的变化率，而经济变量的变化率常常是稳定序列，因此适合于包含在经典回归方程中。

于是：(1) 长期均衡模型

$$Y_t = \alpha_0 + \alpha_1 X_t + \mu_t$$

中的 α_1 可视为Y关于X的长期弹性（long-run elasticity）

(2) 短期非均衡模型

$$Y_t = \beta_0 + \beta_1 X_t + \beta_2 X_{t-1} + \mu Y_{t-1} + \varepsilon_t$$

中的 β_1 可视为Y关于X的短期弹性（short-run elasticity）。

1. 误差修正模型的建立

(1) Granger 表述定理

误差修正模型有许多明显的优点：如

a) 一阶差分项的使用消除了变量可能存在的趋势因素，从而避免了虚假回归问题；

b) 一阶差分项的使用也消除模型可能存在的多重共线性问题；

c) 误差修正项的引入保证了变量水平值的信息没有被忽视；

d) 由于误差修正项本身的平稳性，使得该模型可以使用经典的回归方法进行估计，尤其是模型中差分项可以使用通常的t检验与F检验来进行选取；等等。

因此，一个重要的问题就是：是否变量间的关系都可以通过误差修正模型来表述？

就此问题，Engle 与 Granger 1987年提出了著名的Grange表述定理（Granger representaion theorem）：

如果变量X与Y是协整的，则它们间的短期非均衡关系总能由一个误差修正模型表述：

$$\Delta Y_t = lagged(\Delta Y, \Delta X) - \lambda \mu_{t-1} + \varepsilon_t \quad 0 < \lambda < 1 \quad (*)$$

式中， μ_{t-1} 是非均衡误差项或者说成是长期均衡偏差项， λ 是短期调整参数。

对于(1,1)阶自回归分布滞后模型

$$Y_t = \beta_0 + \beta_1 X_t + \beta_2 X_{t-1} + \mu Y_{t-1} + \varepsilon_t$$

如果 $Y_t \sim I(1)$, $X_t \sim I(1)$ ；那么

$$\Delta Y_t = \beta_1 \Delta X_t - \lambda (Y_{t-1} - \alpha_0 - \alpha_1 X_{t-1}) + \varepsilon_t$$

的左边 $\Delta Y_t \sim I(0)$ ，右边的 $\Delta X_t \sim I(0)$ ，因此，只有Y与X协整，才能保证右边也是 $I(0)$ 。

因此，建立误差修正模型，需要

首先对变量进行协整分析，以发现变量之间的协整关系，即长期均衡关系，并以这种关系构成误差修正项。

然后建立短期模型，将误差修正项看作一个解释变量，连同其它反映短期波动的解释变量一起，建立短期模型，即误差修正模型。

注意，由于

$$\Delta Y = \text{lagged}(\Delta Y, \Delta X) + \lambda \mu_{t-1} + \varepsilon_t \quad 0 < \lambda < 1$$

中没有明确指出Y与X的滞后项数，因此，可以是多个；同时，由于一阶差分项是I(0)变量，因此模型中也允许使用X的非滞后差分项 ΔX_t 。

Granger表述定理可类似地推广到多个变量的情形中去。

(2) Engle-Granger两步法

由协整与误差修正模型的关系，可以得到误差修正模型建立的**E-G**两步法：

第一步，进行协整回归（**OLS**法），检验变量间的协整关系，估计协整向量（长期均衡关系参数）；

第二步，若协整性存在，则以第一步求到的残差作为非均衡误差项加入到误差修正模型中，并用**OLS**法估计相应参数。

需要注意的是：在进行变量间的协整检验时，如有必要可在协整回归式中加入趋势项，这时，对残差项的稳定性检验就无须再设趋势项。

另外，第二步中变量差分滞后项的多少，可以残差项序列是否存在自相关性来判断，如果存在自相关，则应加入变量差分的滞后项。

(3) 直接估计法

也可以采用打开误差修整模型中非均衡误差项括号的方法直接用**OLS**法估计模型。

但仍需事先对变量间的协整关系进行检验。

如对双变量误差修正模型

$$\Delta Y_t = \beta_1 \Delta X_t - \lambda(Y_{t-1} - \alpha_0 - \alpha_1 X_{t-1}) + \varepsilon_t$$

可打开非均衡误差项的括号直接估计下式：

$$\Delta Y_t = \lambda \alpha_0 + \beta_1 \Delta X_t - \lambda Y_{t-1} + \lambda \alpha_1 X_{t-1} + \varepsilon_t$$

这时短期弹性与长期弹性可一并获得。

需注意的是，用不同方法建立的误差修正模型结果也往往不一样。

(β_1 是 y_t 关于 x_t 的短期弹性； α 是 y_t 关于 x_t 的长期弹性)

传统的解决方法

传统方法

一阶差分后进行回归

移除线性趋势

缺点

- 1. 经济理论往往研究的是变量的水平值而不是差分值，差分后的模型不好解释
 - 2. 丢失一些有用的长期信息
-
- 1. 假设序列存在独立的确定性趋势
 - 2. 只能解释变量之间的短期关系

协整

协整

定义：对于两个非平稳时间序列 $\{X_t\}$ 和 $\{Y_t\}$ 如果

① $\{X_t\}$ 和 $\{Y_t\}$ 为 $I(1)$ 序列；

②存在线性组合 $X_t + bY_t$ 使得 $\{X_t + bY_t\}$ 是平稳序列；

则称 $\{X_t\}$ 与 $\{Y_t\}$ 之间具有协整关系。

描述了时间序列之间的 **长期均衡** 关系。

误差修正模型

定义：时间序列 $\{X_t\}$ 和 $\{Y_t\}$ 的误差修正模型表示为：

$$\Delta y_t = \psi_0 + \gamma_1 \hat{z}_{t-1} + \sum_{i=1}^K \psi_{1,i} \Delta x_{t-i} + \sum_{i=1}^L \psi_{2,i} \Delta y_{t-i} + \varepsilon_{1,t},$$

其中 $\{\varepsilon_t\}$ 是平稳随机序列， $\{z_{t-1}\}$ 是误差修正项。

描述了时间序列之间的 **短期波动** 关系。

协整检验

协整检验前提条件：多个时间序列必须是同阶单整的。

协整检验方法：

1. Engle-Granger两步检验法
2. Johansen检验法

Engle-Granger两步检验法：

1. 用OLS估计法对X，Y进行回归估计，得到残差序列 $\{z_t\}$ ，检验残差的平稳性。
2. 根据格兰杰表述定理建立误差修正模型。

EG两步协整检验：第一步（回归方程估计）

```
> exim <- read.table("export and import +of  
China.txt",header=TRUE)  
> ex <- ts(exim[,1],start=1979,end=2008)  
> im <- ts(exim[,2],start=1979,end=2008)  
> lnim <- log(im)  
> lnex <- log(ex)  
> reg <- lm(lnex~lnim)  
> summary(reg)  
> library(lmtest)  
> dw <- dwtest(reg)
```

Durbin-Watson test

data: reg

DW = 0.9589, p-value = 0.0004282

alternative hypothesis: true autocorrelation is
greater than 0

Call:

lm(formula = lnex ~ lnim)

Residuals:

Min	1Q	Median	3Q	Max
-0.40306	-0.05373	0.01223	0.07645	0.23619

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	-0.48270	0.14717	-3.28	0.00278 **
lnim	1.07457	0.02077	51.74	< 2e-16 ***

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 0.1401 on 28 degrees of
freedom

Multiple R-squared: 0.9897, Adjusted R-squared:
0.9893

F-statistic: 2677 on 1 and 28 DF, p-value: < 2.2e-16

协整回归方程: $\ln ex = -0.4827 + 1.0746 \ln im + \varepsilon_t$

EG两步协整检验：第一步（残差单位根检验）

```
> error <- residuals(reg)
> urt.resid <- ur.df(error,type='none',selectlags='AIC')
> summary(urt.resid)
```

Value of test-statistic is: -3.6185

Critical values for test statistics:

	1pct	5pct	10pct
tau1	-2.62	-1.95	-1.61

结论：残差平稳，说明两个时间序列之间存在协整关系。
意味着我国的进口和出口之间具有长期均衡关系，增长或者减少具有协同效应。

EG两步协整检验：第二步（误差修正模型的建立）

```
> error <- residuals(reg)
> error.lagged <- error[-c(29,30)]
> ecm.reg1 <- lm(dlnex~error.lagged+dlnim)
> summary(ecm.reg1)
> dwtest(ecm.reg1)
```

$$\Delta y_t = \psi_0 + \gamma_1 \hat{z}_{t-1} + \sum_{i=1}^K \psi_{1,i} \Delta x_{t-i} + \sum_{i=1}^L \psi_{2,i} \Delta y_{t-i} + \varepsilon_{1,t},$$

Durbin-Watson test

data: ecm.reg1

DW = 2.408, p-value = 0.827

alternative hypothesis: true autocorrelation is greater than 0

Call:

```
lm(formula = dlnex ~ error.lagged + dlnim, data = diff.dat)
```

Residuals:

Min	1Q	Median	3Q	Max
-0.17214	-0.05198	0.01546	0.05053	0.14514

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)	
(Intercept)	0.10648	0.02358	4.516	0.000131	***
error.lagged	-0.29647	0.11430	-2.594	0.015645	*
dlnim	0.33929	0.12470	2.721	0.011681	*

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 0.08024 on 25 degrees of freedom

Multiple R-squared: 0.3008, Adjusted R-squared: 0.2448

F-statistic: 5.377 on 2 and 25 DF, p-value: 0.01142

EG两步协整检验： 第二步（误差修正模型的建立）

误差修正模型：

$$\Delta \ln ex_t = 0.1065 - 0.2965 ecm_{t-1} + 0.3393 \Delta \ln im_t + \varepsilon_t$$

结论：

误差修正项的系数为负，符合误差修正机制，反映了上一期偏离长期均衡的数量将在下一期得到30%的反向修正，这也符合之前证明的协整关系。

总结

中国进出口之间的长期均衡关系：

$$\ln ex_t = -0.4827 + 1.0745 \ln im_t$$

短期波动关系：

$$\Delta \ln ex_t = 0.1065 - 0.2965 ecm_{t-1} + 0.3393 \Delta \ln im_t + \varepsilon_t$$