

金融计量学介绍

本章知识点

1. 了解金融计量学的基本概念与建模步骤。
2. 深入理解金融数据的主要类型、特点和来源。
3. 熟悉收益率的计算方法。
4. 了解统计学与概率的基础知识。

1.1 金融计量学的含义及建模步骤

1.1.1 金融计量学的含义

要了解什么是金融计量学,先要了解什么是计量经济学。计量经济学从字面意思来看,是指经济学中的测量。实质上,计量经济学起源于经济学,是经济学的一个分支学科,是以揭示经济活动中客观存在的数量关系为内容的分支学科。

计量经济学不仅是检验经济理论的一门科学,而且也是预测经济变量未来值的一套工具,如预测股票价格走势、预测公司的销售额等。计量经济学可以利用经济数学模型拟合实际数据的过程,是基于历史数据给予政府和企业以数值化或定量化政策建议的一门科学和艺术。

一般认为金融计量学是计量经济学的一个分支,是计量经济学在金融学中的一种应用。在西方经济中,一般认为金融计量学是指金融市场的计量分析,特别是统计技术在处理金融问题中的应用。而本书所指的金融计量学除了包括以上计量经济学基本理论,还涵盖了其在金融市场的应用,如检验金融市场信息是否有效,检验资本资产定价模型(CAPM)是否决定风险资产收益率的优良模型,测量和预测债券收益率的波动性,检验关于变量相互关系的假设,考察经济状况变化对金融市场的影响,等等,后面相应章节将做详细介绍。

1.1.2 金融计量建模步骤

通过各种不同的方法对金融市场进行描述和模拟就形成了各种各样的金融计量模型,这些模型都具有一些相同的主要步骤。根据克里斯·布鲁克斯(Chris Brooks)提出的金融模型步骤,对金融计量建模的基本步骤总结如图 1-1 所示。

步骤一:问题的概述。该步骤主要涉及金融或经济理论的形成,一般来自某种理论的认识或对某种理论的假设,根据理论建立模型用数学公式表示出来。其具体包括选择变量、确定变量之间的数学关系、拟定模型估计参数的数值范围。选择变量需要注意以下几点:

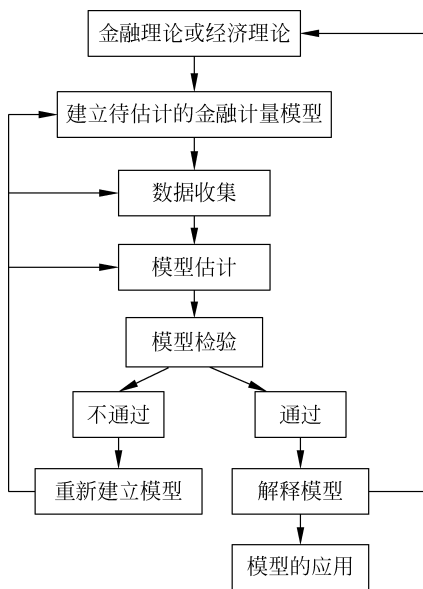


图 1-1 金融计量建模的基本步骤

①正确理解和把握所研究的经济现象中暗含的经济学理论和经济行为规律；②处在不同的环境中，要研究不同的行业，选择的变量也不同；③要考虑数据的可得性；④选择变量要考虑所有入选变量之间的关系，使得每一个解释变量都是独立的。

步骤二：收集样本数据。样本数据的收集与整理，是建立金融计量学模型过程中最基础的工作，也是对模型质量影响极大的一项工作。从工作程序上讲，它是在模型建立之后进行，但实际上经常是同时进行的，因为能否收集到合适的样本观测值是决定变量取舍的主要因素之一。对于金融数据的类型、特点、来源将在 1.2 节详细介绍。

步骤三：选择合适的估计方法。根据模型提出的假设和建立的数学表达式，确定数据的类型选择一元回归还是多元回归，选择单一方程还是联立方程。

步骤四：模型的检验。在步骤三之后初步得到估计结果，但需要进一步检验估计的结果是否满足我们的需要，是否合理地描述数据，是否具有经济学上的意义。一般检验包括三个方面：统计检验、计量经济学检验以及经济、金融意义检验。统计检验的目的在于检验模型参数估计值的可靠性，包括模型拟合优度检验、变量显著性检验、方程显著性检验等；计量经济学检验是否符合计量经济学理论知识，包括序列相关性检验、异方差检验、多重共线性检验以及协整检验等。经济、金融意义检验是将计量检验的结果与相应的经济理论或金融理论相比较，确定两者是否相符。估计的参数通过检验则可进行模型的解释，并应用在实际中，否则回到步骤一到步骤三，要么放宽假设条件，要么改变计算方法，要么重新建立模型，收集更多数据。

步骤五：模型的应用。当模型通过检验，结果获得合理的解释，步骤一的理论得到支撑，就可以将模型用来检验步骤一提出的理论、进行预测或者提出建议。金融计量学模型的应用很广，主要分为以下几个方面：①结构分析。对经济现象中变量之间相互关系进行研究，如研究某一因素变化对经济状况的影响。②金融、经济预测。根据模型变量的设定，预测金融经济变量，如一元回归模型中被解释变量的预测、未来资产价值以及波动率的预测。③政策评价。研究不同的经济政策产生的不同结果，评估这些结果，或是从不同的政策结果中选择最优的政策方案。④检验与发展经济理论。实践是检验真理的唯一标准，任何经济理论都是要通过检验，金融计量学模型是一种很好的科学方法，对于探索经济规律提供了很大的帮助。

1.2 金融数据的主要类型、特点和来源

1.2.1 金融数据的主要类型

在金融问题的分析中，主要有三类数据可供使用，即时间序列数据(time series data)、截面数据(cross-sectional data)、面板数据(panel data)。

(1) 时间序列数据。时间序列数据是指一个实体在不同时期内的观测数据。例如,中国各年的通货膨胀率,一家公司当年每个季度的营业额,每天的股票价格,等等。时间序列数据进行回归时应注意数据的一致性,确保各变量数据的频率相同,应注意模型随机误差项有可能产生序列相关,还应注意数据的平稳性问题,同时也要注意数据的可比性,消除不同时点的数据受通货膨胀因素的影响。时间序列数据可用于研究变量随时间推移的发展变化,并预测这些变量的未来值。如研究某个国家股票指数价格如何随着国家宏观经济基础变量的变化而变化、贸易赤字上升对该国汇率的影响、预测某股票的波动率等。

(2) 截面数据。截面数据是指多个实体在某一时点上的观测数据。例如亚洲各个国家2015年的人均GDP(国内生产总值),某一时点上深圳证券交易所所有股票的收益率,中国各个省份2015年一年的税收情况等。分析截面数据时,可能出现异方差情况,整理数据时应注意消除异方差。这类数据主要用于研究某一时期内不同的人、公司或者其他经济实体之间的差异,从而了解变量之间的关系。例如研究公司规模对其进行股票投资的回报率之间的关系、一国GDP水平与其主权债务违约率之间的关系。

(3) 面板数据。面板数据是指时间序列数据和截面数据相结合的数据,即多个实体在多个时期内的观测数据。例如中国国内所有银行过去3年的贷款数据、所有蓝筹股2010年到2015年每日收盘价等。可以应用面板数据研究不同实体的经历和根据每个实体的变量随时间变化的发展来了解经济关系。

对于以上数据的统计和建模分析大多采用的是点值模型(Point-valued Model)。随着大数据时代的到来,统计学家和计量经济学家开始采用区间模型(Interval-valued Model)对相关数据进行分析。本书不对区间模型进行介绍,有兴趣的读者可参阅相关文献。

1.2.2 金融数据的特点

金融计量学主要研究计量经济学在金融市场中的应用。相比宏观经济数据,金融数据在频率、准确性、周期性等方面具有自己特有的性质。下面选择金融数据关注的两个重要指标即抽样频率和时间跨度逐一介绍。

金融数据可以是低频的、高频的和超高频的,随着技术的进步,研究的数据由以前的月、周、日到现在的10分钟、5分钟、1分钟,可用于计量的数据巨大,这是宏观经济数据不能比拟的。

与宏观经济数据相比,金融数据统计错误和数据修正问题会较少,宏观经济数据通常是测算和估计出来的,难免会有差错,而且新公布的统计数据也有可能出错,都会给计量分析带来困难。而金融数据虽然形式多样,但一般来说价格和其他金融变量都是在交易时准确记录下来的,当然也存在数据测量错误、记录错误的可能性。但总体上,金融学中的误差和修正问题不像经济学中的那么严重。

金融数据特别是时间序列数据,一般都是不平稳的,较难区分是随机游走、趋势或其他特征。金融数据通常还有许多其他特征,特别是高频数据,包含太多的噪声,很难分离出背后的趋势;另外,金融数据大部分不服从正态分布,一般回归方法很难适用,但这也推动了金融经济方法的发展。

1.2.3 金融数据的来源

金融数据主要来源于试验或者对现实世界的实际观测,有以下三个渠道。

(1) 政府部门和国际组织的出版物及网站,如可以从国家统计局网站([www. stats. gov. cn](http://www.stats.gov.cn))获得宏观金融数据,可以从世界银行网站(www. worldbank. org)或国际货币基金组织网站(www. imf. org)获得世界各国的经济数据。

(2) 专业数据公司和信息公司。一些具体公司的市盈率、资产、负债数据需要到特定的公司网站上面下载,有些公司只有本公司网站才会提供这些数据。另外,还有一些信息公司通过收集某方面的数据,建立专业型数据,通过有偿方式来满足客户的需要。当然,也可以在股票软件,如同花顺等,下载不同股票的收盘价、最高(低)价、换手率等方面的交易数据。

(3) 抽样调查。如果分析人员需要一些特定的数据,往往只能通过抽样调查来获得。特别是大数据,很难对总体数据进行分析,或者无法获得总体数据,只能通过抽样来估计总体。例如要对中国股市的投资者信心进行建模,就必须通过设置调查问卷,对不同的投资群体进行数据采集。

1.3 收益率计算

1.3.1 单期收益率

金融许多问题都是从价格的时间序列开始的,如股票每天的收盘价、黄金期货每日价格等。在金融计量上,用得较多的是把价格时间序列转化为收益率时间序列,因为收益率序列统计特性良好,而且收益率还具有无量纲单位的优点,如年收益率为10%,那么投资者投资100元一年后就会得到110元,投资1000元就会得到1100元。

计算收益率一般有两种方法:简单收益率(simple return)和连续复合收益率(continuous compounding return),计算方法如下。

简单收益率计算公式:

$$R_t = \frac{P_t - P_{t-1}}{P_{t-1}} \times 100\% \quad (1.1)$$

连续复合收益率计算公式:

$$r_t = \ln\left(\frac{P_t}{P_{t-1}}\right) \times 100\% \quad (1.2)$$

式中, R_t 为在 t 时期的简单收益率; r_t 为在 t 时期的连续复合收益率; P_t 为在 t 时期的资产价格; $\ln$ 为自然对数。

这里计算的收益率指的都是单期收益率,连续复合收益率也称对数收益率。

根据简单收益率和连续复合收益率的定义,可以根据泰勒公式一阶展开得出如下结论:当简单收益率的值接近零时,连续复合收益率与简单收益率几乎是相等的。也就是求极限时,若样本数据的时间间隔越来越小,简单收益率与连续复合收益率将趋于一致。

1.3.2 多期收益率

多期收益率指的是间隔多个时间段求得的收益率。

多期(假设有 k 期)的简单收益率计算公式如下:

$$R_t(k) = \frac{P_t - P_{t-k}}{P_{t-k}} \times 100\% \quad (1.3)$$

多期的连续复合收益率计算公式如下:

$$\begin{aligned} r_t(k) &= \ln\left(\frac{P_t}{P_{t-k}}\right) \\ &= \ln\left[\left(\frac{P_t}{P_{t-1}}\right)\left(\frac{P_{t-1}}{P_{t-2}}\right)\cdots\left(\frac{P_{t-(k-1)}}{P_{t-k}}\right)\right] \\ &= r_t + r_{t-1} + \cdots + r_{t-k+1} \end{aligned} \quad (1.4)$$

式中, $R_t(k)$ 为 k 期的多期简单收益率; $r_t(k)$ 为 k 期的多期连续复合收益率。

多期收益率可以由单期收益率计算得到,如1年有12个交易月,若采用简单收益率,则年总收益率可以由这12个月总收益率的乘积得到;若采用对数收益率,则可以由这12个月对数收益率求和计算得到。

单期与多期是相对的,比如,在用月收益率计算年收益率时,月收益率可被视作单期收益率,而在用日收益率计算月收益率时,月收益率则被视作多期收益率。这表明单期与多期是相对的,既可以将多期视作一个大的单期,也可以将一个单期分解成包括若干个小单期的多期。

从式(1.4)可以看出,对数收益率具有可加性,给计算带来了很大的方便;同时连续复合收益率有良好的性质,在比较资产间的收益率时,通过连续复合收益率的计算,可以消除不同频率的影响。所以金融学术文献上大都采用连续复利收益率公式计算收益率。

对于简单收益率而言,大部分金融资产以有限的负债形式表现,即投资者的最大损失仅仅是他的全部投资,这就意味着可获取的最小收益率是-100%,但是目前国际上常用的描述收益率序列特征的分布函数(如传统的正态分布)所需要的支撑,往往是整个实数轴。因此,简单收益率的有限负债性给金融建模带来了一定的不便。而对数收益率是对简单总收益率取对数得到,它的取值范围为整个实数轴,所以可以不受有限负债的限制。

近年来,年化收益率在金融市场中较为常见。例如,某个6个月的理财产品“年化收益率为4.8%”,其意思是半年期到期收益率对应的是2.4%。对于季节价格指数,年化收益率计算公式近似写成

$$100\% \times \ln\left(\frac{P_t}{P_{t-1}}\right)^4 = 400\% \times \ln\left(\frac{P_t}{P_{t-1}}\right) \quad (1.5)$$

以此类推,对于月度数据,年化收益率计算公式如下:

$$100\% \times \ln\left(\frac{P_t}{P_{t-1}}\right)^{12} = 1200\% \times \ln\left(\frac{P_t}{P_{t-1}}\right) \quad (1.6)$$

上述计算严格来说是一种近似计算,对于简单收益率而言,上述计算公式分别如下:

$$100\% \times \left[\left(\frac{P_t}{P_{t-1}}\right)^4 - 1\right] \quad \text{和} \quad 100\% \times \left[\left(\frac{P_t}{P_{t-1}}\right)^{12} - 1\right]$$

从基本的高数知识可知,在一般情况下,这两种计算方法近似相等。

1.4 常用的统计学与概率知识

1.4.1 随机变量

1. 随机变量的含义

随机变量是一个随机结果的数值概括。例如未来股票的价格、1年GDP值、企业生产总产值等,这些经济方面的观测值都是随机变量。随机变量根据其结果取值的特性,可以分为离散型随机变量和连续型随机变量,若所有可能的结果是有限的或可数无限的,则随机变量是离散型;若所有可能的结果充满一个或若干有限或无限区间,则随机变量是连续型的。

2. 随机变量的概率分布

为了描述随机变量的不确定性,往往将随机结果与概率联系起来。随机变量 X 所有可能取值为 x_i 的概率就是 X 的概率分布。对于离散型变量,其概率分布为

$$P(X = x_i) = p_i \quad (i = 1, 2, 3, \dots, n) \quad (1.7)$$

对于连续型变量,所有可能的结果对应的概率为0,度量该随机变量在某一特定范围或区间内的概率才有实际意义。概率密度函数(PDF)定义为 $f(x)$ 且满足:

$$\begin{aligned} P(a \leq x \leq b) &= \int_a^b f(x) dx \\ f(x) &\geq 0 \\ \int_{-\infty}^{\infty} f(x) dx &= 1 \end{aligned} \quad (1.8)$$

3. 随机变量的累积分布函数

随机变量小于或者等于某个特定值的概率记为 $F(x)$,就是累积分布函数(CDF),定义如下:

$$F(x) = P(X \leq x) \quad (1.9)$$

容易得出离散随机变量 X 的累积分布函数为

$$F(x) = \sum_{x_i \leq x} p_i \quad (1.10)$$

连续型随机变量 X 的累积分布函数为

$$F(x) = \int_{-\infty}^x f(x) dx \quad (1.11)$$

累积分布函数和概率密度函数可以准确而全面地描述随机变量的分布情况。

4. 随机变量分布的矩条件

矩条件是金融计量建模中重要的内容,这里先回归一下随机变量分布矩条件的一些基本概念。假设 X 是一个密度函数为 $f(x)$ 的连续型随机变量,则矩条件分别定义如下。

p 阶原点矩定义为

$$E[X^p] = \int_{-\infty}^{\infty} x^p f(x) dx \quad (1.12)$$

p 阶中心矩定义为

$$E\{[X - E(X)]^p\} = \int_{-\infty}^{\infty} [x - E(X)]^p f(x) dx \quad (1.13)$$

1) 随机变量的期望

期望就是随机变量平均值,度量了随机变量的集中趋势。用 $E(X)$ 来表示随机变量的数学期望,离散型随机变量期望定义为

$$E(X) = \sum_{i=1}^n x_i p_i \quad (1.14)$$

式中, x_i 为随机变量 X 所有可能取值; p_i 为 x_i 发生的概率。

连续型随机变量期望定义为

$$E(X) = \int_{-\infty}^{+\infty} x f(x) dx \quad (1.15)$$

式中, $f(x)$ 为随机变量 X 的概率密度函数。

数学期望有一个重要的性质:

$$E(a + bx) = a + bE(x) \quad (1.16)$$

式中, a 和 b 都为常数。

期望也是随机变量的一阶矩,其在金融中的含义是金融资产平均收益或者期望收益,反映的是金融资产收益的一阶矩风险。

2) 随机变量的方差和标准差

方差是刻画随机变量偏离期望的程度,偏离的量 $X - E(X)$ 有正、有负,为了不使正负偏离彼此抵消,一般考虑 $[X - E(X)]^2$,这也是一个随机变量,取其均值就可以刻画随机变量 X 的波动程度。所以方差 $\sigma^2(X)$ 定义如下:

$$\text{Var}(X) = \sigma^2(X) = \begin{cases} \sum_i [x_i - E(X)]^2 p_i & (X \text{ 是离散型随机变量}) \\ \int_{-\infty}^{+\infty} [x - E(X)]^2 f(x) dx & (X \text{ 是连续型随机变量}) \end{cases} \quad (1.17)$$

由于期末的收益是不确定的,所以期末的资产回报率是随机变量,收益的方差度量收益率的分散趋势程度,经常用方差来衡量经济变量的波动性。标准差是方差的开方,通常用 σ 来表示:

$$\sigma = \sqrt{\sigma^2(X)} \quad (1.18)$$

方差的一个重要性质是

$$\sigma^2(a + bX) = b^2 \sigma^2(X) \quad (1.19)$$

式中, a 和 b 都为常数。

方差或者标准差也是随机变量的二阶矩,其在金融中的含义是金融资产的一般波动(volatility),反映的是金融资产收益的二阶矩风险。

3) 随机变量的偏度和峰度

偏度是用于衡量分布的不对称程度或偏斜程度的指标,用 $E\{[X - E(X)]^3\}$ 来表示,也

称三阶矩。实际中通常用偏度系数 S 表示对称性程度,当 $S=0$ 时,分布对称;当 $S>0$ 时,为正偏斜,有个较长的右尾部,均值大于中位数;当 $S<0$ 时,为负偏斜,有个较长的左尾部,均值小于中位数,其定义如下:

$$S = \frac{E\{[X - E(X)]^3\}}{\sigma^3} \quad (1.20)$$

对于离散型随机变量 X 的一系列观测值 $(x_1, x_2, \dots, x_n)$ 组成的样本,其偏度系数 S 的计算方法是

$$S = \frac{1}{n} \sum_{i=1}^n \frac{(x_i - \bar{x})^3}{[\sigma^2(x_i)]^{3/2}} \quad (1.21)$$

式中, n 为样本的个数; x_i 为样本的观测值; $\bar{x}$ 为观测值的平均值; $\sigma^2(x_i)$ 为观测值的方差。

峰度是随机变量的四阶矩,用 $E\{[X - E(X)]^4\}$ 来表示。峰度是用于衡量分布的集中程度或分布曲线的尖峭程度的指标,也就是反映分布函数尾巴厚度的指标。实际中通常用峰度系数 K 来表示,正态分布的 $K=3$,当 $K>3$ 时,分布呈尖峰状态。如金融时间序列多是这种分布,则具有“尖峰厚尾”(leptokurtosis and fat-tail)的特征; $K<3$ 为扁峰(矮峰)分布,其定义如下:

$$K = \frac{E\{[X - E(X)]^4\}}{\sigma^4} \quad (1.22)$$

同样地,对于离散型随机变量 X 的一系列观测值 $(x_1, x_2, \dots, x_n)$ 组成的样本,其偏度系数 K 的计算方法如下:

$$K = \frac{1}{n} \sum_{i=1}^n \frac{(x_i - \bar{x})^4}{[\sigma^2(x_i)]^2} \quad (1.23)$$

偏度在金融中的含义是衡量金融资产受信息冲击的非对称程度,反映的是金融资产收益的三阶矩风险(偏度风险)。峰度在金融中的含义是测度极端风险发生的可能性,反映的是金融资产收益的四阶矩风险(峰度风险)。三阶矩风险和四阶矩风险统称高阶矩风险。

以上阶矩风险统称矩风险,它们在金融资产定价、投资组合、风险建模与管理等方面有着广泛的应用。

1.4.2 常用概率分布

每个随机变量都有其特定的概率分布特征,下面介绍金融计量学中常用的几种概率分布,即正态分布(normal distribution)、 χ^2 分布(chisquare distribution)、 t 分布(t distribution)、 F 分布(F distribution)以及二项分布(binomial distribution)。

1. 正态分布

服从正态分布的连续型随机变量的概率密度曲线如图 1-2 所示,其概率密度为

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-(x-\mu)^2/(2\sigma^2)} \quad (-\infty < x < +\infty) \quad (1.24)$$

此时称随机变量 X 服从正态分布,记作 $X \sim N(\mu, \sigma^2)$,其中, $\mu, \sigma (\sigma > 0)$ 分别是正态分布的期望和方差。正态分布也称为高斯(Gauss)分布。

对于 $\mu=0, \sigma=1$ 的特殊情况,即如果 $X \sim N(0, 1)$,则称 X 服从标准正态分布,它的概

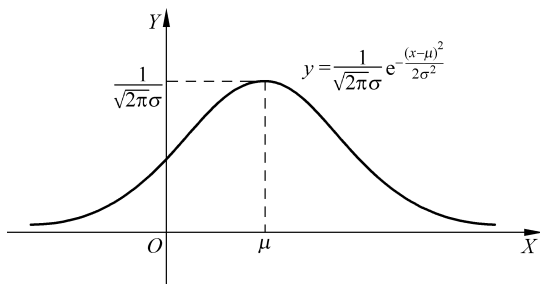


图 1-2 正态分布概率密度曲线

率密度有

$$\phi(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} \quad (1.25)$$

从理论上讲,正态分布具有很多良好的性质,许多概率分布可以用它来近似;还有一些常用的概率分布是由它直接导出的,下面介绍的三种分布都是由正态分布推导得到的。

2. χ^2 分布

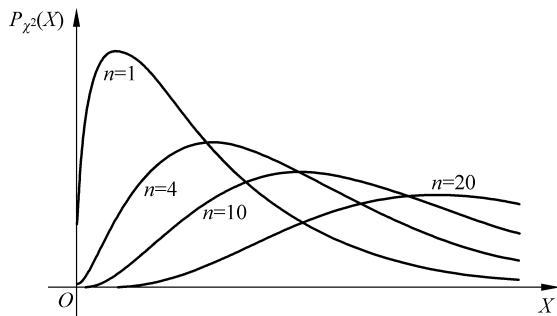
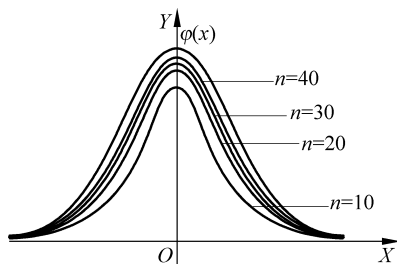
随机变量 $X_1, X_2, \dots, X_n$ 独立同分布于标准正态分布 $N(0, 1)$, 则 $\chi^2 = X_1^2 + X_2^2 + \dots + X_n^2$ 的分布称为自由度为 n 的 χ^2 分布, 记为 $\chi^2 \sim \chi^2(n)$ 。图 1-3 所示为不同的 n 值 χ^2 分布的密度函数。

χ^2 分布在第一象限内, 呈正偏态, 随着参数 n 的增大, χ^2 分布趋近于正态分布。自由度为 n 的 χ^2 分布, 其期望值为 $E(\chi^2) = n$, 方差 $\sigma(\chi^2) = 2n$ 。如果来自方差为 σ^2 的一个正态分布的 N 个观测值的样本方差为 s^2 , 则可以得到

$$(N-1)s^2/\sigma^2 \sim \chi^2(N-1)$$

3. t 分布

t 分布又称学生的 t 分布, 它与正态分布和 χ^2 分布密切相关, 其定义如下: 随机变量 X 和 Y 独立, 且 $X \sim N(0, 1)$, $Y \sim \chi^2(n)$, 则称 $X/\sqrt{Y/n}$ 的分布为自由度为 n 的 t 分布, 记为 $t = X/\sqrt{Y/n} \sim t(n)$ 。图 1-4 所示为不同自由度下 t 分布的密度函数曲线。

图 1-3 不同的 n 值 χ^2 分布的密度函数图 1-4 不同自由度下 t 分布的密度函数曲线

t 分布是一簇曲线,以 O 为中心、左右对称的单峰分布;其形态变化与自由度 n 大小有关。自由度 n 越小, t 分布曲线越低平;自由度 n 越大, t 分布曲线越接近标准正态分布曲线。对应于每一个自由度 n ,都有一条 t 分布曲线,每条曲线都有其曲线下统计量 t 的分布规律,计算较复杂。

t 分布在概率统计中计算置信区间估计、显著性检验等问题时发挥重要作用。还有一个很有用的结论,设 $x_1, x_2, \dots, x_N$ 是来自正态分布 $N(\mu, \sigma^2)$ 的一个样本, N 个观测值的样本方差为 s^2 , 样本均值为 $\bar{x}$, 则有

$$\frac{(\bar{x} - \mu) / (\sigma / \sqrt{N})}{\sqrt{(N-1)s^2 / \sigma^2 (N-1)}} = \frac{\sqrt{N}(\bar{x} - \mu)}{s} \sim t(N-1) \quad (1.26)$$

当然,这里 t 分布是对称的,其实也存在着偏 t 分布,其分布图是非对称的。

4. F 分布

设 X 与 Y 是相互独立的随机变量,且 $X \sim \chi^2(m)$, $Y \sim \chi^2(n)$, 则称 $F = \frac{X/m}{Y/n}$ 的分布是自由度为 m 和 n 的 F 分布,记为 $F \sim F(m, n)$, 其中, m 称为分子自由度或第一自由度, n 称为分母自由度或第二自由度。图 1-5 所示为不同自由度的 F 分布密度函数曲线。

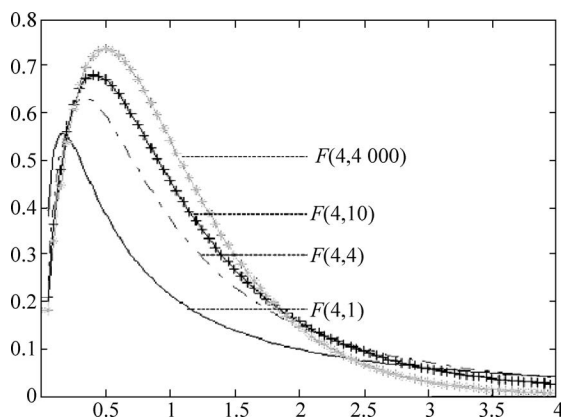


图 1-5 不同自由度的 F 分布密度函数曲线

F 分布是一种非对称分布,另外有两个结论: ①若随机变量 $F \sim F(m, n)$, 则 $\frac{1}{F} \sim F(n, m)$; ②若 $t \sim t(n)$, 则 $t^2 \sim F(1, n)$ 。

5. 二项分布

在金融计量领域中,有一些随机事件是只具有两种互斥结果的离散型随机事件,称为二分变量(dichotomous variable)。二项分布就是对这类只具有两种互斥结果的离散型随机事件的规律性进行描述的一种概率分布。

X 是一个离散型随机变量,其取值的概率分布为二项分布,记为

$$X \sim B(n, p) \quad (1.27)$$

1.4.3 假设检验

1. 假设检验的概念

假设检验是推论统计的重要内容,是先对总体的未知数量特征作出某种假设,然后抽取样本,利用样本信息对假设的正确性进行判断的过程。统计假设分为参数假设、总体分布假设、相互关系假设,如检验两个变量是否独立、两个分布是否相同等。参数假设是对总体参数的一种看法。

假设检验的实质是对可信性的评价,是对一个不确定问题的决策过程,其结果在一定概率上正确,而不是全部。

2. 假设检验的基本原理

假设检验所依据的基本原理是小概率原理,小概率原理就是发生概率很小的随机事件,在一次实验中几乎是不可能发生的。

根据这一原理,可以先假设总体参数的某项取值为真,也就是假设其发生的可能性很大,然后抽取一个样本进行观察,如果样本信息显示出现了与事先假设相反的结果且与原假设差别很大,则说明原来假定的小概率事件在一次实验中发生了,这是一个违背小概率原理的不合理现象,因此有理由怀疑和拒绝原假设;否则不能拒绝原假设。

在假设检验中,我们称检验对象的待检验假设为原假设或零假设,用 H_0 表示。原假设的对立假设称为备择假设或备选假设,用 H_1 表示。在规定了检验的显著性水平 α 后,根据容量为 n 的样本,按照统计量的理论概率分布规律,可以确定据以判断拒绝和接受原假设的检验统计量的临界值。临界值将统计量的所有可能取值区间分为两个互不相交的部分,即原假设的拒绝域和接受域。统计量估计值落在拒绝域就拒绝原假设,落在接受域就接受原假设。

3. 假设检验的两类错误

假设检验是依据样本提供的信息进行判断,有犯错误的可能。所犯错误有以下两种类型(表 1-1)。

表 1-1 假设检验中各种可能结果的概率

H_0 真伪	接受 H_0	接受 H_1
H_0 为真	$1-\alpha$ (正确决策)	α (弃真错误)
H_0 为伪	β (取伪错误)	$1-\beta$ (正确决策)

第一类错误是原假设为真时,检验结果把它当成不真而拒绝了。犯这种错误的概率用 α 表示,也称作 α 错误或弃真错误。

第二类错误是原假设不为真时,检验结果把它当成真而接受了。犯这种错误的概率用 β 表示,也称作 β 错误或取伪错误。

4. 假设检验的步骤

假设检验的步骤如下。

- (1) 根据问题要求,提出原假设和备择假设。
- (2) 选择适当的检验统计量。
- (3) 确定检验的方向性并规定显著性水平。
- (4) 计算检验统计量的值。
- (5) 将统计量的值与临界值对比作出决策。

5. 置信区间法与显著性检验法

置信区间法与显著性检验法是判定原假设的两种互补方法。置信区间法是在给定置信水平下求出置信区间,该置信区间就是原假设的接受域,如果原假设落在置信区间内,则接受原假设;如果原假设没有落在置信区间内,则拒绝原假设。显著性检验法则是把原假设代入统计量中,得到统计量对应的值,将该值与统计量分布在显著性水平下查表对应的临界值相比较,判断是否拒绝原假设。

6. P 值检验法

假设检验的结论通常是很简单的,在给定的显著性水平下,不是拒绝原假设就是接受原假设,然而,也可能出现,在较大的显著性水平下(如 $\alpha=5\%$)得到拒绝原假设,在较小的显著性水平下(如 $\alpha=1\%$)却是接受原假设。显著性水平变小时,拒绝域会变小,原来落在拒绝域中的观测值就落在接受域了。为了更好地检验,引入检验的 P 值,它表示的是统计量拒绝原假设的最小显著性水平。所以只要应用 P 值与给定的显著性水平相比较,若 $\alpha \geq P$,则在显著性水平 α 下拒绝原假设;若 $\alpha < P$,则在显著性水平 α 下接受原假设。

1.4.4 三个常用的检验方法

似然比检验、沃尔德(Wald)检验、拉格朗日乘数(Lagrange Multiplier, LM)检验是三个常用的检验方法,它们都基于极大似然估计。另外,只有当估计量满足一致性和渐进正态性的条件,这三种检验给出的分布结果才是有效的。

在介绍这三种检验方法之前,做如下假设,假设由极大似然估计方法已经估计得到参数 θ ,想要检验线性参数约束条件为

$$H_0: \mathbf{h}(\theta) = \mathbf{0} \quad (1.28)$$

是否显著成立,其备择假设为

$$H_1: \mathbf{h}(\theta) \neq \mathbf{0} \quad (1.29)$$

其中, $\mathbf{h}(\theta)$ 是关于 θ 可导的向量。

1. 似然比检验

似然比检验(LR 检验)是由耶日·奈曼(Jerzy Neyman)和埃贡·皮尔逊(Egon Pearson)于1928年提出的,它的思想是:如果参数约束 $\mathbf{h}(\theta) = \mathbf{0}$ 是有效的,那么加上这样的约束不应该引起似然函数最大值的大幅度降低。也就是说,似然比检验的实质是在比较有约束条件下的似然函数最大值与无约束条件下似然函数最大值。

似然比定义为有约束条件下的似然函数最大值与无约束条件下似然函数最大值之比,即

$$\lambda = \frac{L(\mathbf{x}, \tilde{\theta})}{L(\mathbf{x}, \hat{\theta})} \quad (1.30)$$

其中, $\hat{\theta}$ 为没有约束条件[即不要求满足 $\mathbf{h}(\theta) = \mathbf{0}$ 约束条件]的似然函数取最大值; $L(\mathbf{x}, \hat{\theta})$ 为对应的极大似然估计量; $\tilde{\theta}$ 为加上约束条件 $\mathbf{h}(\theta) = \mathbf{0}$ 的似然函数最大值; $L(\mathbf{x}, \tilde{\theta})$ 为对应的极大似然估计量; λ 为似然比。

两个似然函数取最大值时都是正的, 而且 $L(\mathbf{x}, \tilde{\theta})$ 不可能大于 $L(\mathbf{x}, \hat{\theta})$, 因为有约束条件下的目标函数最大值一定不会大于没有约束条件下目标函数的最大值, 所以 λ 值一定是介于 0 和 1 之间。如果约束是有效的, 则 $L(\mathbf{x}, \tilde{\theta})$ 不会比 $L(\mathbf{x}, \hat{\theta})$ 小很多, λ 值应该接近 1; 如果约束是无效的, 则 $L(\mathbf{x}, \tilde{\theta})$ 可能会比 $L(\mathbf{x}, \hat{\theta})$ 小很多, λ 值应该接近 0。当 λ 值较大时, 我们可以接受约束条件。

似然比检验统计量为

$$\text{LR} = -2\ln \lambda = 2[\ln L(\mathbf{x}, \hat{\theta}) - \ln L(\mathbf{x}, \tilde{\theta})] \sim \chi^2(p) \quad (1.31)$$

式中, 自由度 p 是约束条件的个数。

如果 $\text{LR} > \chi^2(p)$, 则拒绝原假设, 参数约束无效; 如果 $\text{LR} \leq \chi^2(p)$, 则接受原假设, 参数约束有效。

2. 沃尔德检验

似然比检验的一个实际缺点是它要求估计约束和无约束参数向量的极大似然估计量的最大值都已知, 而在复杂的模型中, 有约束条件的极大似然估计很难实现, 可以尝试沃尔德检验。沃尔德检验的思想是: 如果约束是有效的, 那么在没有约束的情况下, 估计出来的估计量应该渐进地满足约束条件, 即若约束条件 $\mathbf{h}(\hat{\theta}) = \mathbf{0}$ 是有效的, 在没有约束条件情况下估计出的 $\hat{\theta}$, 应该渐进地满足 $\mathbf{h}(\hat{\theta}) \approx \mathbf{0}$, 因为极大似然估计是一致的。

以无约束估计量为基础可以构造一个沃尔德检验统计量, 这个统计量也服从卡方分布:

$$\mathbf{W} = \mathbf{h}'(\hat{\theta}) \{ \text{Var}[\mathbf{h}'(\hat{\theta})] \}^{-1} \mathbf{h}(\hat{\theta}) \sim \chi^2(p) \quad (1.32)$$

式中, 自由度 p 是约束条件的个数, $\text{Var}[\mathbf{h}'(\hat{\theta})]$ 为 $\mathbf{h}'(\hat{\theta})$ 的方差-协方差矩阵。如果 $\mathbf{W} > \chi^2(p)$, 则拒绝原假设, 参数约束无效; 如果 $\mathbf{W} \leq \chi^2(p)$, 则接受原假设, 参数约束有效。

3. 拉格朗日乘数检验

当没有约束条件的极大似然估计很难实现, 而有约束条件的极大似然估计比较容易得到时, 可以应用拉格朗日乘数检验。假设在约束条件 $\mathbf{h}(\theta) = \mathbf{0}$ 下, 最大化对数似然函数, 设 λ 是一个拉格朗日乘子向量, $\tilde{\theta}$ 表示加上约束条件 $\mathbf{h}(\theta) = \mathbf{0}$ 的似然函数取最大值 $L_R(\mathbf{x}, \tilde{\theta})$ 时对应的极大似然估计量, 最大化拉格朗日函数 $\ln L_R(\mathbf{x}, \theta)$ 可以得到

$$\begin{aligned} \frac{\partial \ln L_R(\mathbf{x}, \theta)}{\partial \theta} &= \frac{\partial \ln L(\mathbf{x}, \theta)}{\partial \theta} + \frac{\partial \mathbf{h}(\theta)}{\partial \theta} \lambda = 0 \\ \Rightarrow \frac{\partial \ln L(\mathbf{x}, \tilde{\theta})}{\partial \theta} &= - \frac{\partial \mathbf{h}(\theta)}{\partial \theta} \tilde{\lambda} \end{aligned} \quad (1.33)$$

若约束 $h(\theta) = \mathbf{0}$ 是有效的,约束条件不会使似然函数最大值发生太大的偏差,则导数向量中的第二项应该很小,特别的是, $\tilde{\lambda}$ 值应该趋近于 0。所以,拉格朗日乘数检验就是在有约束估计量 $\tilde{\theta}$ 处,通过检验得出向量 $\frac{\partial \ln L(\mathbf{x}, \tilde{\theta})}{\partial \theta}$ 是否趋于零来检验约束是否有效。

这里构造一个 LM 统计量:

$$LM = \left[\frac{\partial \ln L(\mathbf{x}, \tilde{\theta})}{\partial \theta} \right]' \left\{ \text{Var} \left[\frac{\partial \ln L(\mathbf{x}, \tilde{\theta})}{\partial \theta} \right] \right\}^{-1} \left[\frac{\partial \ln L(\mathbf{x}, \tilde{\theta})}{\partial \theta} \right] \sim \chi^2(p) \quad (1.34)$$

式中,自由度 p 是约束条件的个数。如果 $LM > \chi^2(p)$,则拒绝原假设,参数约束无效;如果 $LM \leq \chi^2(p)$,则接受原假设,参数约束有效。

在满足大样本的条件下,三种检验方法是渐进等价的,但是在小样本情况下,它们可能表现得相当不同。对于似然比检验,既需要估计有约束的模型,也需要估计无约束的模型;对于沃尔德检验,只需要估计无约束模型;对于拉格朗日乘数检验,只需要估计有约束的模型。一般情况下,由于估计有约束模型相对更复杂,所以沃尔德检验最为常用。对于小样本而言,似然比检验的渐进性最好,拉格朗日乘数检验也较好,沃尔德检验有时会拒绝原假设,其小样本性质不尽如人意。

1.5 常用金融计量软件介绍

1.5.1 常用金融计量软件简介

估计金融计量学模型可能用到各种各样的软件包,随着科技的发展,可用软件包的数量越来越多,所有软件包都在可用的技术所允许的范围内得到改进。找到一款合适的软件包处理金融时间序列数据相当重要,这里介绍几种主要的金融计量软件。

1. EViews 软件

EViews 是美国 QMS 公司研制的在 Windows 下专门从事数据分析、回归分析和预测的工具。使用 EViews 可以迅速地从数据中寻找出统计关系,并用得到的关系去预测数据的未来值。金融计量学上,EViews 是完成设计模型、估计模型、检验模型、应用模型的一项重要工具。EViews 是一款使用简便的交互式计算机软件,通过菜单很容易实现表格、条形图、序列图和回归结果,使用广泛。正是由于 EViews 等计量经济学软件的出现,计量经济学才取得了长足的进步,发展成为一门较为实用与严谨的经济学科。

2. SPSS 软件

SPSS(Statistical Package for the Social Sciences,社会科学统计程序包)是世界公认的最优秀的统计分析软件之一。20 世纪 60 年代末,美国斯坦福大学的 3 位研究生开发了最早的统计分析软件 SPSS,并于 1975 年在芝加哥成立了 SPSS 公司。“易学,易用,易普及”已成为 SPSS 软件最大的竞争优势之一,也是广大数据分析人员对其偏爱有加的主要原因。SPSS 操作简单,被广泛应用于统计学分析运算、数据挖掘、预测分析和决策支持任务等。其缺点是没有纳入最新统计方法;采用 VB 编制,计算速度慢;输出结果不能和常用文字处理软件直接兼容。

3. SAS 软件

SAS(Statistics Analysis System)最早由北卡罗来纳大学的两位生物统计学研究生编制,并于1976年成立了SAS软件研究所,正式推出了SAS软件。SAS软件研究所一直致力于为金融、医药研发、保险、电信、制造、政府以及科研教育等部门,在SAS应用数据仓库,统计分析、联机分析处理系统进行质量管理、财务管理、生产优化、风险管理、市场调查等业务。用户应用SAS软件还可以通过对数据集的一连串加工,实现更为复杂的统计分析。此外,SAS还提供了各类概率分析函数、分位数函数、样本统计函数和随机数生成函数,满足用户特殊统计要求。但是,这使得初学者在使用SAS时必须学习SAS语言,入门比较困难。

4. Matlab 软件

Matlab是MathWorks公司于1982年推出的一套高性能的数值计算和可视化软件。它集数值分析、矩阵运算、信号处理和图形显示于一体,构成了一个方便、界面良好的用户环境。它还包括Toolbox(工具箱)的各类问题的求解工具,可用来求解特定学科的问题。例如控制系统设计与分析、图像处理、信号处理与通信、金融建模和分析等。另外,它还有一个配套软件包Simulink,提供了一个可视化开发环境,常用于系统模拟、动态/嵌入式系统开发等方面。

5. Statistica 软件

Statistica是一个整合数据分析、图表绘制、数据库管理与自订应用发展系统环境的专业软件。它的图形库种类非常丰富,有同步报告输出功能,另外,Statistica和Excel、C++、Java等多种外部应用工具兼容,通过强大且直观的询问工具,可以连接并处理存储于任何地方的大型资料,而无须事先将资料复制到当前电脑中,就可以执行理想的大型资料提取和资料的分析任务。Statistica软件使用简便,操作界面、操作方式和Office软件类似,友好的操作界面方便客户操作并满足客户定制化的需求。

6. Stata 软件

Stata软件功能强大,是一套为使用者提供数据分析、数据管理以及绘制完整的专业图表及整合性统计软件。它除了传统的统计分析方法外,还收集了近20年发展起来的新方法,如Cox比例风险回归,指数与Weibull回归,多类结果与有序结果的logistic回归,Poisson回归,负二项回归及广义负二项回归,随机效应模型等。Stata的作图模块丰富,有八种基本图形的制作:直方图(histogram)、条形图(bar)、百分条图(oneway)、百分饼图(pie)、散点图(twoway)、散点图矩阵(matrix)、星形图(star)、分位数图。这些图形的巧妙应用,可以满足绝大多数用户的统计作图要求。

7. R 软件

R软件是一个有着统计分析功能及强大作图功能的软件系统,由奥克兰大学统计学系的Ross Ihaka和Robert Gentleman共同创立。R是一个庞大的体系,可以通过编写R语言实现,语言简单,容易实现,主要的统计方法有贝叶斯推断、聚类分析、机器学习、空间统计、稳健统计等。而这些方法又通过相应的R Packages扩展,可以说学习R是一件没有尽头的事情。

1.5.2 EViews 软件使用介绍及操作步骤简介

1. EViews 菜单栏功能介绍

EViews 菜单栏提供了 10 个选项:“File”“Edit”“Object”“View”“Proc”“Quick”“Options”“Add-ins”“Windows”与“Help”。单击这些选项,下方就会出现不同的下拉菜单,这里介绍它们的主要功能。

File 选项主要为用户提供有关文件(工作文件、数据库文件、EViews 程序)的常规操作选项,如文件的建立、打开、保存、关闭、导入数据、读出数据、打印、运行程序等。Edit 选项提供的是编辑功能,如剪切、复制、粘贴、删除、查找、替换、合并等。Object 提供了有关 EViews 对象的各种操作,如建立新的对象、复制、删除、给对象命名、删除、视图选择等。View 和 Proc 二者会根据不同的窗口发生改变,主要涉及用户对对象实行的运算过程和对对象的多种显示方式。Quick 选项包含一些数据的分析命令,如抽取一定范围的样本、生成新的序列、显示某一观测值、创建图形、生成一个新的序列组、做序列描述性统计、估计方程、估计 VAR(向量自回归)模型等。Options 选项是参数设定选项,可以对 EViews 中图像、字体、方程估计等默认设置进行修改。Add-ins 选项类似软件的加载工具包,扩展 EViews 计量模型的功能。Windows 选项主要是对打开的窗口进行切换以及关闭对象等功能。Help 提供 EViews 软件的帮助服务,其中详细地列出 EViews 8.0 的基本操作指南,还包括一些对象的参考、基本命令参考、函数参考等。

2. EViews 工作文件基础

EViews 软件的具体操作在 Workfile 中进行。如果想用 EViews 进行数据分析,必须先新建一个 Workfile 工作文件或打开一个已经储存在硬盘(或软盘)上的 Workfile,然后才能够进行定义变量、输入数据、建造模型等操作。这里用的是 EViews 8.0 软件进行演示,新建工作文件步骤如下。

打开 EViews 软件,在 EViews Workfile 窗口(图 1-6)单击 Create a new EViews workfile;或者在主菜单上依次单击“File”→“New Workfile”,即选择新建对象的类型为工作文件,将弹出一个对话框(图 1-7)。

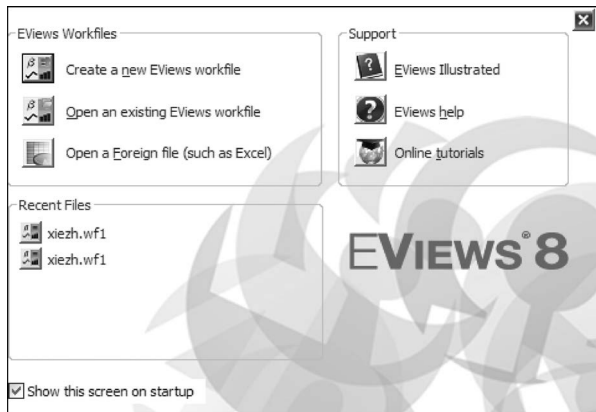


图 1-6 EViews Workfile 窗口

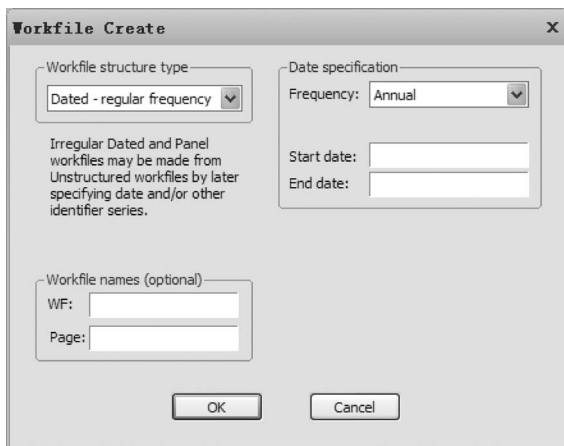


图 1-7 Workfile Create 对话框

在图 1-7 中,对话框中的“Workfile structure type”项用于设置工作文件的数据结构类型,依顺序分别是非结构/非时间数据、时间频率数据、平衡面板数据。处理时间序列数据,一般选择“Dated-regular frequency”,表示创建规则的时间序列结构类型的文件,根据实际需要选择数据的时间频率(frequency)、起始期和终止期。选定时间频率之后,输入相应的日期(如 1978 和 2013,表示从 1978 年到 2013 年每年的数据)。输入季度、月份日期时用符号“:”或“.”分隔年、季或月。另外两个数据类型,只要输入观测值的个数即可。输入后单击 OK 按钮,就创建了一个 Workfile 工作文件。

3. 创建工作文档中的对象

EViews 的工作文件最主要的是对象,对象的类型有很多种,如序列(series)对象,指的是一序列特定变量的观测值,包括时间序列和非时间序列;数组(group)由若干序列组合而成,便于对多个变量执行操作;图形(graph)主要类型有折线图(趋势图)、相关图、柱状图、饼状图、分布图、QQ 图、箱线图,可以在图像窗口定义图形参数;方程(equation)指含有变量之间相互关系的信息集合。

打开建立的 Workfile 工作文件,选择 Object,选择 New Object,出现图 1-8 所示的对象定义设置对话框。从上到下对象的类型依次为: Equation(方程)、Factor(因子)、Graph(图形)、Group(数组)等 22 种对象。右边 Name for Object 框是对对象的命名窗口,EViews 对象的命名不区分大小写。

4. 数据的导入

建立工作文件和创建对象之后,接下来的工作就是导入数据。数据导入的方法有很多,如果数据不多,可以选择直接从键盘输入数据,打开建立的对象,

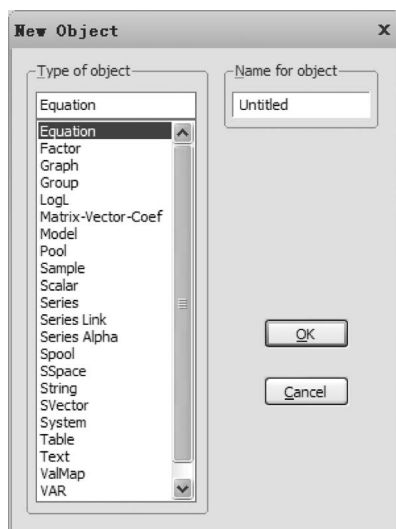


图 1-8 New Object 窗口

出现图 1-9 所示的窗口,在工具栏中单击“Edit+/-”按钮,就可以进入数据编辑状态,再输入数值之后,按 Enter 键,就完成数据的输入操作。

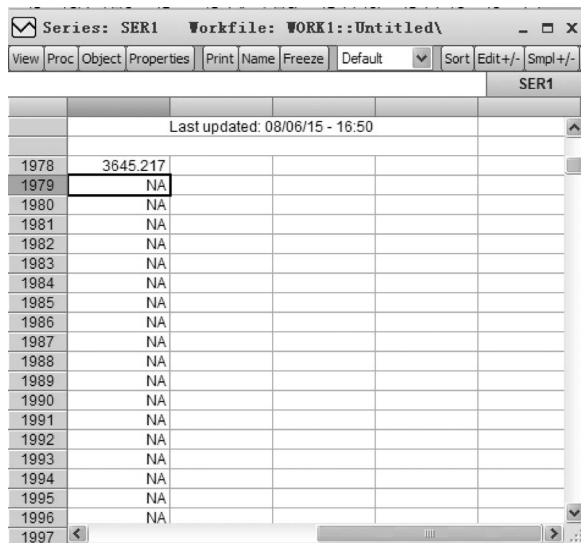


图 1-9 数据输入示意图

导入数据的方法还有直接复制/粘贴输入、直接复制数据、在工具栏中单击“Edit+/-”按钮后直接粘贴数据。也可以从外部直接导入数据,接受调用数据的格式有 ASCII、Lotus、Excel 工作表。选择“Procs”→“Import”→“Read Text-Lotus-Excel”命令,打开目标文件,注意数据频率的一致性。

5. 绘制图示

EViews 中提供了折线图(趋势图)、相关图、柱状图、饼状图、分布图、QQ 图、箱线图等图形,在输入数据之后,可以绘制各种图形,下面演示绘制 1978—2013 年的中国 GDP 数据的折线图和柱状图。数据已经输入保存在 ser1 对象中,打开数据,选择“View”→“Graph”,出现图 1-10 所示的窗口,在“Specific”窗口中选择“Line”(折线图),该窗口包含其他多种图

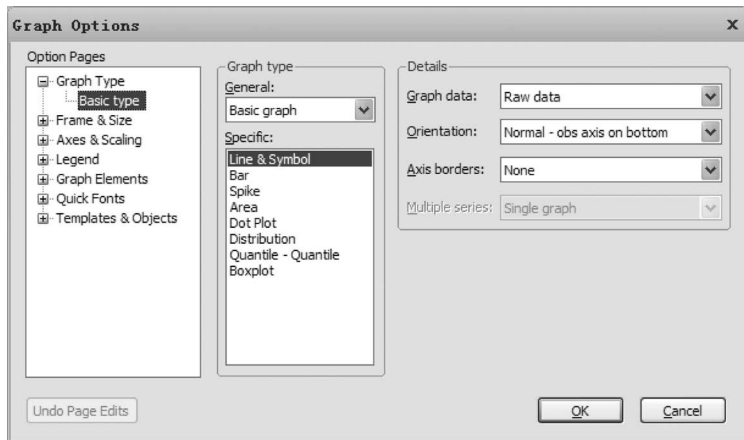


图 1-10 图形选择窗口

形,选择 Bar 就会出现柱状图,单击 OK 按钮,就绘制出 GDP 数据的折线图 and 柱状图,如图 1-11 和图 1-12 所示。

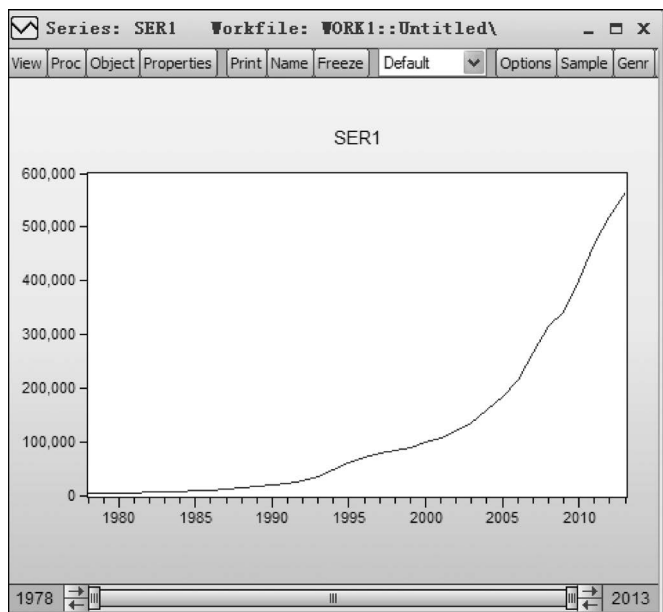


图 1-11 1978—2013 年 GDP 折线图

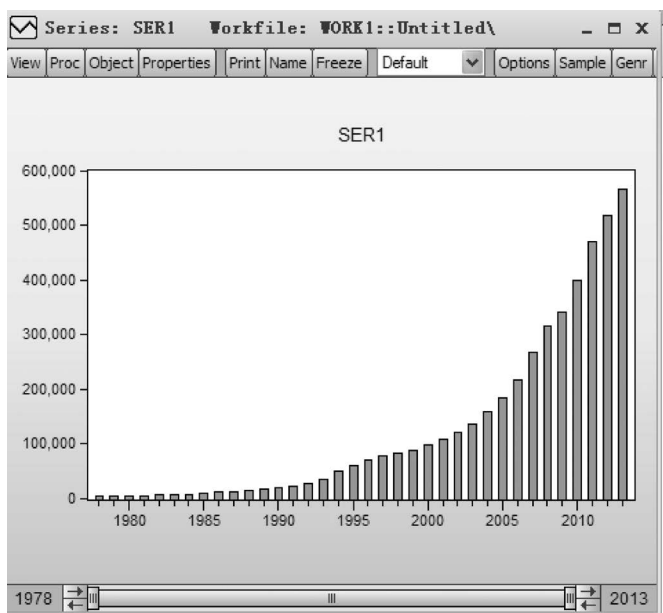


图 1-12 1978—2013 年 GDP 柱状图

图 1-13 和图 1-14 分别作出的是沪深 300 指数从 2012 年 1 月 4 日到 2015 年 7 月 31 日对数收益率数据(868 个数据)的时间序列图和分布直方图,具体操作后面章节会具体介绍。对对数收益率的描述性统计在金融计量学中占重要地位。

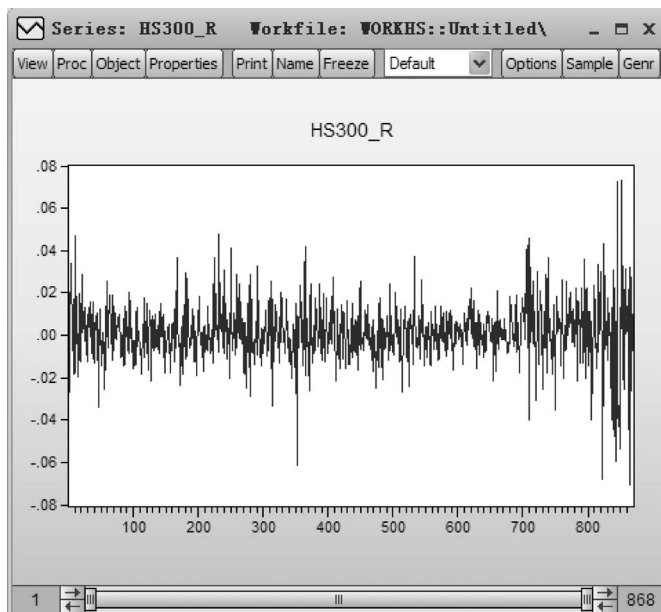


图 1-13 沪深 300 指数收益率时间序列图

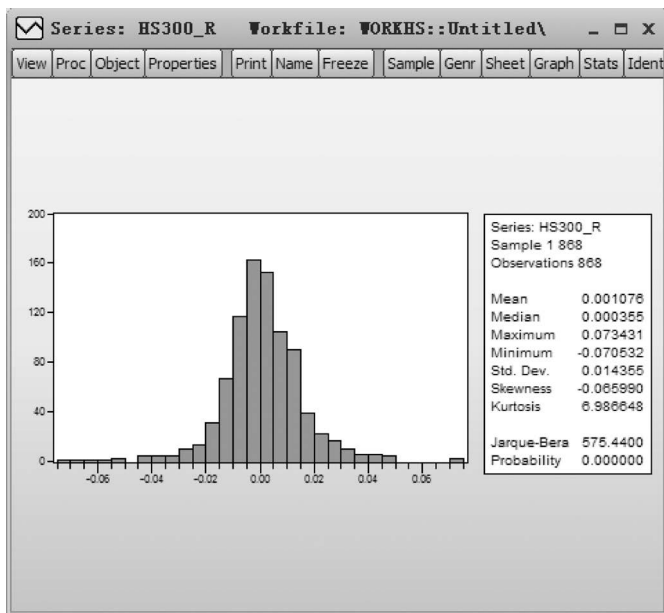


图 1-14 沪深 300 指数收益率的分布直方图

复习思考题

1. 什么是金融计量学?
2. 单期收益率与多期收益率的区别是什么?

3. 金融数据的主要来源是什么?
4. 简述假设检验的步骤。
5. 简要概括三个常用的检验方法。

即 测 即 练

