

1

Introduction

Learning econometrics is in many ways like learning a new language. To begin with, nothing makes sense and it is as if it is impossible to see through the fog created by all the unfamiliar terminology. While the way of writing the models – the *notation* – may make the situation appear more complex, in fact it is supposed to achieve the exact opposite. The ideas themselves are mostly not so complicated, it is just a matter of learning enough of the language that everything fits into place. So if you have never studied the subject before, then persevere through this preliminary chapter and you will hopefully be on your way to being fully fluent in econometrics!

Learning outcomes

In this chapter, you will learn how to

- Compare nominal and real series and convert one to the other
- Distinguish between different types of data
- Describe the key steps involved in building an econometric model
- Calculate asset price returns
- Deflate series to allow for inflation
- Construct a workfile, import data and accomplish simple tasks in EViews

The chapter sets the scene for the book by discussing in broad terms the questions of what econometrics is, and what the ‘stylised facts’ are describing financial data that researchers in this area typically try to capture in their models. Some discussion is presented on the kinds of data we encounter in finance and how to work with them. Finally, the chapter collects together a number of preliminary issues relating to the construction of econometric models in finance and introduces the software that will be used in the remainder of the book for estimating the models.

Box 1.1 Examples of the uses of econometrics

- (1) Testing whether financial markets are weak-form informationally efficient
- (2) Testing whether the capital asset pricing model (CAPM) or arbitrage pricing theory (APT) represent superior models for the determination of returns on risky assets
- (3) Measuring and forecasting the volatility of bond returns
- (4) Explaining the determinants of bond credit ratings used by the ratings agencies
- (5) Modelling long-term relationships between prices and exchange rates
- (6) Determining the optimal hedge ratio for a spot position in oil
- (7) Testing technical trading rules to determine which makes the most money
- (8) Testing the hypothesis that earnings or dividend announcements have no effect on stock prices
- (9) Testing whether spot or futures markets react more rapidly to news
- (10) Forecasting the correlation between the stock indices of two countries.

1.1 What is econometrics?

The literal meaning of the word econometrics is ‘measurement in economics’. The first four letters of the word suggest correctly that the origins of econometrics are rooted in economics. However, the main techniques employed for studying economic problems are of equal importance in financial applications. As the term is used in this book, financial econometrics will be defined as the *application of statistical techniques to problems in finance*. Financial econometrics can be useful for testing theories in finance, determining asset prices or returns, testing hypotheses concerning the relationships between variables, examining the effect on financial markets of changes in economic conditions, forecasting future values of financial variables and for financial decision-making. A list of possible examples of where econometrics may be useful is given in box 1.1.

The list in box 1.1 is of course by no means exhaustive, but it hopefully gives some flavour of the usefulness of econometric tools in terms of their financial applicability.

1.2 Is financial econometrics different from ‘economic econometrics’?

As previously stated, the tools commonly used in financial applications are fundamentally the same as those used in economic applications, although the emphasis and the sets of problems that are likely to be encountered when analysing the two

sets of data are somewhat different. Financial data often differ from macroeconomic data in terms of their frequency, accuracy, seasonality and other properties.

In economics, a serious problem is often a *lack of data at hand* for testing the theory or hypothesis of interest – this is sometimes called a ‘small samples problem’. It might be, for example, that data are required on government budget deficits, or population figures, which are measured only on an annual basis. If the methods used to measure these quantities changed a quarter of a century ago, then only at most twenty-five of these annual observations are usefully available.

Two other problems that are often encountered in conducting applied econometric work in the arena of economics are those of *measurement error* and *data revisions*. These difficulties are simply that the data may be estimated, or measured with error, and will often be subject to several vintages of subsequent revisions. For example, a researcher may estimate an economic model of the effect on national output of investment in computer technology using a set of published data, only to find that the data for the last two years have been revised substantially in the next, updated publication.

These issues are usually of less concern in finance. Financial data come in many shapes and forms, but in general the prices and other entities that are recorded are those at which trades *actually took place*, or which were *quoted* on the screens of information providers. There exists, of course, the possibility for typos or for the data measurement method to change (for example, owing to stock index re-balancing or re-basing). But in general the measurement error and revisions problems are far less serious in the financial context.

Similarly, some sets of financial data are observed at much *higher frequencies* than macroeconomic data. Asset prices or yields are often available at daily, hourly or minute-by-minute frequencies. Thus the number of observations available for analysis can potentially be very large – perhaps thousands or even millions, making financial data the envy of macro-econometricians! The implication is that more powerful techniques can often be applied to financial than economic data, and that researchers may also have more confidence in the results.

Furthermore, the analysis of financial data also brings with it a number of new problems. While the difficulties associated with handling and processing such a large amount of data are not usually an issue given recent and continuing advances in computer power, financial data often have a number of additional characteristics. For example, financial data are often considered very ‘noisy’, which means that it is more difficult to separate *underlying trends or patterns* from random and uninteresting features. Financial data are also almost always not normally distributed in spite of the fact that most techniques in econometrics assume that they are. High frequency data often contain additional ‘patterns’ which are the result of the way that the market works, or the way that prices are recorded. These features need to be considered in the model-building process, even if they are not directly of interest to the researcher.

One of the most rapidly evolving areas of financial application of statistical tools is in the modelling of market microstructure problems. ‘Market microstructure’ may broadly be defined as the process whereby *investors’ preferences and desires are translated into financial market transactions*. It is evident that microstructure effects

Box 1.2 Time series data

<i>Series</i>	<i>Frequency</i>
Industrial production	Monthly or quarterly
Government budget deficit	Annually
Money supply	Weekly
The value of a stock	As transactions occur

are important and represent a key difference between financial and other types of data. These effects can potentially impact on many other areas of finance. For example, market rigidities or frictions can imply that current asset prices do not fully reflect future expected cashflows (see the discussion in chapter 10 of this book). Also, investors are likely to require compensation for holding securities that are illiquid, and therefore embody a risk that they will be difficult to sell owing to the relatively high probability of a lack of willing purchasers at the time of desired sale. Measures such as volume or the time between trades are sometimes used as proxies for market liquidity.

A comprehensive survey of the literature on market microstructure is given by Madhavan (2000). He identifies several aspects of the market microstructure literature, including price formation and price discovery, issues relating to market structure and design, information and disclosure. There are also relevant books by O'Hara (1995), Harris (2002) and Hasbrouck (2007). At the same time, there has been considerable advancement in the sophistication of econometric models applied to microstructure problems. For example, an important innovation was the autoregressive conditional duration (ACD) model attributed to Engle and Russell (1998). An interesting application can be found in Dufour and Engle (2000), who examine the effect of the time between trades on the price-impact of the trade and the speed of price adjustment.

1.3 Types of data

There are broadly three types of data that can be employed in quantitative analysis of financial problems: time series data, cross-sectional data and panel data.

1.3.1 Time series data

Time series data, as the name suggests, are data that have been collected over a period of time on one or more variables. Time series data have associated with them a particular frequency of observation or frequency of collection of data points. The frequency is simply a measure of the *interval over*, or the *regularity with which*, the data are collected or recorded. Box 1.2 shows some examples of time series data.

A word on ‘As transactions occur’ is necessary. Much financial data does not start its life as being *regularly spaced*. For example, the price of common stock for a given company might be recorded to have changed whenever there is a new trade or quotation placed by the financial information recorder. Such recordings are very unlikely to be evenly distributed over time – for example, there may be no activity between, say, 5 p.m. when the market closes and 8.30 a.m. the next day when it reopens; there is also typically less activity around the opening and closing of the market, and around lunch time. Although there are a number of ways to deal with this issue, a common and simple approach is to select an appropriate frequency, and use as the observation for that time period the last prevailing price during the interval.

It is also generally a requirement that all data used in a model be of the *same frequency of observation*. So, for example, regressions that seek to estimate an arbitrage pricing model using monthly observations on macroeconomic factors must also use monthly observations on stock returns, even if daily or weekly observations on the latter are available.

The data may be *quantitative* (e.g. exchange rates, prices, number of shares outstanding), or *qualitative* (e.g. the day of the week, a survey of the financial products purchased by private individuals over a period of time, a credit rating, etc.).

Problems that could be tackled using time series data:

- How the value of a country’s stock index has varied with that country’s macroeconomic fundamentals
- How the value of a company’s stock price has varied when it announced the value of its dividend payment
- The effect on a country’s exchange rate of an increase in its trade deficit.

In all of the above cases, it is clearly the time dimension which is the most important, and the analysis will be conducted using the values of the variables over time.

1.3.2 Cross-sectional data

Cross-sectional data are data on one or more variables collected at a single point in time. For example, the data might be on:

- A poll of usage of internet stockbroking services
- A cross-section of stock returns on the New York Stock Exchange (NYSE)
- A sample of bond credit ratings for UK banks.

Problems that could be tackled using cross-sectional data:

- The relationship between company size and the return to investing in its shares
- The relationship between a country’s GDP level and the probability that the government will default on its sovereign debt.

1.3.3 Panel data

Panel data have the dimensions of both time series and cross-sections, e.g. the daily prices of a number of blue chip stocks over two years. The estimation of panel regressions is an interesting and developing area, and will be examined in detail in chapter 11.

Fortunately, virtually all of the standard techniques and analysis in econometrics are equally valid for time series and cross-sectional data. For time series data, it is usual to denote the individual observation numbers using the index t , and the total number of observations available for analysis by T . For cross-sectional data, the individual observation numbers are indicated using the index i , and the total number of observations available for analysis by N . Note that there is, in contrast to the time series case, no natural ordering of the observations in a cross-sectional sample. For example, the observations i might be on the price of bonds of different firms at a particular point in time, ordered alphabetically by company name. So, in the case of cross-sectional data, there is unlikely to be any useful information contained in the fact that Barclays follows Banco Santander in a sample of bank credit ratings, since it is purely by chance that their names both begin with the letter 'B'. On the other hand, in a time series context, the ordering of the data is relevant since the data are usually ordered chronologically.

In this book, the total number of observations in the sample will be given by T even in the context of regression equations that could apply either to cross-sectional or to time series data.

1.3.4 Continuous and discrete data

As well as classifying data as being of the time series or cross-sectional type, we could also distinguish them as being either continuous or discrete, exactly as their labels would suggest. *Continuous* data can take on any value and are not confined to take specific numbers; their values are limited only by precision. For example, the rental yield on a property could be 6.2%, 6.24% or 6.238%, and so on. On the other hand, *discrete* data can only take on certain values, which are usually integers (whole numbers), and are often defined to be count numbers.¹ For instance, the number of people in a particular underground carriage or the number of shares traded during a day. In these cases, having 86.3 passengers in the carriage or 5857¹/₂ shares traded would not make sense. The simplest example of a discrete variable is a *Bernoulli* or binary random variable, which can only take the values 0 or 1 – for example, if we repeatedly tossed a coin, we could denote a head by 0 and a tail by 1.

¹ Discretely measured data do not necessarily have to be integers. For example, until they became 'decimalised', many financial asset prices were quoted to the nearest 1/16 or 1/32 of a dollar.

1.3.5 Cardinal, ordinal and nominal numbers

Another way in which we could classify numbers is according to whether they are cardinal, ordinal or nominal. *Cardinal* numbers are those where the actual numerical values that a particular variable takes have meaning, and where there is an equal distance between the numerical values. On the other hand, *ordinal* numbers can only be interpreted as providing a position or an ordering. Thus, for cardinal numbers, a figure of 12 implies a measure that is ‘twice as good’ as a figure of 6. Examples of cardinal numbers would be the price of a share or of a building, and the number of houses in a street. On the other hand, for an ordinal scale, a figure of 12 may be viewed as ‘better’ than a figure of 6, but could not be considered twice as good. Examples of ordinal numbers would be the position of a runner in a race (e.g. second place is better than fourth place, but it would make little sense to say it is ‘twice as good’) or the level reached in a computer game.

The final type of data that could be encountered would be where there is no natural ordering of the values at all, so a figure of 12 is simply different to that of a figure of 6, but could not be considered to be better or worse in any sense. Such data often arise when numerical values are arbitrarily assigned, such as telephone numbers or when codings are assigned to qualitative data (e.g. when describing the exchange that a US stock is traded on, ‘1’ might be used to denote the NYSE, ‘2’ to denote the NASDAQ and ‘3’ to denote the AMEX). Sometimes, such variables are called *nominal* variables. Cardinal, ordinal and nominal variables may require different modelling approaches or at least different treatments, as should become evident in the subsequent chapters.

1.4 Returns in financial modelling

In many of the problems of interest in finance, the starting point is a time series of prices – for example, the prices of shares in Ford, taken at 4 p.m. each day for 200 days. For a number of statistical reasons, it is preferable not to work directly with the price series, so that raw price series are usually converted into series of returns. Additionally, returns have the added benefit that they are unit-free. So, for example, if an annualised return were 10%, then investors know that they would have got back £110 for a £100 investment, or £1,100 for a £1,000 investment, and so on.

There are two methods used to calculate returns from a series of prices, and these involve the formation of simple returns, and continuously compounded returns, which are respectively

Simple returns

$$R_t = \frac{p_t - p_{t-1}}{p_{t-1}} \times 100\% \tag{1.1}$$

Continuously compounded returns

$$r_t = 100\% \times \ln \left(\frac{p_t}{p_{t-1}} \right) \tag{1.2}$$

Box 1.3 Log returns

- (1) Log-returns have the nice property that they can be interpreted as *continuously compounded returns* – so that the frequency of compounding of the return does not matter and thus returns across assets can more easily be compared.
- (2) Continuously compounded returns are *time-additive*. For example, suppose that a weekly returns series is required and daily log returns have been calculated for five days, numbered 1 to 5, representing the returns on Monday through Friday. It is valid to simply add up the five daily returns to obtain the return for the whole week:

Monday return	$r_1 = \ln (p_1 / p_0) = \ln p_1 - \ln p_0$
Tuesday return	$r_2 = \ln (p_2 / p_1) = \ln p_2 - \ln p_1$
Wednesday return	$r_3 = \ln (p_3 / p_2) = \ln p_3 - \ln p_2$
Thursday return	$r_4 = \ln (p_4 / p_3) = \ln p_4 - \ln p_3$
Friday return	$r_5 = \ln (p_5 / p_4) = \ln p_5 - \ln p_4$
Return over the week	$\ln p_5 - \ln p_0 = \ln (p_5 / p_0)$

where: R_t denotes the simple return at time t , r_t denotes the continuously compounded return at time t , p_t denotes the asset price at time t and $\ln$ denotes the natural logarithm.

If the asset under consideration is a stock or portfolio of stocks, the total return to holding it is the sum of the capital gain and any dividends paid during the holding period. However, researchers often ignore any dividend payments. This is unfortunate, and will lead to an underestimation of the total returns that accrue to investors. This is likely to be negligible for very short holding periods, but will have a severe impact on cumulative returns over investment horizons of several years. Ignoring dividends will also have a distortionary effect on the cross-section of stock returns. For example, ignoring dividends will imply that ‘growth’ stocks with large capital gains will be inappropriately favoured over income stocks (e.g. utilities and mature industries) that pay high dividends.

Alternatively, it is possible to adjust a stock price time series so that the dividends are added back to generate a *total return index*. If p_t were a total return index, returns generated using either of the two formulae presented above thus provide a measure of the total return that would accrue to a holder of the asset during time t .

The academic finance literature generally employs the log-return formulation (also known as log-price relatives since they are the log of the ratio of this period’s price to the previous period’s price). Box 1.3 shows two key reasons for this.

There is, however, also a disadvantage of using the log-returns. The simple return on a portfolio of assets is a weighted average of the simple returns on the



individual assets

$$R_{pt} = \sum_{i=1}^N w_i R_{it} \tag{1.3}$$

But this does not work for the continuously compounded returns, so that they are not additive across a portfolio. The fundamental reason why this is the case is that the log of a sum is not the same as the sum of a log, since the operation of taking a log constitutes a *non-linear transformation*. Calculating portfolio returns in this context must be conducted by first estimating the value of the portfolio at each time period and then determining the returns from the aggregate portfolio values. Or alternatively, if we assume that the asset is purchased at time $t - K$ for price p_{t-K} and then sold K periods later at price p_t , then if we calculate simple returns for each period, $R_t, R_{t+1}, \dots, R_K$, the aggregate return over all K periods is

$$\begin{aligned} R_{Kt} &= \frac{p_t - p_{t-K}}{p_{t-K}} = \frac{p_t}{p_{t-K}} - 1 = \left[\frac{p_t}{p_{t-1}} \times \frac{p_{t-1}}{p_{t-2}} \times \dots \times \frac{p_{t-K+1}}{p_{t-K}} \right] - 1 \\ &= [(1 + R_t)(1 + R_{t-1}) \dots (1 + R_{t-K+1})] - 1 \end{aligned} \tag{1.4}$$

In the limit, as the frequency of the sampling of the data is increased so that they are measured over a smaller and smaller time interval, the simple and continuously compounded returns will be identical.

1.4.1 Real versus nominal series and deflating nominal series

If a newspaper headline suggests that ‘house prices are growing at their fastest rate for more than a decade. A typical 3-bedroom house is now selling for £180,000, whereas in 1990 the figure was £120,000’, it is important to appreciate that this figure is almost certainly in *nominal* terms. That is, the article is referring to the actual prices of houses that existed at those points in time. The general level of prices in most economies around the world has a general tendency to rise almost all of the time, so we need to ensure that we compare prices on a like-for-like basis. We could think of part of the rise in house prices being attributable to an increase in demand for housing, and part simply arising because the prices of all goods and services are rising together. It would be useful to be able to separate the two effects, and to be able to answer the question, ‘how much have house prices risen when we remove the effects of general inflation?’ or equivalently, ‘how much are houses worth now if we measure their values in 1990-terms?’ We can do this by *deflating* the nominal house price series to create a series of *real* house prices, which is then said to be in *inflation-adjusted terms* or *at constant prices*.

Deflating a series is very easy indeed to achieve: all that is required (apart from the series to deflate) is a *price deflator series*, which is a series measuring general price levels in the economy. Series like the consumer price index (CPI), producer price index (PPI) or the GDP Implicit Price Deflator, are often used. A more detailed discussion of which is the most relevant general price index to use is beyond the

Table 1.1 How to construct a series in real terms from a nominal one

Year	Nominal house prices	CPI (2004 levels)	House prices (2004 levels)	House prices (2013) levels
2001	83,450	97.6	85,502	105,681
2002	93,231	98.0	95,134	117,585
2003	117,905	98.7	119,458	147,650
2004	134,806	100.0	134,806	166,620
2005	151,757	101.3	149,810	185,165
2006	158,478	102.1	155,218	191,850
2007	173,225	106.6	162,500	200,850
2008	180,473	109.4	164,966	165,645
2009	150,501	112.3	134,017	173,147
2010	163,481	116.7	140,086	167,162
2011	161,211	119.2	135,244	155,472
2012	162,228	121.1	133,962	165,577
2013	162,245	123.6	131,266	162,245

Notes: All prices in British pounds; house price figures taken in January of each year from Nationwide (see appendix 1 for the source). CPI figures are for illustration only.

scope of this book, but suffice to say that if the researcher is only interested in viewing a broad picture of the real prices rather than a highly accurate one, the choice of deflator will be of little importance.

The real price series is obtained by taking the nominal series, dividing it by the price deflator index, and multiplying by 100 (under the assumption that the deflator has a base value of 100)

$$real\ series_t = \frac{nominal\ series_t}{deflator_t} \times 100$$

(1.5)

It is worth noting that deflation is only a relevant process for series that are measured in money terms, so it would make no sense to deflate a quantity-based series such as the number of shares traded or a series expressed as a proportion or percentage, such as the rate of return on a stock.

Example: Deflating house prices

Let us use for illustration a series of average UK house prices, measured annually for 2001–13 and taken from Nationwide (see Appendix 1 for the full source) given