

Министерство образования Республики Беларусь
Учреждение образования
«Белорусский государственный университет
информатики и радиоэлектроники»

Кафедра экономической информатики

А. Э. Алёхина, С. А. Поттосина

ЭКОНОМЕТРИКА

*Рекомендовано УМО по образованию в области информатики
и радиоэлектроники для направления специальности 1-40 01 02-02
«Информационные системы и технологии (в экономике)»
в качестве учебно-методического пособия*

Минск БГУИР 2013

УДК 519.862.6(076)
ББК 22.18я7
А49

Р е ц е н з е н т ы:

кафедра математического и программного обеспечения экономических систем
учреждения образования
«Гродненский государственный университет им. Я. Купалы»
(протокол №5 от 17 мая 2011 г.);

профессор кафедры теории вероятностей и математической статистики
Белорусского государственного университета,
доктор физико-математических наук, профессор Г. А. Медведев

Алёхина, А. Э.

А49 Эконометрика : учеб.-метод. пособие / А. Э. Алёхина, С. А. Поттосина. – Минск : БГУИР, 2013. – 98 с. : ил.
ISBN 978-985-488-814-9.

Пособие включает основные разделы эконометрического моделирования: модели и методы регрессионного анализа при соблюдении и нарушении традиционных модельных предположений, модели и методы анализа экономических временных рядов, системы одновременных уравнений, а также эконометрические приложения.

Адресовано студентам специальности «Информационные системы и технологии (в экономике)», а также студентам, обучающимся экономическим специальностям всех форм обучения.

УДК 519.862.6(076)
ББК 22.18я7

ISBN 978-985-488-814-9

© Алёхина А. Э., Поттосина С. А., 2013
© УО «Белорусский государственный университет информатики и радиоэлектроники», 2013

СОДЕРЖАНИЕ

Раздел 1. Особенности эконометрического анализа	5
1.1. Понятие эконометрики	5
1.2. Классификация эконометрических моделей	6
1.2.1. Типы данных	6
1.2.2. Общий вид эконометрической модели	7
1.2.3. Основные типы моделей	8
1.3. Этапы построения эконометрических моделей	9
Раздел 2. Множественный регрессионный анализ	11
2.1. Основная модель линейной регрессии	11
2.2. Метод наименьших квадратов. Теорема Гаусса – Маркова	13
2.3. Проверка качества уравнения регрессии	14
2.4. Точечный и интервальный прогноз по уравнению регрессии	18
Раздел 3. Модель множественной регрессии в условиях нарушения традиционных предпосылок	19
3.1. Методы построения модели в условиях мультиколлинеарностных факторов	19
3.1.1. Содержание и причины мультиколлинеарности	19
3.1.2. Последствия мультиколлинеарности	19
3.1.3. Методы устранения мультиколлинеарности	20
3.2. Гетероскедастичность	21
3.2.1. Содержание гетероскедастичности	21
3.2.2. Последствия гетероскедастичности	23
3.2.3. Обнаружение гетероскедастичности	23
3.2.4. Обобщенный метод наименьших квадратов	26
3.3. Автокорреляция остатков	29
3.3.1. Суть и причины автокорреляции	29
3.3.2. Последствия автокорреляции	31
3.3.3. Обнаружение автокорреляции	31
3.3.4. Оценивание в модели с авторегрессией	34
3.4. Регрессионные модели с переменной структурой	35
3.4.1. Фиктивные переменные	35
3.4.2. Тест Чоу для проверки наличия структурных изменений в регрессионных коэффициентах	38
Раздел 4. Модели и методы анализа экономических временных рядов	39
4.1. Проблемы представления экономических временных рядов	39
4.2. Использование модели регрессии с детерминированными факторами для моделирования временного ряда	43
4.3. Динамические модели с распределенными лагами	48
4.4. Модели стационарных временных рядов	56
4.4.1. Понятие стационарных и нестационарных временных рядов	56

4.4.2. Процесс авторегрессии	59
4.4.3. Процесс скользящего среднего.....	63
4.4.4. Смешанный процесс авторегрессии – скользящего среднего (процесс авторегрессии с остатками в виде скользящего среднего).....	66
4.4.5. Методы построения и тестирования модели ARMA.....	69
4.5. Модели нестационарных временных рядов.....	71
Раздел 5. Системы одновременных уравнений	78
5.1. Структурная и приведенная формы уравнений.....	79
5.2. Мультипликаторы приведенной формы	81
5.3. Идентификация систем линейных одновременных уравнений.....	83
5.4. Методы решения систем одновременных уравнений	84
Раздел 6. Практические приложения: построение и оценивание производственных функций.....	85
6.1. Свойства производственной функции	85
6.2. Оценивание параметров производственных функций. Проблемы и пути их решения	91
6.3. Способы эмпирического оценивания производственных функций разных видов	93
Литература.....	97

Раздел 1. Особенности эконометрического анализа

1.1. Понятие эконометрики

В современной научной литературе существует несколько определений эконометрики.

Наиболее точно, на наш взгляд, объяснил сущность эконометрики один из основателей этой науки Р. Фриш, который и ввел этот термин в научный обиход в 1926 г.: «Эконометрика – это не то же самое, что экономическая статистика. Она не идентична и тому, что мы называем экономической теорией, хотя значительная часть этой теории носит количественный характер. Эконометрика не является синонимом приложений математики к экономике. Как показывает опыт, каждая из трех отправных точек – статистика, экономическая теория и математика – необходимое, но не достаточное условие для понимания количественных соотношений в современной экономической жизни. Это единство всех трех составляющих. И это единство образует эконометрику»¹.

В общем случае эконометрика (Econometrics) – это самостоятельная научная дисциплина, объединяющая совокупность методов анализа связей между различными экономическими показателями (факторами) на основании реальных статистических данных с использованием аппарата теории вероятностей и математической статистики. При помощи этих методов можно выявлять новые, ранее не известные связи, уточнять или отвергать гипотезы о существовании определенных связей между экономическими показателями, предлагаемые экономической теорией.

Эконометрика сформировалась в результате синтеза трех направлений, а именно:

1) *экономической теории*. Предмет исследования – экономические явления. Но в отличие от экономической теории эконометрика делает упор на количественные аспекты, а не на качественные. Например, экономическая теория утверждает, что спрос на товар с ростом цены убывает. При этом практически не исследованным остается вопрос, как быстро и по какому закону происходит это убывание. Эконометрика отвечает на этот вопрос для каждого конкретного случая;

2) *математических методов*. Изучение экономических процессов (взаимосвязей) в эконометрике осуществляется через математические модели. Но в отличие от математической экономики, которая строит эти модели без использования реальных числовых значений, эконометрика концентрируется на изучении моделей на базе эмпирических данных;

3) *статистической теории*. Экономическая статистика обеспечивает исследователя реальными экономико-статистическими данными, обработанными и представленными в наглядной форме. Эти данные обрабатываются с помощью методов математической статистики. Однако в силу специфики статистических данных в экономике (например, в экономике невозможно проведение управляемых

¹ Frisch R. Editorial. *Econometrica*. – 1933. – №1. – P. 2.

экспериментов), людям, занимающимся эконометрикой, приходится разрабатывать свои собственные методы, которые в математической статистике не встречаются.

Таким образом, **цель эконометрики** заключается в *придании конкретного количественного выражения общим (качественным) закономерностям экономической теории на базе данных статистических наблюдений с использованием математико-статистического инструментария.*

Эконометрическое моделирование реальных социально-экономических процессов и систем обычно решает следующие прикладные задачи [11]:

- 1) анализ причинно-следственных связей между экономическими переменными;
- 2) прогноз экономических и социально-экономических показателей, характеризующих состояние и развитие анализируемой системы;
- 3) имитация различных возможных сценариев социально-экономического развития анализируемой системы (многовариантные сценарные расчеты, ситуационное моделирование).

Развитие современных информационных технологий и прикладного эконометрического программного обеспечения позволяет применять эконометрические модели и методы практически во всех направлениях экономических исследований, включая следующие:

- макроэкономика (эконометрические модели как отдельных макроэкономических показателей, так и национальных экономик в целом);
- монетарная экономика (эконометрическое моделирование денежно-кредитной системы);
- международная экономика (эконометрические модели региональной и мировой экономики);
- финансовые рынки (эконометрические модели курсов и доходностей финансовых активов и их применение для принятия оптимальных инвестиционных решений);
- микроэкономика (эконометрические модели показателей хозяйственной и финансовой деятельности фирмы).

1.2. Классификация эконометрических моделей

1.2.1. Типы данных

При моделировании различных экономических процессов возможны следующие типы данных:

- пространственные данные (*cross-sectional data*);
- временные ряды (*time series data*).

Пространственные данные – это совокупность значений $x_1, x_2, \dots, x_n$, анализируемой экономической переменной, полученных для некоторой группы из n ($n > 1$) объектов исследования в фиксированный момент (период) времени.

Примеры: показатели и экономической, и финансовой деятельности предприятий за определенный период времени, данные о расходах и доходах

потребителей за конкретный интервал времени, данные кредитной истории заемщика банка.

Временные ряды – это упорядоченное множество, характеризующее изменение показателя во времени. В отличие от пространственных данных временные ряды характеризуют динамику изменения анализируемых переменных во времени.

Значения временного ряда регистрируются с фиксированным интервалом наблюдения. По интервалу наблюдения различают такие основные типы данных, как годовые, квартальные, ежемесячные, ежедневные. Им соответствуют годовые, квартальные и другие эконометрические модели, поскольку для построения конкретной эконометрической модели используются данные с одним и тем же интервалом наблюдения. Для анализа данных типа «временной ряд» применяются методы статистического анализа временных рядов.

Примеры: годовые или квартальные значения макроэкономических показателей объема выпуска промышленной и сельскохозяйственной продукции, например, ВВП; квартальные и ежемесячные значения показателей денежно-кредитной системы, например денежных агрегатов, процентных ставок, кредиторской задолженности, индексов цен; ежедневные значения обменных курсов валют, а также курсов ценных бумаг.

1.2.2. Общий вид эконометрической модели

Всякая модель есть упрощение действительности, она описывает существенные характеристики изучаемого объекта или явления и абстрагируется от несущественных деталей. Общий вид эконометрической модели можно представить следующим образом.

Пусть состояние некоторого исследуемого объекта или процесса в момент времени t характеризуется некоторой переменной y_t , называемой эндогенной (зависимой) переменной. Состояние процесса, а следовательно, значения переменной y_t могут зависеть от различных факторов, среди которых можно выделить две группы [11]:

1) систематические контролируемые (наблюдаемые) факторы, представляемые *объясняющими переменными*, значения которых считаются известными к моменту времени t ;

2) случайные неконтролируемые факторы, приводящие к *случайным отклонениям* ε_t значений эндогенной переменной y_t от ожидаемых значений.

В свою очередь к объясняющим переменным могут относиться:

– *лаговые (предопределенные) переменные* $y_{t-1}, y_{t-2}, \dots, y_{t-l}$ ($1 \leq l \leq m$), характеризующие воздействия на исследуемый процесс со стороны внешних факторов;

– *экзогенные переменные* $x_{t1}, x_{t2}, \dots, x_{tk}$ ($1 \leq k \leq m$), характеризующие воздействия на исследуемый процесс со стороны внешних факторов.

Целью эконометрического моделирования рассматриваемого процесса является построение по эмпирическим данным $\{y_t\}, \{x_t\}$ ($t = 1, 2, \dots, n$) статистической модели зависимости переменной y_t от переменных $y_{t-1}, y_{t-2}, \dots, y_{t-L}$ и $x_{t1}, x_{t2}, \dots, x_{tm}$ вида

$$y_t = f(y_{t-1}, y_{t-2}, \dots, y_{t-L}, x_{t1}, x_{t2}, \dots, x_{tk}; \beta) + \varepsilon_t, \quad (1.1)$$

где $f(\cdot; \beta)$ – функция, определенная с точностью до неизвестных параметров β ; $\{\varepsilon_t\}$ – *случайные ошибки наблюдения* эндогенной переменной, обусловленные действием неучтенных в модели случайных нерегулярных факторов.

При построении модели вида (1.1) приходится решать следующие задачи:

- 1) выбор и экономическое обоснование вида зависимости $f(\cdot; \beta)$, а также предопределенных и экзогенных переменных;
- 2) статистическое оценивание параметров модели β ;
- 3) статистическая проверка качества построенной модели.

Модель вида (1.1) может использоваться:

- для анализа зависимости эндогенной переменной от включенных в модель экзогенных переменных;
- для прогнозирования значений эндогенной переменной y_t по заданным значениям переменной $x_{t1}, x_{t2}, \dots, x_{tk}$;
- для выбора значений $x_{t1}, x_{t2}, \dots, x_{tk}$ экзогенных переменных, обеспечивающих достижение заданных целевых значений эндогенной переменной y_t^* .

1.2.3. Основные типы моделей

В [11] предложена следующая классификация эконометрических моделей по таким признакам, как размерность модели, учет фактора времени и вид функциональной зависимости.

1. **Размерность модели** определяется числом совместно анализируемых эндогенных переменных, т. е. числом уравнений, входящих в модель. По этому признаку модели делятся на:

– *одномерные модели* (регрессионные уравнения, состоящие из одного уравнения, трендовые модели временных рядов, одномерные модели временных рядов типа авторегрессии и скользящего среднего);

– *многомерные модели*. Выделяют структурные и неструктурные модели. Структурные модели содержат эндогенные переменные как в левых, так и в правых частях уравнений, образующих некоторую систему уравнений. Неструктурные модели содержат эндогенные переменные только в левых частях уравнений.

Примером многомерной неструктурной модели может быть модель множественной линейной регрессии. В качестве многомерной неструктурной модели могут рассматриваться системы одновременных уравнений.

2. **Фактор времени**. В зависимости от того, учитывается в модели фактор времени или нет, различают:

– *статические модели*, которые включают переменные, относящиеся к одному и тому же моменту времени. Как правило, такие модели возникают при анализе пространственных данных, а также при совместном анализе коинтегрируемых временных рядов;

– *динамические модели* – модели временных рядов, включающие лаговые значения анализируемых переменных. Динамические модели позволяют анализировать динамику изменения эндогенных переменных во времени.

3. Вид зависимости. По виду функциональной зависимости модели подразделяются на:

- *линейные относительно параметров*;
- *нелинейные относительно параметров*;
- *внутренне линейные* – нелинейные по параметрам, которые можно привести к линейным с помощью определенного функционального преобразования.

Примеры моделей

Модель спроса и предложения кейнсианского типа:

$$\begin{cases} Q_t^s = \alpha_1 + \alpha_2 P_t + \alpha_3 P_{t-1} + \varepsilon_1; \\ Q_t^d = \beta_1 + \beta_2 P_t + \beta_3 Y_t + \varepsilon_2; \\ Q_t^s = Q_t^d, \end{cases}$$

где Q_t^s – спрос на товар в момент времени t ; Q_t^d – предложение товара в момент времени t ; P_t – цена товара в момент времени t ; Y_t – доход в момент времени t .

Производственная функция Кобба – Дугласа применяется для оценки эластичности выпуска продукции по отдельным факторам производства:

$$P = aL^\alpha K^{1-\alpha} \varepsilon,$$

где P – выпуск продукции; L – затраты труда; K – объем капитала; $0 < \alpha < 1$ – коэффициент эластичности.

Модель временного ряда типа ARMA:

$$x_t = \alpha_1 x_{t-1} + \alpha_2 x_{t-2} + \dots + \alpha_p x_{t-p} + \varepsilon_t + \beta_1 \varepsilon_{t-1} + \dots + \beta_q \varepsilon_{t-q}.$$

1.3. Этапы построения эконометрических моделей

Последовательность разработки эконометрических моделей можно описать следующим образом [11].

Этап 1. Экономическое обоснование модели:

- формулируются задачи и цели исследования;
- формируется состав экзогенных и эндогенных экономических переменных для включения в модель;

- оцениваются возможности получения статистических данных определенного объема и интервала наблюдения (месяц, квартал, год и т. д.) для данного набора экономических переменных;

- проводится экономический анализ априорных взаимосвязей (существующих и предполагаемых) между анализируемыми переменными, формулируется ряд гипотез и исходных допущений об исследуемом явлении и разрабатывается общая структура модели.

Этап 2. Подготовка статистических данных:

- сбор и накопление необходимого объема статистических данных;
- представление значений экономических переменных в требуемом виде (например, в номинальных или реальных ценах, в виде темпов роста или темпов прироста, переход к логарифмам и т. д.);

- предварительный анализ с целью установления особенностей моделей (например, неоднородности для пространственных данных, нестационарности, наличия выбросов, скачков, сезонных эффектов и т. д.).

Этап 3. Построение и анализ адекватности модели:

- спецификация моделей зависимости между эндогенными и экзогенными переменными, получение общего вида модельных соотношений, связывающих между собой исследуемые переменные;

- статистическое оценивание параметров моделей по имеющимся данным в рамках определенного класса моделей;

- тестирование адекватности построенных моделей на основе тестовых статистик и статистических критериев;

- тестирование гипотез, предлагаемых экономической теорией, содержательная экономическая интерпретация свойств полученных моделей и оценка возможности их использования для решения задач исследования и прогнозирования. Например, величина и знак коэффициентов в модели должны согласовываться с теорией.

Этап 4. Использование модели:

- интерпретация полученных результатов, т. е. перевод их с формализованного языка математики на содержательный язык рекомендаций по принятию управленческих решений;

- использование модели для прогноза и проведения экономической политики.

Процесс построения эконометрической модели является итерационным, допускающим возврат на более ранние этапы с целью учета новой информации и корректировки модели. Для решения задач третьего этапа часто требуется специальное программное обеспечение (Econometric Views, Statistica, STATA, SPSS, СЭМП и др.).

Раздел 2. Множественный регрессионный анализ

2.1. Основная модель линейной регрессии

Регрессионные модели применяются для исследования зависимости среднего значения анализируемой зависимой переменной от ряда независимых переменных (факторов). Модель множественной регрессии, или многомерная регрессионная модель, описывается соотношением

$$y_t = \beta_1 g_1(x_{t1}) + \beta_2 g_2(x_{t2}) + \dots + \beta_k g_k(x_{tk}) + \varepsilon_t, \quad (2.1)$$

где $\{\beta_p\} (p=1, \dots, k)$ – параметры модели; $\{y_t\}$ – значение зависимой (эндогенной) переменной в наблюдении t ; $\{x_{tp}\}$ – значение фактора x_p в наблюдении t (экзогенная переменная); ε_t – случайные ошибки наблюдения; $\{g_p(\cdot)\}$ – известные детерминированные функции. Параметры модели $\beta_1, \beta_2, \dots, \beta_k$, считаются неизвестными и подлежат оцениванию на основе методов статистического анализа.

Модель (2.1) является линейной по параметрам и может быть нелинейной по объясняющим переменным. При построении модели существенным является лишь предположение относительно линейности по параметрам, т. к. нелинейность по факторам может быть устранена с помощью замены переменных. Это позволяет рассматривать эконометрические модели в виде

$$y_t = \beta_1 x_{t1} + \beta_2 x_{t2} + \dots + \beta_k x_{tk} + \varepsilon_t. \quad (2.2)$$

В аналитических исследованиях удобно использовать векторно-матричное представление модели. Введем следующие обозначения:

$$y = \begin{pmatrix} y_1 \\ y_2 \\ \dots \\ y_n \end{pmatrix} \in R^n, \quad \beta = \begin{pmatrix} \beta_1 \\ \beta_2 \\ \dots \\ \beta_k \end{pmatrix} \in R^k, \quad \varepsilon = \begin{pmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \dots \\ \varepsilon_n \end{pmatrix} \in R^n, \quad X = \begin{pmatrix} x_{11} & x_{12} & \dots & x_{1k} \\ x_{21} & x_{22} & \dots & x_{2k} \\ \dots & \dots & \dots & \dots \\ x_{n1} & x_{n2} & \dots & x_{nk} \end{pmatrix}.$$

Пусть y обозначает $(n \times 1)$ -матрицу (вектор-столбец) $(y_1, y_2, \dots, y_n)^T$; $\beta = (\beta_1, \beta_2, \dots, \beta_k)^T$ – $(k \times 1)$ -вектор коэффициентов; $\varepsilon = (\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n)^T$ – $(n \times 1)$ -вектор ошибок, элементы которого являются случайными величинами; X – $(n \times k)$ -матрицу объясняющих переменных.

Тогда модель (2.2) можно представить в виде

$$y = X\beta + \varepsilon. \quad (2.3)$$

Множественная регрессия применяется в ситуациях, когда из множества факторов, влияющих на результативный признак, нельзя выделить один доминирующий фактор и необходимо учитывать влияние нескольких факторов. На-

пример, объем выпуска продукции определяется величиной основных и оборотных средств, численностью персонала, уровнем менеджмента и т. д., уровень спроса зависит не только от цены, но и от имеющихся у населения денежных средств.

Основная цель множественной регрессии – построить модель с несколькими факторами и определить при этом влияние каждого фактора в отдельности, а также их совместное воздействие на изучаемый показатель.

Часто полагают $x_{i1} \equiv 1$. При этом имеем модель со свободным членом:

$$y_i = \beta_1 + \beta_2 x_{i2} + \dots + \beta_k x_{ik} + \varepsilon_i.$$

В частном случае, когда используется лишь одна объясняющая переменная, имеет место модель простой (парной) линейной регрессии:

$$y_i = \beta_0 + \beta_1 x_{i1} + \varepsilon_i.$$

Построение модели по реальным данным основывается на решении следующих основных задач [11]:

- построение оценок неизвестных параметров модели $\beta_1, \beta_2, \dots, \beta_k$ с помощью методов статистического оценивания параметров по эмпирическим данным;
- проверка адекватности модели с помощью методов статистической проверки гипотез;
- выбор модели из числа альтернативных вариантов адекватных моделей на основе тестовых статистик.

Использование конкретных методов оценивания параметров и проверки гипотез зависит от «модельных» предпосылок, которые даются относительно случайных ошибок наблюдения и используемых в модели объясняющих переменных [7].

Традиционные предпосылки модели:

1. $y = X\beta + \varepsilon$ – спецификация модели выражает наше представление о механизме зависимости y от x и сам выбор объясняющих переменных.

2. Объясняющие переменные являются детерминированными, т. е. значения факторов $\{x_{ip}\}$ неслучайные (фиксированные величины). Число наблюдений $n > k + 1$, а X – матрица значений факторов. X – матрица полного ранга k , т. е. столбцы матрицы – линейно независимые векторы. В этом случае матрица является невырожденной.

3. а) математическое ожидание ошибок наблюдений равно нулю: $E(\varepsilon_i) = 0$;

б) дисперсия случайных величин $\{\varepsilon_i\}$ постоянная для всех $i = 1, 2, \dots, n$. $E(\varepsilon_i^2) = D(\varepsilon_i) = \sigma^2$. Данное свойство называется «гомоскедастичностью» (однородностью) дисперсии ошибок;

в) случайные ошибки статистически независимы (некоррелированы) для разных наблюдений: $E(\varepsilon_i, \varepsilon_s) = \text{Cov}(\varepsilon_i, \varepsilon_s) = 0$, для $i \neq s$. В случае невыполнения данного условия говорят об автокорреляции ошибок.

С учетом предпосылок а, в ковариационная матрица случайного вектора принимает диагональный вид:

$$E(\varepsilon\varepsilon') = \text{Cov}(\varepsilon, \varepsilon) = \sigma^2 I_n,$$

где I_n – единичная $(n \times n)$ -матрица.

4. Ошибки ε_t , $t = 1, 2, \dots, n$ имеют совместное нормальное распределение:

$$\varepsilon_t \sim N(0, \sigma^2).$$

Модель (2.3), удовлетворяющая приведенным предпосылкам 1–4, называется классической нормальной линейной моделью множественной регрессии. Модель (2.3), удовлетворяющая приведенным предпосылкам 1–3, называется классической линейной моделью множественной регрессии.

2.2. Метод наименьших квадратов. Теорема Гаусса – Маркова

Для нахождения оценок неизвестных параметров β и σ используются:

- метод наименьших квадратов (МНК), если выполняются условия 1–3;
- метод максимального правдоподобия, если выполняются условия 1–4.

МНК-оценка $\hat{\beta}$ вектора параметров β находится из условия минимума суммы квадратов отклонений наблюдаемых значений эндогенной переменной y_t от модельных (теоретических) значений $\hat{y}_t$, определяемых функцией регрессии (2.3): $ESS = \sum e_t^2 = e'e \rightarrow \min$.

Выразим $e'e$ через X и β :

$$\begin{aligned} e'e &= \sum_{t=1}^n (y_t - \hat{y}_t)^2 = \sum_{t=1}^n (y_t - \beta'x_t)^2 = (y - X\beta)'(y - X\beta) = \\ &= y'y - y'X\hat{\beta} - \hat{\beta}'X'y + \hat{\beta}'X'X\hat{\beta} = y'y - 2\hat{\beta}'X'y + \hat{\beta}'X'X\hat{\beta}. \end{aligned} \quad (2.4)$$

Необходимые условия минимума функции получаются дифференцированием (2.4) по вектору β :

$$\frac{\partial ESS}{\partial \hat{\beta}} = -2X'y + 2X'X\hat{\beta} = 0. \quad (2.5)$$

Учитывая обратимость матрицы $X'X$ в силу предпосылки 2, находим оценку метода наименьших квадратов:

$$\hat{\beta} = (X'X)^{-1} X'y. \quad (2.6)$$

Теорема Гаусса – Маркова. Если относительно модели (2.3) выполняются предпосылки 1–3, то оценка метода наименьших квадратов $\hat{\beta}$ является

наиболее эффективной в классе несмещенных линейных по y оценок вектора параметров β .

Таким образом, оценка $\hat{\beta}$ – наилучшая линейная несмещенная оценка вектора β . Это означает, что среди всех возможных оценок из указанного класса МНК-оценки имеют минимальные дисперсии, определяемые по формуле

$$D(\hat{\beta}_i) = \sigma^2 (X'X)^{-1}_{ii},$$

где $(X'X)^{-1}_{ii}$ – i -й диагональный элемент матрицы $(X'X)^{-1}$.

Среднеквадратичное отклонение случайной величины $\hat{\beta}_i$, определяемое как $\sqrt{D(\hat{\beta}_i)}$, называется *стандартной ошибкой* параметра $\hat{\beta}_i$.

Отклонения наблюдаемых значений эндогенной переменной y_t от «модельных» значений $\hat{y}_t$ называются остатками. Для них используют следующие обозначения: $e_t = y_t - \hat{y}_t$ – значение остатка, соответствующее t -му наблюдению; $\hat{e} = y - X\hat{\beta}$ – вектор остатков; $ESS = \sum_{t=1}^n e_t^2$ – сумма квадратов остатков (error sum of squares).

МНК-оценка s^2 дисперсии ошибки σ^2 определяется как нормированная сумма квадратов остатков $s^2 = \frac{\sum_{t=1}^n e_t^2}{n-k}$. Квадратный корень s представляет собой стандартное отклонение зависимой переменной относительно функции регрессии и называется *стандартной ошибкой регрессии*.

Перечислим некоторые свойства МНК-оценок модели [11].

Теорема. Пусть имеют место условия 1–4, тогда справедливы утверждения:

1. Случайный вектор $\hat{\beta}$ является гауссовским, причем $\hat{\beta} \sim N_k(\beta, \sigma^2 (X'X)^{-1})$.
2. Нормированная сумма квадратов имеет χ^2 -распределение с $(n-k)$ -числом степеней свободы: $\frac{e'e}{\sigma^2} \sim \chi^2(n-k)$.
3. Случайные векторы $\hat{\beta}$ и e независимы.
4. Случайный вектор $\hat{\beta}$ и случайная величина s^2 независимы.

2.3. Проверка качества уравнения регрессии

Анализ качества модели множественной регрессии включает выполнение следующих шагов:

1. Проверка статистической значимости параметров модели с целью включения в модель (или исключения из модели) тех или иных факторов.

2. Анализ доли вариации зависимой переменной, объясняемой с помощью включенных в модель факторов, а также проверка общего качества уравнения.

3. Проверка выполнимости модельных предположений относительно случайных ошибок наблюдений.

Анализ вариации зависимой переменной в регрессии. Коэффициент детерминации

Вариацию $\sum (y_t - \bar{y})$ можно разбить на две части: объясненную уравнением регрессии и необъясненную (т. е. связанную с ошибками ε).

Определим полную сумму квадратов (total sum of squares) отклонений y_t от выборочного среднего значения $\bar{y}$:

$$TSS = \sum_{t=1}^n (y_t - \bar{y})^2.$$

Для модели (2.3) полную сумму квадратов можно представить как

$$TSS = RSS + ESS,$$

где $RSS = \sum_{t=1}^n (\hat{y}_t - \bar{y})^2$ – сумма квадратов, обусловленная включенными в модель

объясняющими переменными (regression sum of squares); $ESS = \sum_{t=1}^n (y_t - \hat{y}_t)^2$ – сумма квадратов остатков (error sum of squares); $\bar{y}$ – выборочное среднее наблюдений переменной y_t .

Коэффициент детерминации определяется выражением

$$R^2 = \frac{RSS}{TSS} = 1 - \frac{ESS}{TSS}.$$

Коэффициент детерминации указывает на долю вариации зависимой переменной, объясняемую включенными в модель факторами. Он принимает значения между 0 и 1. Если $R^2 = 0$, то это означает, что модель не улучшает качество предсказания y_t по сравнению с тривиальным $\hat{y}_t = \bar{y}$. Если $R^2 = 1$, то это означает точную подгонку уравнения, т. е. все $e_t = 0$. Однако на R^2 нельзя ориентироваться как на главный критерий при сравнении двух различных структур модели. Коэффициент детерминации целесообразно использовать только совместно с дополнительным анализом регрессионного уравнения.

Коэффициент детерминации обладает рядом недостатков.

1. При добавлении в модель объясняющих переменных увеличивается значение R^2 .

2. При $k \rightarrow n$ статистика R^2 принимает значения, «неправдоподобно» близкие к 1.

3. Статистика R^2 не может использоваться для выбора модели из некоторого набора альтернативных вариантов, получающихся при преобразовании зависимой переменной.

4. Статистика R^2 принимает близкие к 1 значения при построении регрессионных моделей по нестационарным временным рядам, что затрудняет распознавание ложных регрессионных зависимостей.

При сравнении альтернативных регрессионных моделей, отличающихся количеством объясняющих переменных, предпочтительнее использовать скорректированный коэффициент детерминации:

$$R_{adj}^2 = 1 - \frac{ESS / (n - k)}{TSS / (n - 1)}.$$

Можно указать следующие свойства данной статистики:

1. $R_{adj}^2 = 1 - (1 - R^2) \frac{(n - 1)}{(n - k)}.$

2. $R^2 \geq R_{adj}^2, k > 1.$

3. $R_{adj}^2 \leq 1$, но может принимать значения < 0 .

Проверка гипотезы о значениях коэффициентов регрессии

Оценки коэффициентов регрессии зависят от используемой выборки значений переменных x и y и являются случайными величинами. Сопоставляя оценки параметров и их стандартные ошибки, можно сделать вывод о надежности (точности) полученных оценок.

Требуется проверить предположение о том, что истинное значение коэффициента регрессии равно некоторому заданному значению, т. е. выдвигается нулевая гипотеза:

$$H_0 : \beta_i = \beta_{i0}$$

при альтернативе

$$H_1 : \beta_i \neq \beta_{i0}.$$

Для тестирования нулевой гипотезы используется t -статистика:

$$t = \frac{\hat{\beta}_i - \beta_{i0}}{s_{\hat{\beta}_i}},$$

где $s_{\hat{\beta}_i}$ – среднеквадратичное отклонение оценки $\hat{\beta}_i$.

Если нулевая гипотеза верна, то статистика t имеет t -распределение Стьюдента с $(n - k)$ -степенями свободы. Гипотеза H_0 не отклоняется на заданном уровне значимости α , если $|t| \leq t_{\alpha/2}(n - k)$; здесь $t_{\alpha/2}(n - k)$ есть

100($\alpha / 2$) %-ная точка распределения Стьюдента с $(n - k)$ -степенями свободы. Здесь n – число наблюдений, k – число оцениваемых параметров.

В частном случае, когда $\beta_{i0} = 0$, а гипотезы $H_0 : \beta_i = 0$, $H_1 : \beta_i \neq 0$, соответствующая статистика принимает вид $t = \frac{\hat{\beta}_i}{s_{\hat{\beta}_i}}$. Такая гипотеза известна как гипотеза о

значимости коэффициентов регрессии. Если гипотеза H_0 отклоняется на уровне значимости α , то значение оценки $\hat{\beta}_i$ значимо отлично от нуля и, следовательно, i -й фактор оказывает существенное влияние на зависимую переменную. Это может быть основанием для ее исключения из модели. Отклонение гипотезы H_0 в пользу H_1 , наоборот, означает наличие статистически линейной зависимости между анализируемыми переменными.

t -критерий Стьюдента применяется в процедуре принятия решения о целесообразности включения фактора в модель. Если коэффициент при факторе в уравнении регрессии оказывается незначимым, то включать данный фактор в модель не рекомендуется. Отметим, что это правило не является абсолютным и бывают ситуации, когда включение в модель статистически незначимого фактора определяется экономической целесообразностью.

Интервальные оценки коэффициентов регрессии

Наряду с точечной МНК-оценкой параметров может быть найдена интервальная оценка в виде доверительного интервала $\left[\hat{\beta}_i - t_{\alpha/2}(n - k)s_{\hat{\beta}_i}, \hat{\beta}_i + t_{\alpha/2}(n - k)s_{\hat{\beta}_i} \right]$ с доверительной вероятностью $(1 - \alpha)$, определяемой соотношением

$$P(\beta_{i0} \in [\hat{\beta}_i - s_{\hat{\beta}_i} t_{\alpha/2}(n - k), \hat{\beta}_i + s_{\hat{\beta}_i} t_{\alpha/2}(n - k)]) = 1 - \alpha.$$

Такой интервал называется доверительным интервалом для β с уровнем доверия (доверительной вероятностью) $1 - \alpha$, или $(1 - \alpha)$ -доверительным интервалом, или 100($1 - \alpha$)-процентным доверительным интервалом для β .

Проверка гипотезы об адекватности модели в целом

Для модели рассмотрим задачу проверки гипотезы вида

$$H_0 : \beta_2 = \beta_3 = \dots = \beta_k = 0 \quad H_1 : \bar{H}_0$$

на заданном уровне значимости α . Здесь $\bar{H}_0$ – альтернативное утверждение по отношению к H_0 .

Статистика критерия определяется по формуле

$$F = \frac{R^2 / (k - 1)}{(1 - R^2) / (n - k)} = \frac{RSS (n - k)}{ESS (k - 1)}.$$

Гипотеза H_0 отклоняется на заданном уровне значимости α , если $F > F_c$, где $F_c = F_{\alpha}(k - 1, n - k)$ – 100 α %-ная точка распределения Фишера.

Отклонение гипотезы означает, что среди множества объясняющих факторов есть хотя бы один, который оказывает существенное влияние на эндогенную переменную. Если гипотеза не отклоняется, то следуют признать модель неадекватной, т. е. ни одна из объясняющих переменных не оказывает существенного влияния на объясняемую переменную.

2.4. Точечный и интервальный прогноз по уравнению регрессии

Одна из важнейших целей моделирования заключается в прогнозировании поведения исследуемого объекта. Как правило, термин «прогнозирование» используется в тех ситуациях, когда требуется предсказать состояние системы в будущем. Для регрессионных моделей он имеет более широкое значение. Данные могут не иметь временной структуры, однако и в этом случае может возникнуть задача оценить значение зависимой переменной для некоторого набора независимых переменных. Именно в этом смысле – смысле построения оценок зависимой переменной – следует понимать прогнозирование в эконометрике.

Предположим, что вектор независимых переменных x_{n+1} известен точно, $\hat{\beta}, \hat{s}^2$ – МНК-оценки множественной регрессии.

Точечный прогноз заключается в получении прогнозного значения y_{n+1} , которое определяется путем подстановки в уравнение регрессии

$$\hat{y} = x'_{n+1} \hat{\beta}.$$

Поскольку $E(\hat{\beta}) = \beta$, то $E(\hat{y}) = E(y_{n+1})$, т. е. оценка $\hat{y}$ является несмещенной. В классе линейных (по y) несмещенных оценок она обладает наименьшей среднеквадратичной ошибкой. Среднеквадратичная ошибка прогноза

$$E(\hat{y} - y_{n+1})^2 = \sigma^2 (1 + x'_{n+1} (X'X)^{-1} x_{n+1}).$$

Заменим σ^2 на ее оценку s^2 и обозначим

$$\delta = \sqrt{s^2 (1 + x'_{n+1} (X'X)^{-1} x_{n+1})}.$$

Если ошибки $(\varepsilon, \varepsilon_{n+1})$ имеют совместное нормальное распределение, то случайная величина $(\hat{y} - y_{n+1}) / \delta$ имеет распределение Стьюдента с $(n - k)$ -степенями свободы. Поэтому доверительным интервалом для y_{n+1} с уровнем доверия α будет интервал $(\hat{y} - \delta t_{\alpha/2}, \hat{y} + \delta t_{\alpha/2})$, где $t_{\alpha/2}$ – 100($\alpha/2$)-ная точка распределения Стьюдента с $(n - k)$ -степенями свободы.

Раздел 3. Модель множественной регрессии в условиях нарушения традиционных предпосылок

3.1. Методы построения модели в условиях мультиколлинеарностных факторов

3.1.1. Содержание и причины мультиколлинеарности

Мультиколлинеарность – это высокая степень корреляции между объясняющими переменными. Если факторы, включаемые в модель, связаны строгой линейной зависимостью, имеет место совершенная (строгая) мультиколлинеарность. В этом случае нарушается предпосылка 2 модели множественной регрессии и матрица X является вырожденной, поэтому матрица (XX) необратима, что порождает проблему идентифицируемости модели.

На практике проблема мультиколлинеарности смягчается в силу погрешности вычислений, ошибок округлений, в результате чего определитель матрицы X отличается от нуля, но принимает близкие к нему значения. В этом случае говорят, что матрица X является плохо обусловленной и имеет место *несовершенная мультиколлинеарность*. В любом случае мультиколлинеарность затрудняет разделение влияния объясняющих факторов на поведение зависимой переменной и делает оценки коэффициентов регрессии ненадежными.

Причинами мультиколлинеарности факторов являются следующие погрешности в спецификации модели:

1. Включение в модель экзогенной переменной, представляющей собой линейную комбинацию других факторов, например, суммарного показателя наряду с его составляющими.
2. Использование для моделирования сезонных изменений фиктивных переменных, число которых равно периоду сезонности.
3. Применение в модели лаговых значений экзогенных переменных.

3.1.2. Последствия мультиколлинеарности

Выделяют следующие последствия мультиколлинеарности:

1. Большие дисперсии оценок. Это затрудняет нахождение истинных значений оцениваемых величин, расширяет интервальные оценки, снижает точность оценок.
2. Уменьшаются t -статистики коэффициентов, что может привести к неоправданному выводу о существенности влияния соответствующей объясняющей переменной на зависимую переменную.
3. Оценки коэффициентов модели и их стандартные отклонения становятся очень чувствительными к малейшим изменениям выборочных данных, т. е. они становятся неустойчивыми.

4. Затрудняется определение вклада каждой из объясняющих переменных в объясняемую уравнением регрессии дисперсию зависимой переменной.

5. Возможно получение неверного знака у коэффициента регрессии.

Определение мультиколлинеарности. Установить наличие мультиколлинеарности можно по следующим признакам:

1. Малая значимость оценок коэффициентов регрессии, их большие стандартные ошибки, в то время как модель в целом является значимой, т. е. имеет высокий коэффициент детерминации R^2 и соответствующее значение F -статистики.

2. Высокие частные коэффициенты корреляции.

В последнем случае измеряется теснота линейной связи между двумя переменными, очищенная от влияния на них других факторов.

Проверка наличия мультиколлинеарности основывается на анализе матрицы парных корреляций между факторами:

$$R = \begin{bmatrix} r_{11} & r_{12} & \dots & r_{1p} \\ r_{12} & r_{22} & \dots & r_{2p} \\ \dots & \dots & \dots & \dots \\ r_{p1} & r_{p2} & \dots & r_{pp} \end{bmatrix} = \begin{bmatrix} 1 & r_{12} & \dots & r_{1p} \\ r_{12} & 1 & \dots & r_{2p} \\ \dots & \dots & \dots & \dots \\ r_{p1} & r_{p2} & \dots & 1 \end{bmatrix},$$

где r_{ij} – выборочные парные коэффициенты корреляции факторов, элементы корреляционной матрицы вектора экзогенных переменных.

3. Сильная вспомогательная (дополнительная) регрессия. Мультиколлинеарность может иметь место тогда, когда какая-либо из объясняющих переменных является линейной (или близкой к ней) комбинацией других объясняющих переменных. Для данного анализа строятся уравнения регрессии каждой из объясняющей переменной на оставшиеся независимые переменные. Вычисляются соответствующие коэффициенты детерминации R^2 и F -статистики. Если отвергается гипотеза об адекватности модели в целом ($F < F_c$), то соответствующая переменная x не является линейной комбинацией других и ее можно оставить в модели. В противном случае следует считать, что переменная x существенно зависит от остальных независимых переменных и имеет место мультиколлинеарность.

3.1.3. Методы устранения мультиколлинеарности

Вопрос об устранении мультиколлинеарности зависит от целей исследования. Если основная задача моделирования – прогноз будущих значений зависимой переменной, то при достаточно больших значениях коэффициента детерминации наличие мультиколлинеарности практически не сказывается на прогнозных качествах модели. Однако это утверждение справедливо, если линейная зависимость между зависимыми факторами сохранится в будущем.

Если же целью исследования является определение степени влияния каждой из независимых переменных на зависимую, то наличие мультиколлинеарности, приводящее к увеличению стандартных ошибок, исказит истинные зависимости между переменными. В этом случае мультиколлинеарность представляет собой серьезную проблему [2].

Можно выделить следующие методы устранения мультиколлинеарности:

1. Исключение из модели одной или ряда коррелированных переменных. Однако на практике не всегда ясно, какие из переменных являются лишними в указанном смысле. Поэтому исключение независимых переменных связано с риском возникновения ошибки спецификации.

2. Получение дополнительных данных или новой выборки. Ослабить проблему мультиколлинеарности можно за счет увеличения статистических данных. Увеличение данных сокращает дисперсии коэффициентов регрессии и тем самым увеличивает их статистическую значимость. Однако данный подход может усилить проблему автокорреляции остатков.

3. Изменение спецификации модели. Здесь речь идет либо об изменении формы модели, либо о добавлении новых объясняющих переменных, неучтенных в первоначальной модели, но существенно влияющих на зависимую переменную. Использование данного подхода уменьшает сумму квадратов отклонений, сокращая стандартную ошибку регрессии. Это приводит к уменьшению стандартных ошибок коэффициентов.

4. Переход с помощью линейного преобразования к новым некоррелирующим независимым переменным. Например, переход к главным компонентам вектора исходных объясняющих переменных (что позволяет также уменьшить количество рассматриваемых факторов), переход к последовательным разностям во временных рядах $\Delta x_{it} = x_{it} - x_{it-1}$.

3.2. Гетероскедастичность

3.2.1. Содержание гетероскедастичности

Оценки коэффициентов линейной множественной регрессии (2.3) являются эффективными (имеющими минимальную дисперсию в классе линейных несмещенных оценок) только при выполнении основных модельных предположений. Нарушение предпосылок 2, 3 ведет к утере эффективности оценок, т. е. существуют оценки с меньшей дисперсией (с меньшим разбросом значений оценок).

Выполнимость предпосылки модели (3, б) называется *гомоскедастичностью* (постоянством дисперсии отклонений). Невыполнимость данной предпосылки называется *гетероскедастичностью* (непостоянством дисперсий отклонений). Математически гомоскедастичность и гетероскедастичность определяются следующим образом:

Гомоскедастичность: $D(\varepsilon_i) = \sigma^2$, постоянная для всех наблюдений.

Гетероскедастичность: $D(\varepsilon_i) \neq \sigma_i^2$, необязательно одинакова для всех i .

Гетероскедастичность представляет собой существенную проблему, когда значения переменных в уравнении регрессии значительно различаются в разных наблюдениях. Проблема гетероскедастичности может быть проиллюстрирована на следующих примерах.

1. При изучении бюджетов потребителей можно заметить, что дисперсии остатков увеличиваются с ростом доходов.

2. При изучении зависимости между расходами на образование и ВВП в различных странах. Доля государственных расходов на образование в валовом внутреннем продукте обычно находится в диапазоне 3–9 %. Однако в абсолютном выражении в странах с большим ВВП изменение расходов на 1 % будет выражаться значительно большими цифрами, чем при малом.

3. Непостоянство дисперсии может наблюдаться при анализе временных рядов. Оно связано с эффектом масштаба.

Наличие гетероскедастичности можно наглядно видеть из поля корреляции (рис. 3.1).

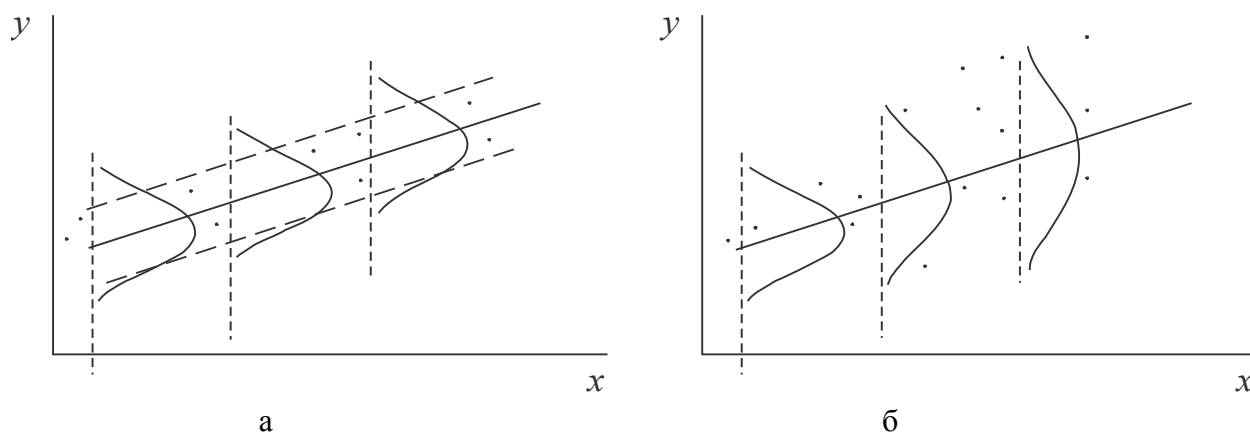


Рис. 3.1. Различия между гомоскедастичностью и гетероскедастичностью:
а – остатки гомоскедастичны; б – остатки гетероскедастичны

В обоих случаях с ростом независимой переменной растет среднее значение анализируемого признака. Но если на рис. 3.1, а дисперсия остается одной и той же для различных уровней переменной x , то на рис. 3.1, б дисперсия не остается постоянной, а увеличивается с ростом фактора.

Таким образом, проблема гетероскедастичности связана с неоднородностью исходной совокупности данных.

3.2.2. Последствия гетероскедастичности

При невыполнимости предпосылки модели (3, б) (наличие гетероскедастичности остатков) последствия применения МНК будут следующими.

1. Оценки коэффициентов по-прежнему остаются несмещенными и линейными.

2. Оценки не будут эффективными (т. е. они не будут иметь наименьшую дисперсию по сравнению с другими оценками данного параметра). Они не будут даже асимптотически эффективными. Увеличение дисперсии оценок снижает вероятность получения максимально точных оценок.

3. Дисперсии оценок будут рассчитываться со смещением. Смещенность появляется вследствие того, что необъясненная уравнением регрессии дисперсия $s^2 = \sum e_i^2 / (n - k)$ (k – число оцениваемых параметров модели), которая используется при вычислении оценок дисперсий всех коэффициентов, не является более несмещенной.

4. Вследствие вышесказанного все выводы, получаемые на основе соответствующих t - и F -статистик, а также интервальные оценки будут ненадежными. Следовательно, статистические выводы, получаемые при стандартных проверках качества оценок, могут быть ошибочными и приводить к неверным заключениям по построенной модели. Вполне вероятно, что стандартные ошибки коэффициентов будут занижены, а следовательно, t -статистики будут завышены. Это может привести к признанию статистически значимыми коэффициентов, таковыми на самом деле не являющихся.

3.2.3. Обнаружение гетероскедастичности

Иногда появление гетероскедастичности можно предвидеть заранее, основываясь на знании характера данных. В таких случаях можно предпринять соответствующие действия по устранению этого эффекта на этапе спецификации модели. Однако в большинстве случаев данную проблему приходится решать после построения уравнения регрессии. Не существует какого-либо однозначного метода определения гетероскедастичности.

Для проверки остатков на гетероскедастичность разработано большое число тестов, в которых делаются различные предположения о зависимости между величиной случайной ошибки и объясняющими переменными: графический анализ остатков, тест ранговой корреляции Спирмена, тест Глейзера, тест Голдфелда – Квандта.

Графический анализ остатков позволяет определить наличие гетероскедастичности. Для этого строятся графики зависимости e_i или e_i^2 от значений переменной x_i (либо от линейной комбинации значений переменных $y = \sum_{i=1}^k \beta_i x_i$ (рис. 3.2)).

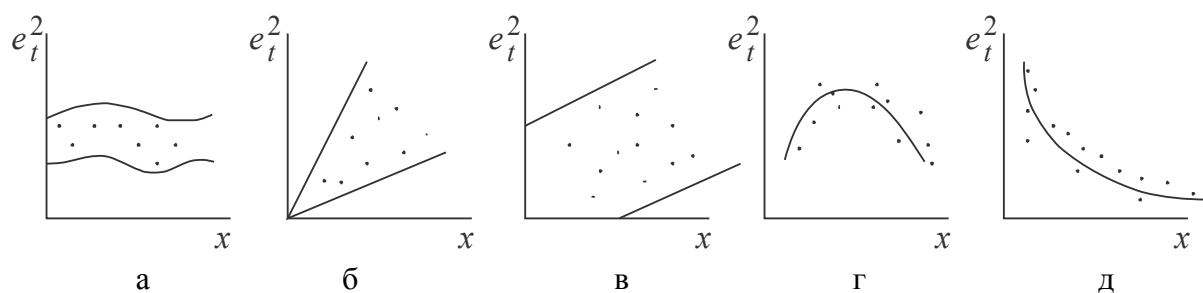


Рис. 3.2. Графический анализ остатков:
а – гомоскедастичность дисперсии шума; б–д – с большой вероятностью
существует гетероскедастичность шума

На рис. 3.2, а все отклонения находятся внутри полуполосы постоянной ширины, параллельной оси абсцисс. Это говорит о независимости дисперсий от значений переменной x и их постоянстве, т. е. в этом случае мы находимся в условиях гомоскедастичности. На рис. 3.2, б–г наблюдаются некие систематические изменения в соотношениях между значениями переменной x и квадратами отклонений. На рис. 3.2, в отражена линейная зависимость; 3.2, г – квадратичная; 3.2, д – гиперболическая зависимость между квадратами отклонений и значениями объясняющей переменной x . Другими словами, ситуации, представленные на рис. 3.2, б–д, отражают большую вероятность наличия гетероскедастичности для рассматриваемых статистических данных.

Отметим, что графический анализ отклонений является удобным и достаточно надежным в случае парной регрессии. При множественной регрессии графический анализ возможен для каждой из объясняющих переменных x_i , $i = 1, 2, \dots, (k - 1)$ отдельно. Чаще же вместо объясняющих переменных x_i по оси абсцисс откладывают значения y_i , получаемые из эмпирического уравнения регрессии. Такой анализ целесообразен при большом количестве объясняющих переменных.

Однако для строго формального вывода о наличии или отсутствии гетероскедастичности остатков следует воспользоваться формальным статистическим тестом.

Во всех тестах на гетероскедастичность проверяется нулевая гипотеза $H_0 : \sigma_1^2 = \sigma_2^2 = \dots = \sigma_n^2$ об отсутствии гетероскедастичности против альтернативной гипотезы $H_1 : \sigma_1^2 \neq \sigma_2^2 \neq \dots \neq \sigma_n^2$.

Тест ранговой корреляции Спирмена проверяет наличие монотонной зависимости между дисперсией ошибки и величиной фактора. Наблюдения (значения фактора x_i и остатка e_i) упорядочиваются по величине фактора x и вычисляется коэффициент ранговой корреляции Спирмена:

$$r_{x,e} = 1 - 6 \sum \frac{d_i^2}{n(n^2 - 1)},$$

где d_i – разность между рангами значений x_i и e_i в i -наблюдении.

Коэффициент ранговой корреляции считается значимым на уровне значимости α при $n > 10$, если выполняется условие

$$|t| = \frac{|r_{x,e}| \sqrt{n-2}}{\sqrt{1-r_{x,e}^2}} > t_{\alpha/2, n-2},$$

где $t_{\alpha/2, n-2}$ – табличное значение t -критерия Стьюдента на уровне значимости α и при числе степеней свободы $(n-2)$.

В этом случае необходимо отклонить гипотезу о равенстве нулю коэффициента корреляции $r_{x,e}$, а следовательно, и об отсутствии гетероскедастичности. В противном случае гипотеза об отсутствии гетероскедастичности не отклоняется. Если в модели регрессии больше чем одна объясняющая переменная, то проверка гипотезы может осуществляться с помощью t -статистики для каждой из них отдельно.

Тест Парка

При реализации теста Парка предполагается, что дисперсия $\sigma_i^2 = \sigma^2(e_i)$ является функцией i -го значения x_i объясняющей переменной. Р. Парк предложил следующую функциональную зависимость

$$\sigma_i^2 = \sigma^2 x_i^\beta e^{v_i}. \quad (3.1)$$

Прологарифмировав (3.1), получим

$$\ln \sigma_i^2 = \ln \sigma^2 + \beta \ln x_i + v_i.$$

Так как дисперсии σ_i^2 обычно неизвестны, то их заменяют оценками квадратов отклонений e_i^2 .

Критерий Парка включает следующие этапы:

1. Строится уравнение регрессии $y_i = \beta_0 + \beta_1 x_i + \varepsilon_i$.
2. Для каждого наблюдения определяется $\ln e_i^2 = \ln(y_i - \hat{y}_i)^2$.
3. Строится регрессия

$$\ln e_i^2 = \beta_0 + \beta_1 \ln x_i + v_i, \quad (3.2)$$

где $\beta_0 = \ln \sigma^2$.

В случае множественной регрессии зависимость (3.2) строится для каждой объясняющей переменной.

4. Проверяется статистическая значимость коэффициента β уравнения (3.2) на основе t -статистики: $t = \frac{\hat{\beta}}{S_{\hat{\beta}}}$. Если коэффициент β статистически значим,

то это означает наличие связи между $\ln e_i^2$ и $\ln x_i$, т. е. наличие гетероскедастичности в статистических данных.

Отметим, что использование в критерии Парка конкретной функциональной зависимости (3.2) может привести к необоснованным выводам (например, коэффициент β статистически незначим, а гетероскедастичность имеет место). Возможна еще одна проблема. Для случайного отклонения v_i , в свою очередь,

может иметь место гетероскедастичность. Поэтому критерий Парка дополняется другими тестами.

Тест Глейзера позволяет не только выявить наличие гетероскедастичности остатков, но и сделать определенные выводы о характере зависимости дисперсии остатков σ_i от значений фактора x_i . В тесте проверяется существование функциональной зависимости следующего вида:

$$\sigma_i = \alpha + \beta x_i^\gamma.$$

По полученным остаткам уравнения регрессии осуществляются регрессии

$$|e_i| = \alpha + \beta x_i^\gamma$$

при различных значениях параметра γ (например, -1 ; $-0,5$; $0,5$; 1 ; $1,5$; 2 ; ...) и выбирается зависимость с наиболее значимым коэффициентом β . Если все коэффициенты β незначимы, то нет оснований говорить о гетероскедастичности остатков.

Тест Голдфелда – Квандта. При проведении этого теста предполагается, что стандартное отклонение σ_i пропорционально значению некоторой независимой переменной x_i в этом наблюдении, а также, что ε_i имеет нормальное распределение и отсутствует автокорреляция остатков. Тест выполняется в следующей последовательности:

1. Все наблюдения упорядочиваются по убыванию той независимой переменной, относительно которой есть предположение на гетероскедастичность.

2. Исключаются d средних наблюдений, $d \approx \frac{1}{4}n$.

3. Проводятся две независимые регрессии первых $\frac{n}{2} - \frac{d}{2}$ наблюдений и последних $\frac{n}{2} - \frac{d}{2}$ наблюдений и строятся соответствующие остатки e_1, e_2 .

4. Вычисляется статистика $F = \frac{e_1^T e_1}{e_2^T e_2}$. Числитель и знаменатель в выра-

жении F следуют разделить на соответствующее число степеней свободы, однако для данной реализации они одинаковы.

Превышение полученной статистики критического значения на заданном уровне значимости α свидетельствует об отсутствии гомоскедастичности.

3.2.4. Обобщенный метод наименьших квадратов

Оценки (2.6) коэффициентов линейной множественной регрессии (2.3) являются эффективными (имеющими минимальную дисперсию в классе линейных несмещенных оценок) только при выполнении основных предпосылок. Нарушение предпосылок 3, б–в ведет к утере эффективности оценок (2.6), т. е. существуют оценки с меньшей дисперсией (с меньшим разбросом значений

оценок). Следствием предпосылок о постоянстве дисперсии и отсутствии автокорреляции случайных ошибок является диагональная структура матрицы ковариаций $V(\varepsilon)$ случайного члена ε_t с одинаковыми диагональными элементами σ^2 (дисперсия случайного члена ε_t):

$$E(\varepsilon\varepsilon') = \text{Cov}(\varepsilon, \varepsilon) = \sigma^2 I_n,$$

где I_n – единичная матрица размерностью n .

При нарушении указанных предпосылок $\text{Cov}(\varepsilon, \varepsilon)$ перестает иметь диагональную структуру. Обозначим ее через Ω .

Для получения эффективной оценки используется обобщенный метод наименьших квадратов (ОМНК).

Теорема Айткена. В классе линейных несмещенных оценок вектора β для обобщенной регрессионной модели оценка

$$\hat{\beta}^* = (X'\Omega^{-1}X)^{-1}X'\Omega^{-1}y$$

имеет наименьшую матрицу ковариаций.

Следует отметить, что ввиду сложности определения матрицы ковариаций Ω этот результат имеет в основном теоретический характер. Тем не менее при определенных предположениях о структуре Ω теорема имеет практическое значение [7].

Обобщенный метод наименьших квадратов в случае гетероскедастичности остатков

Предположим, что нарушается только предпосылка о постоянстве дисперсии случайного члена $\sigma_1^2 \neq \sigma_2^2 \neq \sigma^2$. При выполнении предпосылки о постоянстве дисперсии матрицы Ω и Ω^{-1} являются диагональными:

$$\Omega = \begin{bmatrix} \sigma_1^2 & 0 & \dots & 0 \\ 0 & \sigma_2^2 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & \sigma_n^2 \end{bmatrix}, \quad \Omega^{-1} = \begin{bmatrix} \frac{1}{\sigma_1^2} & 0 & \dots & 0 \\ 0 & \frac{1}{\sigma_2^2} & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & \frac{1}{\sigma_n^2} \end{bmatrix}.$$

Система нормальных уравнений ОМНК (2.6) имеет вид

$$(X'\Omega^{-1}X)\beta = X'\Omega^{-1}Y, \quad (3.3)$$

или в координатной форме

$$\sum \frac{y_t}{\sigma_t^2} = \beta_1 \sum \frac{1}{\sigma_t^2} + \beta_2 \sum \frac{x_{2t}}{\sigma_t^2} + \beta_3 \sum \frac{x_{3t}}{\sigma_t^2} + \dots + \beta_k \sum \frac{x_{kt}}{\sigma_t^2},$$

$$\sum \frac{y_t x_{2t}}{\sigma_t^2} = \beta_1 \sum \frac{x_{2t}}{\sigma_t^2} + \beta_2 \sum \frac{x_{2t}^2}{\sigma_t^2} + \beta_3 \sum \frac{x_{3t} x_{2t}}{\sigma_t^2} + \dots + \beta_k \sum \frac{x_{kt} x_{2t}}{\sigma_t^2},$$

.....

$$\sum \frac{y_t x_{kt}}{\sigma_t^2} = \beta_1 \sum \frac{x_{kt}}{\sigma_t^2} + \beta_2 \sum \frac{x_{2t} x_{kt}}{\sigma_t^2} + \beta_3 \sum \frac{x_{3t} x_{kt}}{\sigma_t^2} + \dots + \beta_k \sum \frac{x_{kt}^2}{\sigma_t^2}. \quad (3.4)$$

Система уравнений (3.4) соответствует модели, определяемой соотношениями

$$\frac{y_t}{\sigma_t} = \beta_1 \frac{1}{\sigma_t} + \beta_2 \frac{x_{2t}}{\sigma_t} + \beta_3 \frac{x_{3t}}{\sigma_t} + \dots + \beta_k \frac{x_{kt}}{\sigma_t} + u_t, \quad (3.5)$$

которые получаются, если исходное уравнение множественной регрессии (2.3), записанное для каждого наблюдения, разделить на среднеквадратичное отклонение σ_t случайного члена ε_t в t -наблюдении. Случайный член $u_t = \frac{\varepsilon_t}{\sigma_t}$ в модели

(3.5) имеет постоянную для всех наблюдений дисперсию $\sigma_{\varepsilon t}^2 = 1$.

Запись модели в виде (3.5) соответствует уравнению линейной множественной регрессии (без свободного члена)

$$y' = \beta_1 x'_1 + \beta_2 x'_2 + \dots + \beta_k x'_k + u,$$

записанному в новых переменных $y', x'_1, x'_2, \dots, x'_k$, значения которых определяются по формулам

$$y'_t = \frac{y_t}{\sigma_t}, \quad x'_{1t} = \frac{1}{\sigma_t}, \quad x'_{2t} = \frac{x_{2t}}{\sigma_t}, \dots, \quad x'_{kt} = \frac{x_{kt}}{\sigma_t}, \quad u_t = \frac{\varepsilon_t}{\sigma_t}.$$

Следует отметить, что величины σ_t^2 практически никогда не известны и вместо них следует использовать состоятельные оценки $\hat{\sigma}_t^2$.

При практическом использовании ОМНК используется какое-либо предположение относительно зависимости дисперсии σ_t^2 случайного члена ε от наблюдения или величины факторов x_t .

Метод взвешенных наименьших квадратов

Данный метод основывается на знании ковариационной матрицы вектора ошибок. Предположим, что ковариационная матрица ошибок диагональная, $D(\varepsilon_t) = \sigma_t^2$, кроме того, $\sigma_t^2 = \sigma^2 \omega_t$, где числа ω_t нормированы таким образом, что $\sum_t \omega_t = n$. Вспомогательная система получается делением каждого уравнения на соответствующее значение σ_t :

$$\frac{y_t}{\sigma_t} = \sum_{j=1}^k \beta_j \frac{x_{tj}}{\sigma_t} + u_t,$$

где $u_t = \frac{\varepsilon_t}{\sigma_t}$ некоррелированы и имеют постоянную дисперсию. Применяя к преобразованной системе стандартный метод наименьших квадратов, получаем наилучшие линейные несмещенные оценки. Содержательный смысл такого преобразования очевиден: при взвешивании каждого наблюдения с помощью коэффициента $\frac{1}{\sigma_t}$ устраняется неоднородность данных: наблюдениям с меньшей дисперсией придается больший вес и наоборот. Поэтому данный подход называют методом взвешенных наименьших квадратов.

Стандартное отклонение ошибки пропорционально независимой переменной. Если можно считать, что стандартное отклонение ошибки прямо пропорционально одной из независимых переменных $\sigma_t^2 = \sigma^2 x_{tk}^2$, то применяются следующие преобразования: $x_{ij}^* = \frac{x_{ij}}{x_{ik}}$, $y_t^* = \frac{y_t}{x_{tk}}$. МНК-оценки коэффициентов преобразованной системы дают непосредственно оценки исходной модели.

3.3. Автокорреляция остатков

3.3.1. Суть и причины автокорреляции

На практике может нарушаться предпосылка о независимости значений случайного члена ε_t и ε_s в различных наблюдениях $\text{Cov}(\varepsilon_t, \varepsilon_s) = 0$ ($t \neq s$). В этом случае говорят, что случайные ошибки – *автокоррелированные*.

Автокорреляция (последовательная корреляция) определяется как корреляция между наблюдаемыми показателями, упорядоченными во времени (временные ряды) или в пространстве (пространственные данные). Автокорреляция случайных ошибок чаще встречается в регрессионном анализе при использовании данных временных рядов. Однако нередко проявления данного явления и при моделировании на пространственных данных.

Случайная ошибка в уравнении регрессии подвергается воздействию тех переменных, влияющих на зависимую переменную, которые не включены в уравнение регрессии. Если значение ε_t в любом наблюдении должно быть независимым от его значений в предыдущем наблюдении, то и значение любой переменной, «скрытой» в ε_t , должно быть некоррелируемо с ее значением в предыдущем наблюдении.

Постоянная направленность воздействия невключенных в уравнение переменных является наиболее частой причиной положительной автокорреляции (рис. 3.3) – ее обычный тип для экономического анализа ($\text{Cov}(\varepsilon_{t-1}, \varepsilon_t) > 0$).

Пример. Спрос на мороженое обусловлен состоянием погоды, которое в свою очередь является скрытым в ε фактором. Вероятно, при измерении будем иметь несколько последовательных наблюдений, когда теплая погода будет

благоприятствовать увеличению спроса на мороженое, после этого – несколько последовательных наблюдений противоположного характера.

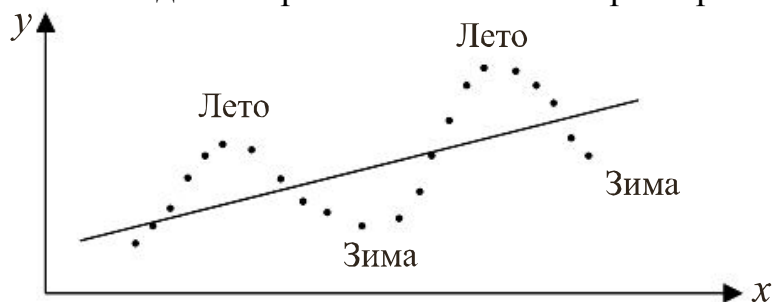


Рис. 3.3. Положительная автокорреляция

Отрицательная автокорреляция (рис. 3.4) означает, что за положительным отклонением следует отрицательное и наоборот ($\text{Cov}(\varepsilon_{t-1}, \varepsilon_t) < 0$). Отрицательная автокорреляция встречается редко. Но иногда она проявляется при преобразовании первоначальной спецификации модели.

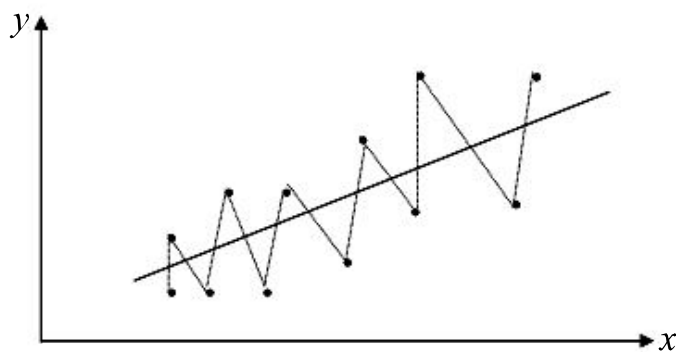


Рис. 3.4. Отрицательная автокорреляция

Среди основных причин, вызывающих появление автокорреляции, можно выделить ошибки спецификации, инерцию в изменении экономических показателей, эффект паутины, сглаживание данных [2].

Ошибки спецификации. Отсутствие в модели какой-либо важной объясняющей переменной либо неправильный выбор формы зависимости обычно приводит к системным отклонениям точек наблюдений от линии регрессии, что может привести к автокорреляции.

Инерция. Многие экономические показатели (например, инфляция, безработица, ВВП и т. п.) обладают определенной цикличностью, связанной с волнообразностью деловой активности. Действительно, экономический подъем приводит к росту занятости, сокращению инфляции, увеличению ВВП и т. д. Этот рост продолжается до тех пор, пока изменение конъюнктуры рынка и ряда экономических характеристик не приведет к замедлению роста, затем – остановке и движению вспять рассматриваемых показателей. В любом случае эта трансформация происходит не мгновенно, а обладает определенной инертностью.

Эффект паутины. Во многих сферах экономики экономические показатели реагируют на изменение экономических условий с запаздыванием (вре-

менным лагом). Например, предложение сельскохозяйственной продукции реагирует на изменение цены с запаздыванием (равным периоду созревания урожая). Большая цена сельскохозяйственной продукции в прошлом году вызовет (скорее всего) ее перепроизводство в текущем году, а следовательно, цена на нее снизится и т. д. В этой ситуации нельзя предполагать случайность отклонений друг от друга.

Сглаживание данных. Зачастую данные по некоторому продолжительному временному периоду получают усреднением данных по составляющим его подинтервалам. Это может привести к определенному сглаживанию колебаний, которые имелись внутри рассматриваемого периода, что в свою очередь может послужить причиной автокорреляции.

3.3.2. Последствия автокорреляции

При применении МНК можно выделить следующие последствия автокорреляции.

1. Оценки параметров остаются линейными и несмещенными, но перестают быть эффективными. Следовательно, они перестают обладать свойствами наилучших линейных несмещенных оценок.

2. Дисперсии оценок являются смещенными, чаще заниженными, что приводит к увеличению t -статистик. Это может привести к признанию статистически значимыми объясняющие переменные, которые в действительности таковыми могут и не являться.

3. Оценка дисперсии регрессии $\hat{\sigma}^2 = \frac{\sum e_t^2}{n - k}$ является смещенной оценкой истинного значения σ^2 , во многих случаях заниженной.

4. В силу вышесказанного выводы по t - и F -статистикам, определяющим значимость коэффициентов регрессии и коэффициента детерминации, возможно, будут неверными. Вследствие этого ухудшаются прогнозные качества модели.

3.3.3. Обнаружение автокорреляции

Выводы о независимости случайных ошибок наблюдений осуществляются на основе оценок e_t , полученных из эмпирического уравнения регрессии. Для обнаружения автокорреляции случайных ошибок используют следующие методы: графический анализ остатков, тест Дарбина – Уотсона.

Графический метод

Одним из способов графического обнаружения автокорреляции является построение последовательно-временных графиков. С этой целью строится график зависимости e_t от его порядкового номера i (рис. 3.5)

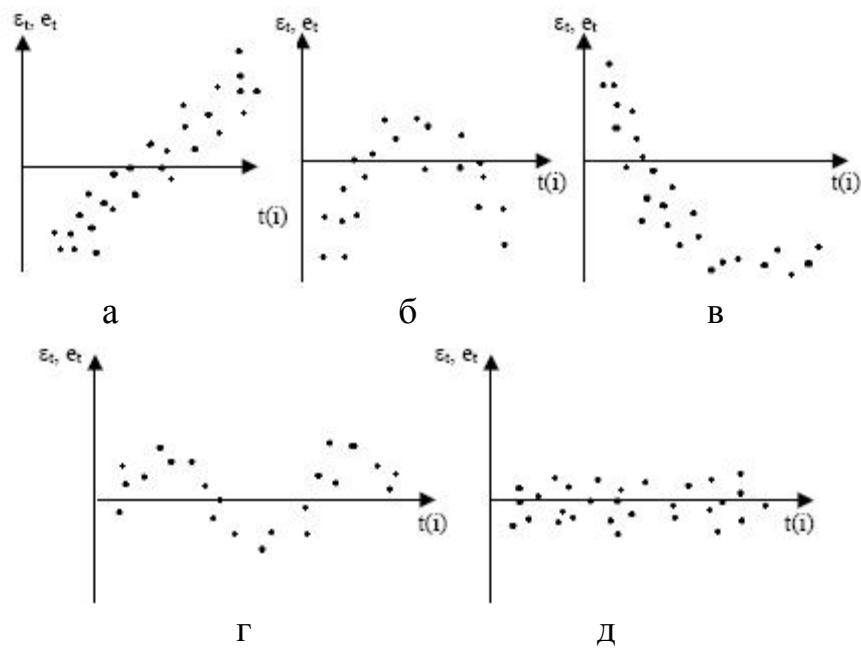


Рис. 3.5. Графический анализ остатков

Можно предположить, что на рис. 3.5, а–г имеются определенные связи между отклонениями, т. е. автокорреляция имеет место. Отсутствие систематической связи между остатками на рис. 3.5, д скорее всего свидетельствует об отсутствии автокорреляции.

Графический анализ остатков не является самостоятельным и убедительным методом обнаружения автокорреляции, а может только иллюстрировать результаты формальных статистических тестов.

Тест Дарбина – Уотсона

Данный тест используется для обнаружения автокорреляции первого порядка. Большинство формальных тестов используют следующую идею: если автокорреляция есть у ошибок ε , то она присутствует и в остатках e , получаемых после применения к модели метода наименьших квадратов. Во всех формальных тестах проверяется гипотеза об отсутствии автокорреляции: $H_0: \rho = 0$. Альтернативным считается предположение H_1 о коррелированности остатков.

Предположим, что остатки в уравнении регрессии $y = X\beta + \varepsilon$ образуют авторегрессионный процесс первого порядка:

$$\varepsilon_t = \rho \varepsilon_{t-1} + v_t,$$

где $|\rho| < 1$ – некоторый параметр, называемый коэффициентом авторегрессии, а $\{v_t\}$ ($E(v_t) = 0, D(v_t) = E(vv^T) = \sigma_v^2$) – независимые в совокупности случайные величины, имеющие одинаковое нормальное распределение $N(0, \sigma^2)$.

Тест Дарбина – Уотсона основан на статистике:

$$DW = \frac{\sum_{t=2}^n (e_t - e_{t-1})^2}{\sum_{t=1}^n e_t^2}.$$

Можно проверить, как эта статистика связана с выборочным коэффициентом корреляции между e_t и e_{t-1} .

$$\begin{aligned} DW &= \frac{\sum_{t=2}^n (e_t - e_{t-1})^2}{\sum_{t=1}^n e_t^2} = \frac{\sum_{t=2}^n e_t^2 + \sum_{t=2}^n e_{t-1}^2 - 2 \sum_{t=2}^n e_t e_{t-1}}{\sum_{t=1}^n e_t^2} = \\ &= \frac{\sum_{t=1}^n e_t^2 - e_1^2 + \sum_{t=1}^n e_t^2 - e_n^2 - 2 \sum_{t=2}^n e_t e_{t-1}}{\sum_{t=1}^n e_t^2} = 2 \left(1 - \frac{\sum_{t=2}^n e_t e_{t-1}}{\sum_{t=1}^n e_t^2} \right) - \frac{e_1^2 + e_n^2}{\sum_{t=1}^n e_t^2}. \end{aligned}$$

Предполагая число наблюдений достаточно большим, можно считать, что приближенно выполнены следующие равенства: $\frac{1}{n-1} \sum_{t=2}^n e_t = -e_1 / (n-1) \approx 0$ и

$\frac{1}{n-1} \sum_{t=2}^n e_{t-1} = -e_n / (n-1) \approx 0$. Поэтому выборочный коэффициент корреляции r между e_t и e_{t-1} можно приближенно представить в виде

$$r \approx \frac{\sum_{t=2}^n e_t e_{t-1}}{\sqrt{\sum_{t=2}^n e_t^2 \sum_{t=1}^{n-1} e_t^2}} \approx \frac{\sum_{t=2}^n e_t e_{t-1}}{\sum_{t=1}^n e_t^2}.$$

Пренебрегая слагаемыми e_1^2, e_n^2 по сравнению с общей суммой $\sum_{t=1}^n e_t^2$, окончательно получим

$$DW \approx 2(1 - r).$$

Содержательный смысл статистики заключается в следующем. Если между e_t и e_{t-1} существует высокая положительная корреляция, то величина статистики DW мала. Значения статистики, близкие к 2, могут означать отсутствие автокорреляции. Если коэффициент r близок к -1 , то величина DW стремится к 4 и можно сделать вывод о наличии отрицательной автокорреляции. Для обоснованного вывода об автокорреляции на заданном уровне значимости α использует пороговые значения d_l, d_u , вычисленные для некоторых фиксирован-

ных значений числа наблюдений n и количества экзогенных переменных в исходной модели. В соответствии с тестом Дарбина – Уотсона принимаются следующие решения:

Таблица 3.1

Принятие решений по статистике Дарбина – Уотсона

Значение статистики DW	Вывод
$4 - d_l < DW < 4$	Гипотеза H_0 отвергается, есть отрицательная корреляция
$4 - d_u < DW < 4 - d_l$	Неопределенность
$d_u < DW < 4 - d_u$	Гипотеза H_0 не отвергается
$d_l < DW < d_u$	Неопределенность
$0 < DW < d_l$	Гипотеза H_0 отвергается, есть положительная корреляция

Наличие зоны неопределенности представляет определенные трудности при использовании теста. Однако она может быть устранена путем увеличения объемов выборочных данных. В этом случае зона неопределенности может быть уменьшена.

Существуют определенные ограничения по использованию теста Дарбина – Уотсона.

1. Критерий DW применяется лишь для тех моделей, которые содержат свободный член.

2. Тест Дарбина – Уотсона построен в предположении о некоррелированности регрессоров x и случайных ошибок ε .

3. Предполагается, что случайные отклонения ε_t определяются авторегрессионной схемой первого порядка.

4. Статистические данные должны иметь одинаковую периодичность (т. е. не должно быть пропусков в наблюдениях).

5. Критерий Дарбина – Уотсона не применим для регрессионных моделей, содержащих лаговые значения эндогенных переменных.

Следует иметь в виду, что отсутствие автокорреляции первого порядка не исключает возможность автокорреляции более высокого порядка. По этим причинам для анализа автокоррелированности остатков модели по временным рядам целесообразен анализ автокорреляционной и частной автокорреляционной функций, а также анализ статистики Льюнга – Бокса, целью которого является проверка гипотезы о том, что остатки описываются случайным процессом белого шума.

3.3.4. Оценивание в модели с авторегрессией

На практике значение коэффициента авторегрессии, как правило, не известно. Поэтому применяются приближенные методы оценивания, например,

итеративная процедура Кохрейна – Оркатта. Начальным шагом является применение обычного МНК к модели множественной регрессии и получение остатков $\{e_t, t = \overline{1, n}\}$. Далее процедура выполняется в следующей последовательности:

1) в качестве приближенного значения ρ берется его МНК-оценка в уравнении (2.1);

2) проводится преобразование $y_t - \rho y_{t-1} = (x_t - \rho x_{t-1})^T \beta + v_t$. Игнорирование первого наблюдения может привести к потере важной информации, поэтому для $t = 1$ используется поправка Прайса – Уинстена: $\sqrt{1 - \rho^2} y_1 = \sqrt{1 - \rho^2} x_1^T \beta + \sqrt{1 - \rho^2} \varepsilon_1$. Затем находятся МНК-оценки параметров β ;

3) вычисляется новый вектор остатков e_t ;

4) процедура повторяется возвращением к п. 1 и продолжается до тех пор, пока последовательные оценки параметров ρ не будут отличаться меньше любого заранее заданного числа, либо фиксируется количество итераций.

3.4. Регрессионные модели с переменной структурой

3.4.1. Фиктивные переменные

Независимые переменные в регрессионных моделях имеют «непрерывные» области изменений (национальный доход, размер заработной платы, объем производства). Однако теория не накладывает никаких ограничений на характер регрессоров, в частности, некоторые переменные могут принимать только два значения, в более общей ситуации – дискретное множество значений. Необходимость рассматривать такие переменные возникает в тех случаях, когда требуется принимать во внимание влияние качественного фактора (фактора, не имеющего количественного выражения), например, пол потребителя, фактор сезонности, наличие государственных программ.

Чтобы учесть влияние качественного фактора в рамках одного регрессионного уравнения, вводятся так называемые *фиктивные переменные* (*dummy variables*) с двумя значениями 0 и 1. Например, изучается зависимость потребления товара y от величины дохода x с учетом пола потребителя. С использованием фиктивной переменной d :

$$d_t = \begin{cases} 0, & \text{потребитель – мужчина;} \\ 1, & \text{потребитель – женщина.} \end{cases}$$

Уравнение регрессии принимает вид

$$y_t = \beta_0 + \beta_1 x_t + \gamma d_t + \varepsilon_t.$$

Тогда ожидаемое значение объемов потребления в зависимости от величины доходов будет

$$E(y|x, d_i = 0) = \beta_0 + \beta_1 x_i - \text{мужчины,}$$

$$E(y|x, d_i = 1) = \beta_0 + \gamma + \beta_1 x_i - \text{женщины.}$$

Потребление в данном случае является линейной функцией от дохода (рис. 3.6). Причем и для мужчин, и для женщин потребление меняется с одним и тем же коэффициентом пропорциональности β_1 . Свободные члены в моделях отличаются на величину γ . Проверка значимости коэффициентов при фиктивном факторе d_i покажет значимость влияния качественного показателя на изучаемый признак и необходимость включения в уравнение регрессии соответствующего члена.

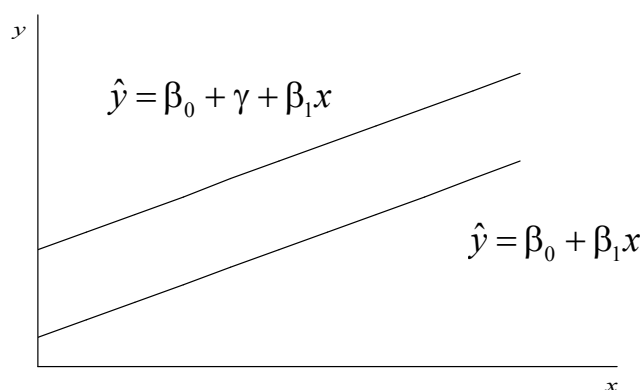


Рис. 3.6. Зависимость потребления товара от величины дохода и пола

Если включаемый в рассмотрение качественный признак имеет не два, а несколько значений, то можно было бы ввести дискретную переменную, принимающую такое же количество значений. Но этого, как правило, не делают из-за сложности дальнейшей интерпретации коэффициентов. В этом случае целесообразнее использовать несколько бинарных переменных. Отметим, что в данной ситуации следует придерживаться следующего правила.

Если качественная переменная имеет k альтернативных значений, то при моделировании используются только $(k - 1)$ фиктивных переменных.

В противном случае сумма фиктивных переменных равна константе, и, как следствие, становится затруднительным оценивание модели с помощью МНК; такая ситуация называется «ловушкой фиктивной переменной» (*dummy trap*).

Фиктивные переменные широко используются для моделирования сезонных колебаний. Для этого вводят три фиктивные переменные:

$d_{i1} = 1$, если наблюдение относится к зимнему периоду, $d_{i1} = 0$ – иначе,

$d_{i2} = 1$, если наблюдение относится к весеннему периоду, $d_{i2} = 0$ – иначе,

$d_{i3} = 1$, если наблюдение относится к летнему периоду, $d_{i3} = 0$ – иначе,

переменная для осеннего периода, исключаемая из модели, будет представлять собой эталонную или базовую категорию. Выбор эталонной категории не ока-

зывает воздействия на сущность уравнения регрессии. Но от этого зависит, какие тесты можно будет выполнить.

Уравнение регрессии, отражающее зависимость потребления y от некоторой количественной переменной x и учитывающее сезонные изменения, примет вид

$$y_t = \beta_0 + \beta_1 x_{t1} + \gamma_1 d_{t1} + \gamma_2 d_{t2} + \gamma_3 d_{t3} + \varepsilon_t.$$

Коэффициенты $\gamma_i, i = 1, 2, 3$ показывают средние сезонные отклонения в объеме потребления по отношению к осенним месяцам.

Иногда необходимо ввести в модель регрессии две и более группы фиктивных переменных. Предположим, что имеется набор статистических данных о бюджетах потребителей за каждый из четырех кварталов. При этом потребление задается функцией, учитывающей сезонные колебания и различия в социальных группах, а также некоторые другие экономические переменные.

Если принять во внимание четыре сезона и три социальные группы, то можно задать эту зависимость в виде

$$q_t = \beta_0 + \beta_1 x_{t1} + \beta_2 x_{t2} + \beta_3 x_{t3} + \gamma_2 Q_{t2} + \gamma_3 Q_{t3} + \gamma_4 Q_{t4} + \delta_2 S_2 + \delta_3 S_3 + \varepsilon_t,$$

где $Q_i = 1$, если наблюдение относится к кварталу $i = 2, 3, 4$; $S_j = 1$, если наблюдение относится к социальным группам $j = 2, 3$, x_i – множество экономических переменных. Среднее ожидаемое потребление в зависимости от принадлежности к социальной группе и времени года может быть выражено следующим образом:

$$E(q|1 \text{ квартал}, 1 \text{ соц. группа}) = \beta_0 + \beta_1 x_{t1} + \beta_2 x_{t2} + \beta_3 x_{t3};$$

$$E(q|1 \text{ квартал}, 2 \text{ соц. группа}) = \beta_0 + \delta_2 + \beta_1 x_{t1} + \beta_2 x_{t2} + \beta_3 x_{t3};$$

$$E(q|1 \text{ квартал}, 3 \text{ соц. группа}) = \beta_0 + \delta_3 + \beta_1 x_{t1} + \beta_2 x_{t2} + \beta_3 x_{t3};$$

$$E(q|2 \text{ квартал}, 1 \text{ соц. группа}) = \beta_0 + \gamma_2 + \beta_1 x_{t1} + \beta_2 x_{t2} + \beta_3 x_{t3};$$

$$E(q|2 \text{ квартал}, 2 \text{ соц. группа}) = \beta_0 + \gamma_2 + \delta_2 + \beta_1 x_{t1} + \beta_2 x_{t2} + \beta_3 x_{t3} \text{ и т. д.}$$

Фиктивные переменные позволяют строить и оценивать кусочно-линейные модели, которые можно применять для исследования структурных изменений. Пусть y – зависимая переменная, объем продукции, выпущенной в период t ; x – независимый признак, отражающий размер основного фонда предприятия в этот же период времени; x, y представлены в виде временных рядов. Из априорных соображений считаем, что в период времени t_0 произошла структурная перестройка, и линия регрессии будет отличаться от первоначальной, но общая линия останется непрерывной (рис. 3.7).

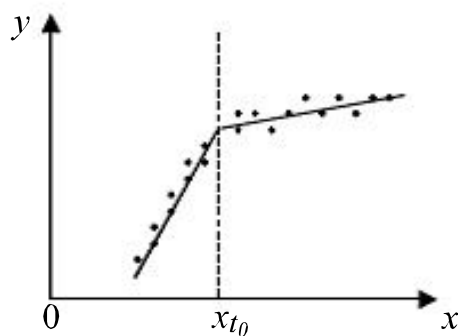


Рис. 3.7. Структурные изменения

Для оценивания такой модели вводится бинарная переменная r : $r_t = 0$, если $t \leq t_0$, и $r_t = 1$, если $t > t_0$. Регрессионное уравнение имеет вид

$$y_t = \beta_0 + \beta_1 x_t + \beta_2 (x_t - x_{t_0}) r_t + \varepsilon_t.$$

Регрессионная линия, соответствующая периоду $t \leq t_0$, имеет коэффициент наклона β_1 и $\beta_1 + \beta_2$ для $t > t_0$. Тестируя гипотезу $\beta_2 = 0$, проверяем предположение о том, что фактически структурного изменения не произошло.

Можно отметить следующие преимущества использования фиктивных переменных:

1. Простой способ проверки того, является ли воздействие качественного фактора значимым.
2. Иногда регрессионные оценки оказываются более эффективными.
3. Описание экономического объекта (явления) только одним уравнением.

3.4.2. Тест Чоу для проверки наличия структурных изменений в регрессионных коэффициентах

Тест Г. Чоу изначально разработан для тех случаев, когда имеется две выборки значений, одна из которых (объемом n) получена при одних условиях, другая (m) – при несколько измененных. При этом необходимо выяснить, действительно ли выборки *однородны* (гипотеза H_0). Критерий также может быть использован при построении регрессионных моделей при воздействии *качественных* признаков, когда имеется возможность разделения совокупности наблюдений по степени воздействия признака на отдельные группы и требуется установить возможность использования единой модели регрессии.

Механизм теста: по каждой выборке строятся регрессионные модели:

$$y_t = \beta'_1 + \sum_{i=2}^k \beta'_i \cdot x_{ti} + \varepsilon'_t, \quad t = \overline{1, n};$$

$$y_t = \beta''_1 + \sum_{i=2}^k \beta''_i \cdot x_{ti} + \varepsilon''_t, \quad t = \overline{n+1, n+m}.$$

Проверяемая гипотеза имеет вид $H_0: \beta' = \beta''; D(\varepsilon') = D(\varepsilon'') = \sigma^2$.

Если гипотеза H_0 верна, то модели можно объединить в одну:

$$y_t = \beta_1 + \sum_{i=2}^k \beta_i \cdot x_{ti} + \varepsilon_t, \quad t = \overline{1, n+m}.$$

Пусть суммы квадратов остатков частных регрессий – ESS_1 и ESS_2 ; ESS_p – сумма квадратов остатков в объединенной регрессии. Равенство между ESS_p и $ESS_1 + ESS_2$ будет иметь место только при совпадении коэффициентов в объединенной и частных моделях регрессий. В общем случае при разделении выборки будет наблюдаться изменение качества уравнения, что можно представить как $ESS_p - (ESS_1 + ESS_2)$.

Нулевая гипотеза H_0 отвергается на уровне значимости α , если выполняется следующее условие:

$$F = \frac{ESS_p - (ESS_1 + ESS_2)}{ESS_1 + ESS_2} \cdot \frac{(n+m-2k)}{k} > F_\alpha(k, n+m-2k).$$

В противном случае гипотеза о совпадении моделей для обеих выборок не отклоняется, т. е. влияние качественного признака на исследуемую выборку несущественно, и их можно объединить в одну.

Раздел 4. Модели и методы анализа экономических временных рядов

4.1. Проблемы представления экономических временных рядов

В сфере экономики встречаются явления или события, которые интересно и важно изучать в их динамике, т. к. они изменяются во времени. Также во времени изменяются цены, экономические условия, режим протекания того или иного производственного процесса. Совокупность измерений подобного рода показателей в течение некоторого периода времени и представляет *временной ряд*.

Цели анализа временных рядов могут быть различными. Можно, например, стремиться предсказать будущее на основании знаний прошлого, пытаться выяснить механизм, лежащий в основе процесса, и управлять им. Необходимо уметь освобождать временной ряд от компонент, которые затемняют его динамику. Часто требуется сжато представлять характерные особенности ряда.

Под *временным рядом* (ВР) понимается упорядоченное множество, характеризующее изменение показателя во времени. Элементы этого множества состоят из численных значений показателя, называемых уровнями ВР, и периодов (моментов, интервалов времени), к которым относятся уровни [5].

Различают два вида временных рядов. Измерение некоторых величин (температуры, напряжения и т. д.) производится непрерывно, по крайней мере теоретически. При этом наблюдения можно фиксировать в виде графика. Но даже в том случае, когда изучаемые величины регистрируются непрерывно,

практически при их обработке используются только те значения, которые соответствуют дискретному множеству моментов времени. Следовательно, если время измеряется непрерывно, временной ряд называется **непрерывным**, если же время фиксируется дискретно, т. е. через фиксированный интервал времени, то временной ряд **дискретен**.

В дальнейшем речь будет идти только о дискретных временных рядах. Дискретные временные ряды получаются двумя способами:

- выборкой из непрерывных временных рядов через регулярные промежутки времени (например, численность населения, величина собственного капитала фирмы, объем денежной массы, курс акции); такие временные ряды называются **моментными**;

- накоплением переменной в течение некоторого периода времени (например, объем производства какого-либо вида продукции, количество осадков, объем импорта); в этом случае временные ряды называются **интервальными**.

В эконометрике принято моделировать временной ряд как **случайный процесс**, называемый также **стохастическим процессом**, под которым понимается статистическое явление, развивающееся во времени согласно законам теории вероятностей. Случайный процесс – это случайная последовательность.

Возможные значения временного ряда в данный момент времени t описываются с помощью случайной величины x_t и связанного с ней распределения вероятностей $p(x_t)$. Тогда наблюдаемое значение x_t временного ряда в момент t рассматривается как одно из множества значений, которые могла бы принять случайная величина x_t в этот момент времени. Следует отметить, однако, что, как правило, наблюдения временного ряда взаимосвязаны, и для корректного его описания следует рассматривать совместную вероятность $p(x_1, \dots, x_T)$, где T – некоторый фиксированный базисный период.

Важнейшим условием правильного формирования временных рядов является сопоставимость уровней, образующих ряд. Уровни ряда, подлежащие изучению, должны быть *однородны по экономическому содержанию* и учитывать существо изучаемого явления и цель исследования. Статистические данные, представленные в виде временных рядов, должны быть *сопоставимы по территории, кругу охватываемых объектов, единицам измерения, моменту регистрации, методике расчета, ценам, достоверности* [1].

Несопоставимость по территории возникает в результате изменений границ стран, регионов, хозяйств и т. п. Для приведения данных к сравнимому виду производится пересчет прежних данных с учетом новых границ.

Полнота охвата различных частей явления – важнейшее условие сопоставимости уровней ряда. Требование одинаковой полноты охвата разных частей изучаемого объекта означает, что уровни ряда за отдельные периоды должны характеризовать размер того или иного явления по одному и тому же кругу входящих в его состав частей.

При определении сравниваемых уровней ряда необходимо использовать *единую методику* их расчета. Особенно часто эта проблема возникает при международных сопоставлениях.

Несопоставимость показателей возникает в силу неодинаковости применяемых *единиц измерения*. С различием применяемых единиц измерения приходится встречаться при изучении динамики: производственных ресурсов, когда они представляются то в стоимостном, то в трудовом исчислении; энергетических мощностей (кВт/ч); атмосферного давления и т. д. Трудности при сравнении данных *по моменту регистрации* возникают из-за сезонных явлений. Уровни при сравнении должны относиться к определенной дате ежегодно.

При анализе показателей в стоимостном выражении следует учитывать, что с течением времени происходит непрерывное *изменение цен*. Причин у этого процесса множество – инфляция, рост затрат, рыночные условия (спрос и предложение) и т. д. В этой связи при характеристике стоимостных показателей объема продукции во времени должно быть устранено влияние изменения цен. Для решения этой задачи количество продукции, произведенное в разные периоды, оценивают в ценах одного периода, которые называют *фиксированными* или в определенных статистических органах – *сопоставимыми* ценами.

Широкое использование в статистических исследованиях выборочного метода требует учитывать *достоверность* количественных и качественных характеристик изучаемых явлений в динамике. Различная репрезентативность выборки по периодам внесет существенные погрешности в величины уровней ряда.

Одним из условий сопоставимости уровней интервального ряда, кроме равенства периодов, за которые приводятся данные, является *однородность этапов*, в пределах которых показатель подчиняется одному закону развития. В этих случаях проводят периодизацию временных рядов, типологическую группировку во времени.

Все вышеназванные обстоятельства следует учитывать при подготовке информации для анализа изменений явлений во времени (динамике).

При решении практических задач эконометрического моделирования часто переходят от уровней временного ряда к экономическим индексам. Формы представления экономических индексов определяют классификацию временных рядов по следующим видам [5]:

- в уровнях (объемном выражении) x_t (t – интервал времени) – значения показателя выражены в некоторых единицах измерения и имеют некоторый масштаб (например, ВВП выражается в миллиардах рублей, экспорт сырой нефти – в тоннах);

- в темпах роста (в индексах) I_t , с использованием данных по отношению к фиксированному периоду T :

$$I_t = \frac{x_t}{x_T}; \quad (4.1)$$

- в темпах роста (в индексах) I_t , с использованием данных x_t по отношению к предыдущему периоду:

$$I_t = \frac{x_t}{x_{t-1}}; \quad (4.2)$$

– в темпах роста I_t , с использованием данных по отношению к соответствующему периоду предыдущего года:

$$I_t = \frac{x_t}{x_{t-k}},$$

где k – число периодов в году ($k = 12$ – для месячных и $k = 4$ – для квартальных данных);

– в темпах роста $I_{i,m}$, с использованием данных нарастающим итогом с начала текущего года по отношению к данным нарастающим итогом с начала предыдущего года:

$$I_{i,m} = \frac{\sum_{j=ik+1}^{ik+m} x_j}{\sum_{j=(i-1)k+1}^{(i-1)k+m} x_j},$$

где $m = 1, 2, \dots, k$, i – номер года, k – число периодов в году.

Временные ряды в уровнях обладают наиболее полной информацией, необходимой для анализа и построения адекватных эконометрических моделей прогнозирования.

Говорят, что показатель (4.1) представлен в базисной форме. Он характеризует динамику соотношений x_t между различными текущими периодами t и некоторым фиксированным базисным периодом T ; является безразмерным, имеет заданный масштаб и обеспечивает сопоставимость динамики ВР. Даже в том случае, когда их уровни выражены в разных единицах измерения.

Говорят, что показатель (4.2) представлен в *цепной форме* в виде темпа роста. Так же, как и базисная, цепная форма является безразмерной величиной и может использоваться для сопоставления разных показателей. При этом переход к ней устраняет экспоненциальный тренд временного ряда. Однако при использовании ряда в темпах роста возникают проблемы, связанные с его интерпретацией и переводом в ряд в уровнях.

4.2. Использование модели регрессии с детерминированными факторами для моделирования временного ряда

Экономические ВР, как и ряды во многих других областях, можно рассматривать в виде совокупности составляющих динамики, т. е. как функцию нескольких рядов. Наблюдаемый ВР предполагается состоящим из ненаблюдаемых составляющих динамики, которые в первом приближении можно считать независимыми. К основным составляющим динамики экономических ВР прежде всего следует отнести трендовую (эволюторную), циклическую, календарную, сезонную и нерегулярную составляющие.

Эти виды динамики могут комбинироваться. Тем самым задается **разложение временного ряда** на составляющие, которые с экономической точки зрения несут разную содержательную нагрузку. Наиболее важные из них:

- **тренд** (эволюторная составляющая) – соответствует медленному изменению, происходящему в некотором направлении, которое сохраняется в течение значительного промежутка времени;

- **циклические колебания** – это более быстрая, чем тенденция, квазипериодическая динамика, в которой имеются фазы возрастания и убывания. Наиболее часто цикл связан с флуктуациями экономической активности;

- **сезонные колебания** – соответствуют изменениям, которые происходят регулярно в течение года, недели или суток. Они связаны с сезонами и ритмами человеческой активности;

- **календарные эффекты** – это отклонения, связанные с определенными предсказуемыми календарными событиями, – такими, как праздничные дни, количество рабочих дней за месяц, високосность года и т. п.;

- **случайная составляющая** – беспорядочные движения относительно большой частоты. Они порождаются влиянием разнородных событий на изучаемую величину (несистематический или случайный эффект). Часто такую составляющую называют **шумом**;

- **выбросы** – это аномальные движения временного ряда, связанные с редко происходящими событиями, которые резко, но лишь очень кратковременно отклоняют ряд от общего закона, по которому он движется;

- **структурные сдвиги** – это аномальные движения временного ряда, связанные с редко происходящими событиями, имеющие скачкообразный характер и меняющие тенденцию.

Первые три составляющие объединяют в одну детерминированную и рассматривают временной ряд в виде

$$x_t = f(t, \beta) + \varepsilon_t,$$

где $f(t, \beta)$ – детерминированная функция времени t , называемая трендом, заданная с точностью до неизвестных параметров β ; $\{\varepsilon_t\}$ – случайные ошибки наблюдений с нулевым математическим ожиданием $E(\varepsilon_t) = 0$ и постоянной дисперсией $D\{\varepsilon_t\} = \sigma^2$.

Анализ временного ряда заключается в выделении и изучении указанных компонент ряда в рамках одной из моделей ряда: аддитивной или мультипликативной.

Модели трендов

Для построения трендов часто применяются следующие функции.

Простая линейная модель:

$$x_t = \beta_0 + \beta_1 t.$$

Этот тип тренда подходит для отображения тенденции примерно равномерных изменений уровней: равных в среднем абсолютных приростов или абсолютных сокращений уровней за равные промежутки времени.

Полиномиальная модель:

$$x_t = \beta_0 + \beta_1 t + \beta_2 t^2 + \dots + \beta_n t^n.$$

Значение степени полинома n в практических задачах редко превышает 5. В частности, тренд, выраженный параболой 2-го порядка, имеет вид

$$x_t = \beta_0 + \beta_1 t + \beta_2 t^2.$$

Параболы 3-го и более высоких порядков редко применимы для выражения тенденции динамики и слишком сложны для получения надежных оценок параметров при ограниченной длине временного ряда.

Значения параметров параболы 2-го порядка таковы: свободный член β_0 – это средний (выровненный) уровень тренда на момент или период, принятый за начало отсчета времени, т. е. $t = 0$; β_1 – это средний за весь период среднегодовой прирост, который уже не является константой, а изменяется равномерно со средним ускорением, равным $2\beta_2$, которое и служит константой, главным параметром параболы 2-го порядка.

Следовательно, тренд в форме параболы 2-го порядка применяется для отображения таких тенденций динамики, которым свойственно примерно постоянное ускорение абсолютных изменений уровней. С одной стороны, процессы такого рода встречаются на практике гораздо реже, чем процессы с равномерным изменением, но, с другой стороны, любое отклонение процесса от строго равномерного прироста (или сокращения) уровней можно интерпретировать как наличие ускорения.

Экспоненциальная модель:

$$x_t = \beta_0 + \beta_1 \exp(\beta_2 t).$$

Свободный член экспоненты β_0 равен выровненному уровню, т. е. уровню тренда в момент или период, принятый за начало отсчета времени, т. е. при $t = 0$. Основным параметр экспоненциального тренда k является постоянным темпом изменения уровней (ценным). Если $k > 1$, имеем тренд с возрастающими уровнями, причем это возрастание не просто ускоренное, а с возрастающим ускорением и возрастающими производными всех более высоких порядков. Если $k < 1$, то имеем тренд, выражающий тенденцию постоянного, но замедляющегося сокращения уровней, причем замедление непрерывно усиливается.

Экспоненциальный тренд характерен для процессов, развивающихся в среде, не создающей никаких ограничений для роста уровня. Из этого следует, что на практике он может развиваться только на ограниченном промежутке времени, т. к. любая среда рано или поздно создает ограничения, любые ресурсы со временем исчерпаемы.

Гиперболическая модель:

$$x_t = \beta_0 + \frac{\beta_1}{t} \text{ — наиболее простой вид гиперболы.}$$

Если основным параметр гиперболы $\beta_1 > 0$, то этот тренд выражает тенденцию замедляющегося снижения уровней, и при $t \rightarrow \infty, x_t \rightarrow \beta_0$. Таким образом, свободный член гиперболы — это предел, к которому стремится уровень тренда.

Такая тенденция наблюдается при изучении процесса снижения затрат любого ресурса (труда, материалов, энергии) на единицу данного вида продукции или ее себестоимости в целом. Затраты ресурса не могут стремиться к нулю, значит, экспонента не соответствует сущности процесса; нужно применить гиперболическую формулу тренда.

Если параметр $\beta_1 < 0$, то с возрастанием t , т. е. с течением времени, уровни тренда возрастают и стремятся к величине β_0 при $t \rightarrow \infty$.

Логарифмическая модель:

$$x_t = \beta_0 + \beta_1 \ln t.$$

Такая форма тренда характерна для процесса, приводящего к замедлению роста какого-то показателя, но при этом рост не прекращается, не стремится к какому-либо ограниченному пределу.

Примером тенденций, соответствующих логарифмическому тренду, может служить динамика рекордных достижений в спорте: известно, что увеличение на 1 см рекорда прыжка в высоту или снижение на 0,1 с времени бега на 200 или 400 м требует все больших и больших затрат времени, каждый рекорд дается все большим и большим трудом. В то же время нет и «вечных» рекордов, все спортивные достижения улучшаются, но медленнее и медленнее, т. е.

по логарифмическому тренду. Нередко такой же характер динамики присущ на отдельных этапах развития динамике урожайности или валового сбора какой-то культуры в данном регионе, пока новое агротехническое достижение не придаст снова тенденции ускорения, что иллюстрирует следующая модель.

Логистическая модель:

$$x_t = \frac{a}{1 + be^{-ct}}.$$

Логистическая форма тренда подходит для описания такого процесса, при котором изучаемый показатель проходит полный цикл развития, начиная, как правило, от нулевого уровня, сначала медленно, но с ускорением возрастая, затем ускорение становится нулевым в середине цикла, т. е. рост происходит по линейному тренду, затем, в завершающей части цикла, рост замедляется по гиперболе по мере приближения к предельному значению показателя. К S-образным процессам можно отнести процесс развития новой отрасли (нового производства).

При выборе модели тренда можно руководствоваться следующими рекомендациями [11].

1. Если значения времени t представляют собой арифметическую прогрессию, а первые разности соответствующих значений временного ряда x_t постоянны, то тренд может быть описан линейной моделью.

2. Если значения времени t представляют собой арифметическую прогрессию, а p -е разности соответствующих значений временного ряда x_t постоянны, то тренд может быть описан полиномом степени p .

3. Если связь между $\ln(x_t)$ и t линейна, то описание тренда целесообразно производить по степенной модели.

Проверка наличия тренда

Для проверки гипотезы о наличии тренда применяются непараметрические критерии, например, знаковый критерий Кокса – Стюарта, а также метод сравнения средних и метод Фостера – Стюарта.

Для временного ряда x_t , $t = 1, \dots, n$, проверяется гипотеза об отсутствии тренда:

$$H_0 : E\{x_t\} = a = \text{const}$$

против альтернативы

$$H_1 : |E\{x_{t+1}\} - E\{x_t\}| \neq 0,$$

т. е. гипотеза о существовании систематического смещения среднего.

Для проверки гипотезы H_0 , H_1 используется следующий алгоритм.

1. Временной ряд разбивается на три части $n = n_1 + n_2 + n_3$ так, что первая и последняя содержат одинаковое число наблюдений $n_1 = n_3$.

2. Пусть $x_i^{(1)}$ и $x_i^{(3)}$ – элементы первой и последней третей временного ряда. Вычисляется разность $x_i^{(3)} - x_i^{(1)}$ и фиксируется ее знак. Пусть n_1^+ – число положительных разностей, n_1^- – число отрицательных разностей:

$$S = \max\{n_1^+, n_1^-\}.$$

Принимается решение о том, что тренд возрастающий, если $n_1^+ > n_1^-$ (убывающий, если $n_1^+ < n_1^-$).

3. Статистика S при $n > 30$ аппроксимируется нормальным распределением со средним $T/6$ и дисперсией $T/12$, тогда статистика

$$T^* = \frac{|S - n/6|}{\sqrt{n/12}}$$

подчинена стандартному нормальному закону распределения $N(0,1)$. При малых объемах выборки ($n \leq 30$) вводится поправка на непрерывность:

$$T^* = \frac{|S - n/6| - 0,5}{\sqrt{n/12}}.$$

4. Для двустороннего критерия гипотеза об отсутствии тренда принимается, если $|T^*| < u_{1-\alpha/2}$, где $u_{1-\alpha/2}$ – квантиль уровня $1-\alpha/2$ стандартного нормального распределения.

Метод сравнения средних. Временной ряд разбивается на две примерно равные части $x_1, x_2, \dots, x_{n_1}$ и $x_{n_1+1}, x_{n_2+2}, \dots, x_{n_1+n_2}$ с количеством уровней n_1 и n_2 , и для каждой части вычисляются средние $(\bar{x}_1, \bar{x}_2)$ и выборочные дисперсии (s_1^2, s_2^2) соответственно.

Далее рассчитывается значение критерия Стьюдента по формуле

$$t = \frac{|\bar{x}_1 - \bar{x}_2|}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}},$$

если предполагается, что значения дисперсий на этих участках не равны между собой, т. е. $\sigma_1^2 \neq \sigma_2^2$ и по формуле

$$t = \frac{|\bar{x}_1 - \bar{x}_2|}{s^2} \sqrt{\frac{n_1 \cdot n_2}{n_1 + n_2}},$$

где s^2 – общая выборочная дисперсия ряда, если предполагается, что дисперсии одинаковы: $\sigma_1^2 = \sigma_2^2$.

Методы построения и тестирования модели

Процесс построения моделей с детерминированным трендом предусматривает процедуру выделения тренда. При выделении полиномиального тренда к регрессионной модели применяется метод наименьших квадратов.

Анализ адекватности модели основан на использовании стандартного набора тестовых статистик с целью решения следующих задач [11]:

- анализ значимости коэффициентов регрессии и адекватности модели в целом (t- и F-статистика, коэффициент детерминации R^2 , сумма квадратов остатков);
- проверка предположения о том, что остатки являются «белым шумом», т. е. некоррелированными (Q-статистика Льюнга – Бокса) и имеют нормальный закон распределения (критерий χ^2 Пирсона, Жака-Бера);
- выбор наиболее простой модели (статистика AIC, SC).

В случае нелинейного тренда, допускающего линеаризацию за счет специальных функциональных преобразований, также применяется МНК по отношению к преобразованному временному ряду. В противном случае используют нелинейный метод наименьших квадратов.

Нестационарные временные ряды, которые после выделения детерминированного тренда с помощью МНК приводят к стационарному ряду остатков, называются «стационарными относительно тренда».

4.3. Динамические модели с распределенными лагами

При изучении поведения экономических процессов на достаточно длительном промежутке времени есть все основания предполагать наличие определенных взаимосвязей между их последовательными состояниями, т. е. состояние экономического явления в данный момент или период времени определяется в том числе и его состояниями, а также состояниями окружающей среды в предшествующие моменты или периоды времени. Данное обстоятельство является следствием наличия запаздывания в действии факторов либо инерционностью изучаемых процессов.

Модели, связывающие состояния экономических явлений в последовательные моменты (периоды) времени, принято называть *динамическими*. Такие модели позволяют изучать явления в динамике, в развитии. Аналитическое

представление динамических моделей включает значения переменных, относящиеся как к текущему, так и к предыдущим моментам (периодам) времени.

Обычно динамические модели подразделяют на два класса.

1. Эконометрические модели, включающие в качестве факторов значения факторных переменных в предыдущие моменты времени, называются *моделями с распределенным лагом*.

$$y_t = \alpha + \beta_0 x_t + \beta_1 x_{t-1} + \beta_2 x_{t-2} + \dots + \beta_p x_{t-p} + \varepsilon_t. \quad (4.3)$$

Моделями этого типа описываются ситуации, когда влияние причины (независимых факторов) на следствие (зависимую переменную) проявляется с некоторым запаздыванием. Например, при изучении зависимости объемов выпуска продукции от величины инвестиций, выручки от расходов на рекламу и т. п.

2. Эконометрические модели, включающие в качестве факторов наряду со значениями независимых переменных также значения результативной переменной в предыдущие моменты времени – *модели авторегрессии*.

$$y_t = \alpha + \beta_0 x_t + \beta_1 x_{t-1} + \beta_2 x_{t-2} + \dots + \gamma_1 y_{t-1} + \gamma_2 y_{t-2} + \dots + \varepsilon_t. \quad (4.4)$$

Отдельную группу динамических моделей составляют модели, учитывающие ожидаемые уровни переменных. Эти уровни определяются экономическими субъектами на основе информации, которой субъекты располагают в текущий и предыдущий моменты времени. Например, модели адаптивных ожиданий или частичной корректировки.

Включенные в модель в качестве факторов значения переменных в предыдущие моменты времени называются *лаговыми переменными*. Примером значений лаговых переменных являются временные ряды исходных уровней, сдвинутые назад на один или более моментов времени. Величина этого сдвига называется *лагом*.

Причин наличия лагов в экономике достаточно много, и среди них можно выделить следующие [2].

Психологические причины, которые обычно выражаются через инерцию в поведении людей. Например, люди тратят свой доход постепенно, а не мгновенно. Привычка к определенному образу жизни приводит к тому, что люди приобретают те же блага в течение некоторого времени даже после падения их реального дохода.

Технологические причины. Например, изобретение персональных компьютеров не привело к мгновенному вытеснению ими больших ЭВМ в силу необходимости замены соответствующего программного обеспечения, которое потребовало продолжительного времени.

Институциональные причины. Например, контракты между фирмами, трудовые договоры требуют определенного постоянства в течение времени контракта (договора).

Механизмы формирования экономических показателей. Например, инфляция во многом является инерционным процессом; денежный мультипликатор (создание денег в банковской системе) также проявляет себя на определенном временном интервале и т. д.

Включение в эконометрическую модель лаговых значений зависимой переменной осложняет проблему получения несмещенных и эффективных оценок ее параметров. Во-первых, наличие нескольких лаговых переменных $y_{t-1}, y_{t-2}, \dots$, либо $x_{t-1}, x_{t-2}, \dots$, зачастую сильно коррелирующих между собой, ведет к потере качества модели вследствие ухудшения точности оценок ее параметров, снижению их эффективности и устойчивости к незначительным колебаниям исходной информации, ошибкам округления.

Во-вторых, как правило, существует сильная корреляционная зависимость между переменными $y_{t-1}, y_{t-2}, \dots$ и ошибкой ε_t , ведущая к появлению смещения в оценках параметров при использовании МНК.

В-третьих, временной ряд ошибки модели ε_t часто характеризуется наличием автокорреляционной связи, вследствие чего оценки параметров модели, полученные непосредственно на основе МНК, являются неэффективными. Отметим, что важным этапом при построении моделей с распределенным лагом и моделей авторегрессии является выбор оптимальной величины лага и определение его структуры.

Для оценки моделей с бесконечным числом лагов разработано несколько методов.

Метод последовательного увеличения количества лагов

По данному методу уравнения рекомендуется оценивать с последовательно увеличивающимся количеством лагов. Признаков завершения процедуры увеличения количества лагов может быть несколько:

- при добавлении нового лага какой-либо коэффициент регрессии β_k при переменной x_{t-k} меняет знак. Тогда в уравнении регрессии оставляют переменные $x_t, x_{t-1}, \dots, x_{t-k+1}$, коэффициенты при которых знак не поменяли;
- при добавлении нового лага коэффициент регрессии β_k при переменной x_{t-k} становится статистически незначимым. Очевидно, что в уравнении будут использоваться только переменные $x_t, x_{t-1}, \dots, x_{t-k+1}$, коэффициенты при которых остаются статистически значимыми.

Однако применение этого метода весьма ограничено в силу постоянно уменьшающегося числа степеней свободы, сопровождающегося увеличением стандартных ошибок и ухудшением качества оценок, а также возможности мультиколлинеарности. Кроме этого, при неправильном определении количества лагов возможны ошибки спецификации.

Модели с распределенным лагом

Рассмотрим модель с распределенным лагом порядка p :

$$y_t = \alpha + \beta_0 x_t + \beta_1 x_{t-1} + \beta_2 x_{t-2} + \dots + \beta_p x_{t-p} + \varepsilon_t.$$

Основную проблему при оценке параметров составляет, как правило, сильная корреляция между факторами $x_t, x_{t-1}, x_{t-2}, \dots$. Для ее преодоления применяется преобразование лаговых переменных, либо делаются определенные предположения о характере коэффициентов регрессии.

Оценка параметров модели с распределенным лагом методом Койка

В методе Койка предполагается, что коэффициенты при лаговых переменных убывают в геометрической прогрессии:

$$\beta_k = \beta_0 \lambda^k, k = 0, 1, \dots,$$

где $0 < \lambda < 1$ характеризует скорость убывания коэффициентов с увеличением лага (с удалением от момента анализа). Такое предположение достаточно логично, если считать, что влияние прошлых значений объясняющих переменных на текущее значение зависимой переменной будет тем меньше, чем дальше по времени эти показатели имели место.

В данном случае уравнение (4.3) преобразуется в уравнение

$$y_t = \alpha + \beta_0 x_t + \beta_0 \lambda x_{t-1} + \beta_0 \lambda^2 x_{t-2} + \dots + \varepsilon_t. \quad (4.5)$$

Модель (4.5) включает три параметра α, β, λ , для определения которых применяется нелинейный метод наименьших квадратов. Согласно этому методу:

а) задаются границы изменения параметра λ (в простейшем случае 0 и 1) и определяются всевозможные значения параметра λ с достаточно малым шагом (например 0,01). Для каждого значения параметра λ вычисляются значения новой переменной z :

$$z_t = x_t + \lambda x_{t-1} + \lambda^2 x_{t-2} + \lambda^3 x_{t-3} + \dots + \lambda^p x_{t-p},$$

где величина p выбирается такой, чтобы воздействием последующих лаговых значений x_{t-p+i} можно было пренебречь;

б) оценивается уравнение регрессии

$$y_t = \alpha + \beta_0 z_t + \varepsilon_t; \quad (4.6)$$

в) выбирается такое значение параметра λ , которому соответствует наибольший коэффициент детерминации R^2 при оценке уравнения (4.6). Получен-

ные при этом оценки параметров a и b принимаются в качестве оценок параметров исходного уравнения (4.5).

Однако более распространенной является схема вычислений на основе *преобразования Койка*.

Вычитая из уравнения (4.5) такое же уравнение, но умноженное на λ и вычисленное для предыдущего периода времени $(t - 1)$

$$\lambda y_{t-1} = \lambda \alpha + \beta_0 \lambda x_{t-1} + \beta_0 \lambda^2 x_{t-1} + \dots + \lambda \varepsilon_{t-1}, \quad (4.7)$$

получим следующее уравнение:

$$y_t - \lambda y_{t-1} = (1 - \lambda)\alpha + \beta_0 x_t + (\varepsilon_t - \lambda \varepsilon_{t-1}),$$

далее

$$y_t = (1 - \lambda)\alpha + \beta_0 x_t + \lambda y_{t-1} + v_t, \quad (4.8)$$

где $v_t = \varepsilon_t - \lambda \varepsilon_{t-1}$ – *скользящая средняя между ε_t и ε_{t-1}* .

Преобразование уравнения (4.3) по данному методу в уравнение (4.8) называется *преобразованием Койка*.

Полученное уравнение представляет собой авторегрессионную модель первого порядка. Оценив параметры этого уравнения, можно получить оценки параметров α , β , λ исходного уравнения (4.5).

При применении преобразования Койка возможны следующие проблемы:

- среди объясняющих переменных появляется переменная y_{t-1} , которая, в принципе, носит случайный характер, что нарушает одну из предпосылок МНК. Кроме того, данная объясняющая переменная скорее всего коррелирует со случайным отклонением v_t ;
- если для случайных отклонений ε_t , ε_{t-1} исходной модели выполняется предпосылка 30 МНК, то для случайных отклонений v_t , очевидно, имеет место автокорреляция. Для ее анализа вместо обычной статистики DW Дарбина – Уотсона необходимо использовать статистику Дарбина;
- при указанных выше проблемах оценки, полученные по МНК, являются смещенными и несостоятельными.

Оценка параметров модели с распределенным лагом методом Алмон

В методе Алмон для преодоления сильной корреляции между факторами x_t , x_{t-1} , x_{t-2} , ... используется переход к $(k + 1)$ новым переменным z_j с меньшей корреляционной зависимостью по формулам

$$z_{jt} = a_{j0} x_t + a_{j1} x_{t-1} + a_{jp} x_{t-p}, \quad (4.9)$$

где коэффициенты подобраны соответствующим образом.

Согласно методу Алмон, коэффициенты представляют в виде полиномов заданной степени k от величины лага j :

$$b_j = c_0 + c_1j + c_2j^2 + \dots + c_kj^k.$$

В частности:

- для полинома первой степени (при $k = 1$): $b_j = c_0 + c_1j$;
- для полинома второй степени (при $k = 2$): $b_j = c_0 + c_1j + c_2j^2$ и т. д.

Выражения для коэффициентов b_j модели (4.3) принимают вид:

$$\begin{aligned} b_0 &= c_0; \\ b_1 &= c_0 + c_1 + \dots + c_k; \\ b_2 &= c_0 + 2c_1 + 4c_2 \dots + 2^k c_k; \\ &\dots \\ b_p &= c_0 + pc_1 + p^2c_2 \dots + p^k c_k. \end{aligned} \tag{4.10}$$

Подставляя в (4.3) найденные соотношения для b_j и вводя новые переменные z_j с помощью соотношения (4.9), представим исходное уравнение (4.3) в следующем виде:

$$y_t = a + c_0z_0 + c_1z_1 + c_2z_2 + \dots + c_kz_k + \varepsilon_t, \tag{4.11}$$

где

$$\begin{aligned} z_0 &= x_t + x_{t-1} + \dots + x_{t-p} = \sum_{j=0}^p x_{t-j}; \\ z_1 &= x_{t-1} + 2x_{t-2} + \dots + px_{t-p} = \sum_{j=1}^p jx_{t-j}; \\ z_2 &= x_{t-1} + 4x_{t-2} + \dots + p^2x_{t-p} = \sum_{j=1}^p j^2x_{t-j}; \\ &\dots \\ z_k &= x_{t-1} + 2^k x_{t-2} + \dots + p^k x_{t-p} = \sum_{j=0}^p j^k x_{t-j}. \end{aligned} \tag{4.12}$$

После определения численных значений параметров c_j модели (4.11) коэффициенты исходной модели b_j находятся из соотношений (4.10). Применение метода Алмон для расчета параметров модели с распределенным лагом предполагает предварительное определение максимальной величины лага p и степени полинома k . Оптимальную величину лага можно приближенно определить на

основе априорной информации экономической теории или проведенных ранее эмпирических исследований. Приближенно в качестве величины лага можно взять значение максимального лага, для которого парный коэффициент корреляции между y и лаговыми переменными $x_t, x_{t-1}, x_{t-2}, \dots$ остается значимым.

Можно также построить несколько уравнений регрессии с разной величиной лага и выбрать наилучшее. Что касается степени полинома k , то на практике обычно ограничиваются рассмотрением полиномов второй и третьей степени. Величину k также можно определять путем сравнения моделей, построенных для различных значений k . Следует отметить, что при наличии сильной корреляционной связи между исходными лаговыми переменными $x_t, x_{t-1}, x_{t-2}, \dots$ переменные z_j , представляющие собой их линейные комбинации, также будут коррелировать между собой. Однако коэффициенты в формулах (4.12) подобраны таким образом, что такая зависимость будет существенно меньше. Метод Алмон имеет следующие достоинства: он достаточно универсален и с помощью введения небольшого количества вспомогательных переменных z_j в уравнении (4.11) ($k = 2, 3$) позволяет построить модели с распределенным лагом любой длины.

Интерпретация параметров

Из соотношения (4.3) следует, что изменение независимой переменной x в каком-либо периоде времени t влияет на значение переменной y в данном периоде и в течение p следующих периодов времени. В последующие периоды это влияние проявляться не будет. Таким образом, временной интервал влияния конечен и ограничен $(p + 1)$ -периодом. Коэффициент регрессии β_0 при переменной x_t называют *краткосрочным мультипликатором*. Он характеризует среднее абсолютное изменение y_t при изменении x_t на одну единицу своего измерения в некотором периоде времени t без учета воздействия лаговых значений фактора x .

Величины $(\beta_0 + \beta_1)$, $(\beta_0 + \beta_1 + \beta_2)$ и т. д. называются *промежуточными мультипликаторами*. Они характеризуют изменение y_t в течение 2, 3 и т. д. периодов после изменения x_t на одну единицу.

Величина $(\beta_0 + \beta_1 + \beta_2 + \dots + \beta_p)$ показывает максимальное суммарное изменение результирующей переменной y , которое будет достигнуто (по окончании текущего и p следующих периодов) под влиянием изменения фактора x на единицу в каком-либо периоде, и называется *долгосрочным мультипликатором*.

Авторегрессионные модели

Рассмотрим два важных примера авторегрессионных моделей в экономике: модель адаптивных ожиданий и модель частичной корректировки.

1. Модель адаптивных ожиданий

Роль ожидания в экономических исследованиях обусловлена тем, что ряд макроэкономических показателей (инвестиции, сбережения, спрос на активы) оказывается чувствительным к ожиданиям относительно будущего. Это в

известной мере затрудняет моделирование соответствующих экономических процессов и осуществление на их базе точных прогнозов развития экономики.

Для преодоления данной проблемы используется модель адаптивных ожиданий. В данной модели происходит постоянная корректировка ожиданий на основе получаемой информации о реализации исследуемого показателя. При этом величина корректировки должна быть пропорциональна разности между реальным и ожидаемым значениями.

Обозначим через x_{t+1}^* ожидаемое (долгосрочное) значение переменной x_t . Будем полагать, что значение величины y_t определяется этим значением:

$$y_t = \alpha + \beta x_{t+1}^* + \varepsilon_t. \quad (4.13)$$

На каждом шаге ожидания пересматриваются в некоторой пропорции от разницы между наблюдаемым значением и ожиданием переменной x на предыдущем шаге, т. е. выдвигается предположение, что эти значения связаны следующим соотношением:

$$(x_{t+1}^* - x_t^*) = \gamma (x_t - x_t^*).$$

Это выражение может быть представлено следующим образом:

$$x_{t+1}^* = \gamma x_t + (1 - \gamma)x_t^*, \quad (4.14)$$

где $0 \leq \gamma < 1$ называется коэффициентом ожидания.

Модель (4.14) называют моделью адаптивных ожиданий. Из (4.14) следует, что значение переменной, ожидаемое в следующий период времени, формируется как взвешенное среднее ее реального и ожидаемого значений в текущем периоде.

Интегрируя (4.14) и подставляя полученные результаты в (4.13), имеем

$$y_t = \alpha + \beta \gamma (x_t + (1 - \gamma)x_{t-1} + (1 - \gamma)^2 x_{t-2} + \dots) + \varepsilon_t.$$

Данная модель по форме аналогична модели Койка. Модель адаптивных ожиданий может использоваться при анализе зависимости потребления от дохода, спроса на деньги либо инвестиций от процентной ставки и в других ситуациях, где экономические показатели оказываются чувствительными в ожидании относительно будущего.

2. Модель акселератора (модель частичной корректировки)

В модели частичной корректировки (модели акселератора) предполагается, что уравнение определяет не фактическое значение зависимой переменной y_t , а желаемое (долгосрочное) значение y_t^* :

$$y_t^* = \alpha + \beta x_t + \varepsilon_t. \quad (4.15)$$

Поскольку гипотетическое значение y_t^* не является фактически существующим, то относительно него выдвигается предположение частичной корректировки:

$$y_t - y_{t-1} = \lambda(y_t^* - y_{t-1}), \quad (4.16)$$

где λ , $0 \leq \lambda \leq 1$ – коэффициент корректировки. Выражение (4.16) можно записать в виде

$$y_t = \lambda y_t^* + (1 - \lambda)y_{t-1}, \quad (4.17)$$

т. е. y_t является взвешенным средним желаемого уровня и фактического значения этой переменной в предыдущий момент времени.

Подставив (4.15) в (4.17), получим следующую модель частичной корректировки:

$$y_t = \lambda\alpha + \lambda\beta x_t + (1 - \lambda)y_{t-1} + \lambda\varepsilon_t.$$

Заметим, что данная модель по форме аналогична модели Койка. Она также включает в себя случайную объясняющую переменную y_{t-1} , но в данном случае эта переменная не коррелирует с текущим значением случайного отклонения ε_t . Поэтому МНК позволяет получить асимптотически несмещенные и эффективные оценки.

4.4. Модели стационарных временных рядов

4.4.1. Понятие стационарных и нестационарных временных рядов

Наблюдения за некоторым явлением (процессом, экономическим показателем), характер которого меняется со временем, порождает упорядоченную последовательность $\{x_t, t = \overline{1, n}\}$, называемую *временным рядом*.

На практике большинство экономических временных рядов являются нестационарными, поскольку их математическое ожидание и дисперсия зависят от времени.

Отправной точкой при разработке моделей экономических временных рядов является понятие стационарного временного ряда. Для описания класса стационарных временных рядов используется модель авторегрессии скользящего среднего.

Временной ряд $\{x_t, t = \overline{1, n}\}$ называется *строго стационарным* (или *стационарным в узком смысле*), если для любого m совместное распределение вероятностей случайных величин $x_{t_1}, \dots, x_{t_m}$ такое же, как и для $x_{t_1+\tau}, \dots, x_{t_m+\tau}$ для любых

$t_1, \dots, t_m, \tau$. В противном случае свойства строго стационарного ряда не изменяются при изменении начала отсчета времени. Из этого следует, что и все числовые характеристики временного ряда (если они существуют), в том числе математическое ожидание $E(x_t) = \mu$ и дисперсия $D(x_t) = \sigma^2$, не зависят от t .

Значение μ определяет постоянный уровень, относительно которого колеблется анализируемый временной ряд x_t , а постоянная σ^2 характеризует размах этих колебаний.

Одно из главных отличий последовательности наблюдений, образующих временной ряд, заключается в том, что члены временного ряда являются статистически взаимозависимыми. Степень тесноты статистической связи между случайными величинами x_t и $x_{t+\tau}$ может быть измерена парным коэффициентом корреляции

$$\text{Corr}(x_t, x_{t+\tau}) = \frac{\text{Cov}(x_t, x_{t+\tau})}{\sqrt{D(x_t)}\sqrt{D(x_{t+\tau})}},$$

где $\text{Cov}(x_t, x_{t+\tau}) = E((x_t - E(x_t))(x_{t+\tau} - E(x_{t+\tau})))$.

Если ряд $\{x_t\}$ стационарный, то значение $\text{Cov}(x_t, x_{t+\tau})$ не зависит от времени, а является функцией только τ , т. е.

$$\gamma(\tau) = \text{Cov}(x_t, x_{t+\tau}), \quad D(x_t) = \text{Cov}(x_t, x_t) \equiv \gamma(0).$$

Очевидно, что для стационарного временного ряда значение коэффициента корреляции также зависит только от τ :

$$\rho(\tau) = \text{Cov}(y_t, y_{t+\tau}) / D(y_t) = \gamma(\tau) / \gamma(0).$$

Временной ряд называется *стационарным в широком смысле (слабо стационарным)*, если его математическое ожидание, дисперсия и ковариация не зависят от момента времени t :

$$\begin{aligned} E(x_t) &= \mu, \\ D(x_t) &= \gamma(0), \\ \text{Cov}(x_t, x_{t+\tau}) &= \gamma(\tau) \text{ для любых } t, \tau. \end{aligned}$$

Ряд $\{x_t\}$, $t = 1, \dots, n$, называется гауссовским, если совместное распределение случайных величин $x_1, \dots, x_n$ является n -мерным нормальным распределением. Для гауссовского ряда понятия стационарности в узком и в широком смысле совпадают. В дальнейшем, говоря о стационарности некоторого ряда x_t , мы (если не оговаривается противное) будем иметь в виду, что этот

ряд стационарен в широком смысле (так, что у него существуют математическое ожидание и дисперсия).

При анализе величины $\rho(\tau)$ в зависимости от τ говорят об *автокорреляционной функции*. Автокорреляционная функция безразмерна, ее значения изменяются в пределах от -1 до 1 , кроме того, $\rho(\tau) = \rho(-\tau)$. График зависимости $\rho(\tau)$ от τ называют *коррелограммой*, которая используется для характеристики свойств механизма, порождающего временной ряд.

Процесс белого шума

Процессом белого шума («белым шумом») называют стационарный временной ряд $\{x_t\}$, для которого

$$E(x_t) \equiv 0, D(x_t) \equiv \sigma^2 > 0 \text{ и } \rho(\tau) = 0 \text{ при } \tau \neq 0.$$

Последнее означает, что при $t \neq s$ случайные величины x_t и x_s , соответствующие наблюдениям процесса белого шума в моменты t и s , некоррелированы.

В случае когда x_t имеет нормальное распределение, случайные величины $x_1, \dots, x_n$ взаимно независимы и имеют одинаковое нормальное $x_t \sim N(0, \sigma^2)$. Такой ряд называют гауссовским белым шумом.

В то же время, в общем случае, даже если некоторые случайные величины $x_1, \dots, x_n$ взаимно независимы и имеют одинаковое распределение, то это еще не означает, что они образуют процесс белого шума, т. к. случайная величина x_0 может просто не иметь математического ожидания и/или дисперсии (в качестве примера можно указать на распределение Коши).

Временной ряд, соответствующий процессу белого шума, ведет себя крайне нерегулярным образом из-за некоррелированности при $t \neq s$ случайных величин x_t и x_s , что иллюстрирует приводимый ниже график смоделированной реализации гауссовского процесса белого шума (рис. 4.1).

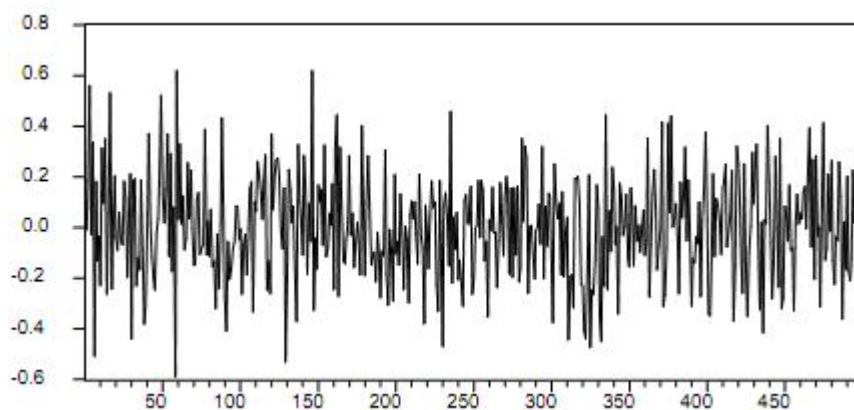


Рис. 4.1. Процесс белого шума

В связи с этим процесс белого шума не подходит для непосредственного моделирования эволюции большинства временных рядов, встречающихся в экономике. В то же время такой процесс является базой для построения более

реалистичных моделей временных рядов, порождающих «более гладкие» траектории ряда. В связи с частым использованием процесса белого шума далее, мы будем отличать этот процесс от других моделей временных рядов, используя для него обозначение ε_t .

4.4.2. Процесс авторегрессии

Одной из широко используемых моделей временных рядов является процесс авторегрессии (модель авторегрессии). В своей простейшей форме модель авторегрессии описывает механизм порождения ряда следующим образом:

$$x_t = ax_{t-1} + \varepsilon_t, \quad t = 1, \dots, n,$$

где ε_t – процесс белого шума, имеющий нулевое математическое ожидание и дисперсию σ_ε^2 ; x_0 – некоторая случайная величина; $a \neq 0$ – некоторый постоянный коэффициент.

При этом $E(x_t) = a E(x_{t-1})$, так, что рассматриваемый процесс может быть стационарным, только если $E(x_t) = 0$ для всех $t = 0, 1, \dots, n$. Отсюда

$$\begin{aligned} x_t &= ax_{t-1} + \varepsilon_t = a(ax_{t-2} + \varepsilon_{t-1}) + \varepsilon_t = a^2x_{t-2} + a\varepsilon_{t-1} + \varepsilon_t = \dots = \\ &= a^t x_0 + a^{t-1}\varepsilon_1 + a^{t-2}\varepsilon_2 + \dots + \varepsilon_t, \\ x_{t-1} &= ax_{t-2} + \varepsilon_{t-1} = a^{t-1}x_0 + a^{t-2}\varepsilon_1 + a^{t-3}\varepsilon_2 + \dots + \varepsilon_{t-1}, \\ x_{t-2} &= ax_{t-3} + \varepsilon_{t-2} = a^{t-2}x_0 + a^{t-3}\varepsilon_1 + a^{t-4}\varepsilon_2 + \dots + \varepsilon_{t-2}, \\ &\dots\dots \\ x_1 &= ax_0 + \varepsilon_1. \end{aligned}$$

Если случайная величина x_0 не коррелирована со случайными величинами $\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n$, то $\text{Cov}(x_0, \varepsilon_1) = 0$, $\text{Cov}(x_1, \varepsilon_2) = 0$, $\dots$, $\text{Cov}(x_{t-2}, \varepsilon_{t-1}) = 0$, $\text{Cov}(x_{t-1}, \varepsilon_t) = 0$ и

$$D(x_0) = D(ax_{t-1} + \varepsilon_t) = a^2 D(x_{t-1}) + D(\varepsilon_t), \quad t = 1, \dots, n.$$

Предполагая, что $D(x_0) = D(x_t) = \sigma_x^2$ для всех $t = 1, \dots, n$, находим: $\sigma_x^2 = a^2 \sigma_x^2 + \sigma_\varepsilon^2$.

Последнее может выполняться только при выполнении условия $a^2 < 1$, т. е. $|a| < 1$.

При этом получаем выражение для σ_x^2 : $\sigma_x^2 = \sigma_\varepsilon^2 / (1 - a^2)$.

Что касается автоковариаций и автокорреляций, то

$$\begin{aligned} \text{Cov}(x_t, x_{t+\tau}) &= \text{Cov}(a^t x_0 + a^{t-2}\varepsilon_2 + \dots + \varepsilon_t), \\ &a^{t+\tau}x_0 + a^{t+\tau-1}\varepsilon_1 + a^{t+\tau-2}\varepsilon_2 + \dots + \varepsilon_{t+\tau} = \\ &= a^{2t+\tau} D(x_0) + a^\tau (1 + a^2 + \dots + a^{2(t-1)}) \sigma_\varepsilon^2 = \\ &= a^\tau [a^{2t} \sigma_\varepsilon^2 / (1 - a^2) + (1 - a^{2t}) / (1 - a^2)] = [a^\tau / (1 - a^2)] \sigma_\varepsilon^2 \end{aligned}$$

и

$$\text{Corr}(x_t, x_{t+\tau}) = a^\tau,$$

т. е. при сделанных предположениях автоковариации и автокорреляции зависят только от того, насколько разнесены по времени соответствующие наблюдения.

Таким образом, механизм порождения последовательных наблюдений, заданный соотношениями

$$x_t = ax_{t-1} + \varepsilon_t, \quad t = 1, \dots, n,$$

порождает стационарный временной ряд, если

- $|a| < 1$;
- $E(x_0) = 0$;
- $D(x_0) = \frac{\sigma_\varepsilon^2}{(1-a^2)}$;
- случайная величина x_0 не коррелирована со случайными ошибками

$\varepsilon_1, \dots, \varepsilon_n$.

При этом $\rho(\tau) = a^\tau$.

Рассмотренная модель порождает (при указанных условиях) стационарный ряд, имеющий нулевое математическое ожидание. Однако ее можно легко распространить и на временные ряды y_t с ненулевым математическим ожиданием $E(y_t) = \mu$, полагая, что указанная модель относится к центрированному ряду $y_t = y_t - \mu$:

$$x_t - \mu = a(x_{t-1} - \mu) + \varepsilon_t, \quad t = 1, \dots, n,$$

поэтому

$$x_t = ax_{t-1} + \delta + \varepsilon_t, \quad t = 1, \dots, n,$$

где $\delta = \mu(1-a)$.

Следовательно, без ограничения общности можно обойтись в текущем рассмотрении моделей авторегрессии, порождающей стационарный процесс с нулевым средним.

Продолжая рассмотрение ранее определенного процесса x_t (с нулевым математическим ожиданием), заметим, что для него

$$\gamma(1) = E(X_t X_{t-1}) = E[(aX_{t-1} + \varepsilon_t)X_{t-1}] = a\gamma(0),$$

при этом

$$\rho(1) = \gamma(1)/\gamma(0) = a,$$

и при значениях $a > 0$, близких к 1, между соседними наблюдениями имеется сильная положительная корреляция, что обеспечивает более гладкий характер поведения траекторий ряда по сравнению с процессом белого шума. При $a < 0$ процесс авторегрессии, напротив, имеет менее гладкие реализации, поскольку в этом случае проявляется тенденция чередования знаков последовательных наблюдений.

Следующие два графика демонстрируют поведение смоделированных реализаций временных рядов, порожденных моделями авторегрессии $x_t = ax_{t-1} + \varepsilon_t$ с $\sigma_\varepsilon^2 = 0,2$ при $a = 0,8$ (первый график) и $a = -0,8$ (второй график).

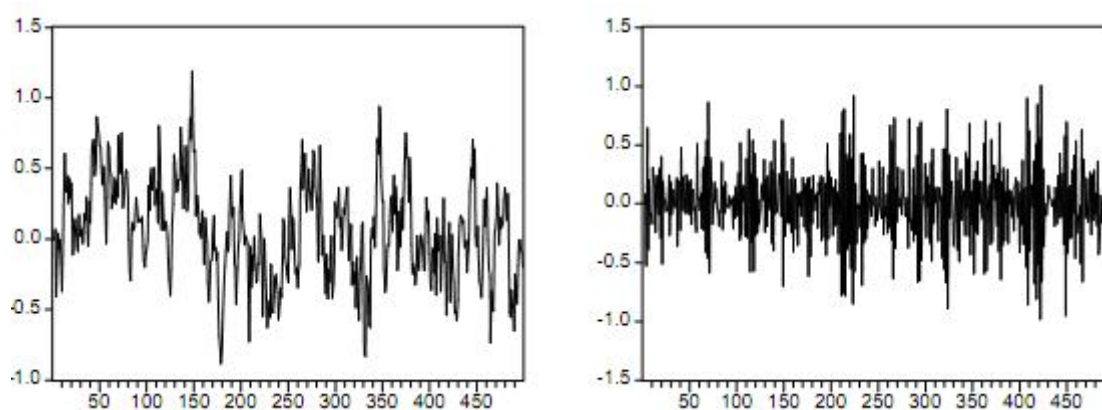


Рис. 4.2. Процесс авторегрессии

Рассмотренную модель $x_t = ax_{t-1} + \varepsilon_t$ (рис. 4.2) называют процессом авторегрессии первого порядка. *Процесс авторегрессии порядка p* (в кратком обозначении – **AR(p)**) определяется соотношениями

$$x_t = a_1x_{t-1} + a_2x_{t-2} + \dots + a_px_{t-p} + \varepsilon_t, \quad a_p \neq 0,$$

где ε_t – процесс белого шума с $D(\varepsilon_t) = \sigma_\varepsilon^2$.

Будем полагать, что $\text{Cov}(x_{t-s}, \varepsilon_t) = 0$ для всех $s > 0$; при этом говорят, что случайные величины ε_t образуют инновационную (обновляющую) последовательность, а случайная величина ε_t называется инновацией для наблюдения в момент t . Такая терминология объясняется тем, что наблюдаемое значение ряда в момент t получается здесь как линейная комбинация p предшествующих значений этого ряда плюс не коррелированная с этими предшествующими значениями случайная составляющая ε_t , отражающая обновленную информацию, скажем о состоянии экономики на момент t , влияющую на наблюдаемое значение x_t .

При рассмотрении процессов авторегрессии и некоторых других моделей удобно использовать *оператор запаздывания* L (*lag operator*), который воздействует на временной ряд и определяется соотношением

$$Lx_t = x_{t-1};$$

иногда его называют оператором обратного сдвига и используют для него обозначение B (backshift operator).

Если оператор запаздывания применяется k раз, что обозначается как L^k , то это дает в результате $L^k x_t = x_{t-k}$.

Выражение $a_1 x_{t-1} + a_2 x_{t-2} + \dots + a_p x_{t-p}$ можно записать теперь в виде $(a_1 L + a_2 L^2 + \dots + a_p L^p) x_t$, а соотношение, определяющее процесс авторегрессии p -го порядка, в виде $a(L)x_t = \varepsilon_t$, где $a(L) = 1 - (a_1 L + a_2 L^2 + \dots + a_p L^p)$.

Для того чтобы такой процесс был стационарным, все корни алгебраического уравнения $a(z) = 0$ (вещественные и комплексные) должны лежать вне единичного круга $|z| > 1$. (В частности, для процесса **AR(1)** имеем $a(z) = 1 - az$, уравнение $a(z) = 0$ имеет корень $z = 1/a$, и условие стационарности $z > 1$ равносильно уже знакомому нам условию $a < 1$.) При этом решение уравнения $a(L)x_t = \varepsilon_t$ можно представить в виде

$$x_t = \frac{1}{a(L)} \varepsilon_t = \sum_{j=0}^{\infty} b_j \varepsilon_{t-j}, \text{ где } \sum_{j=0}^{\infty} |b_j| < \infty,$$

откуда следует, что

$$E(x_t) = E\left(\sum_{j=0}^{\infty} b_j \varepsilon_{t-j}\right) = \sum_{j=0}^{\infty} b_j E(\varepsilon_{t-j}) = 0.$$

Стационарный процесс **AR(p)** с ненулевым математическим ожиданием μ удовлетворяет соотношению $a(L)(x_t - \mu) = \varepsilon_t$, или $a(L)x_t = \delta + \varepsilon_t$, где $\delta = a(L)\mu = \mu(1 - a_1 - a_2 - \dots - a_p) = \mu a(1)$.

При этом решение уравнения $a(L)(x_t - \mu) = \varepsilon_t$ имеет вид

$$x_t = \mu + \frac{1}{a(L)} \varepsilon_t.$$

Таким образом, если стационарный процесс **AR(p)** задан в виде $a(L)x_t = \delta + \varepsilon_t$, то следует помнить о том, что в этом случае математическое ожидание этого процесса равно не δ , а

$$\mu = \frac{\delta}{(1 - a_1 - a_2 - \dots - a_p)}.$$

Для процесса **AR(1)** имеем $a(L) = 1 - aL$, так что (вне зависимости от того, равно μ нулю или нет)

$$x_t - \mu = (1/(1 - aL))\varepsilon_t = (1 + a_1L + a_2L^2 + \dots)\varepsilon_t = \varepsilon_t + a\varepsilon_{t-1} + a^2\varepsilon_{t-2} + \dots.$$

Из последнего выражения видно, что

$$\rho(k) = \text{Corr}(x_t, x_{t+k}) = a^k, \quad k = 0, 1, 2, \dots$$

При $0 < a < 1$ коррелограмма (график функции $\rho(k)$ для $k = 0, 1, 2, \dots$) отражает показательное убывание корреляций с возрастанием интервала между наблюдениями; при $-1 < a < 0$ коррелограмма имеет характер затухающей косинусоиды. Сравним поведение коррелограмм стационарного процесса **AR(1)** при $a = \pm 0,8$ (рис. 4.3):

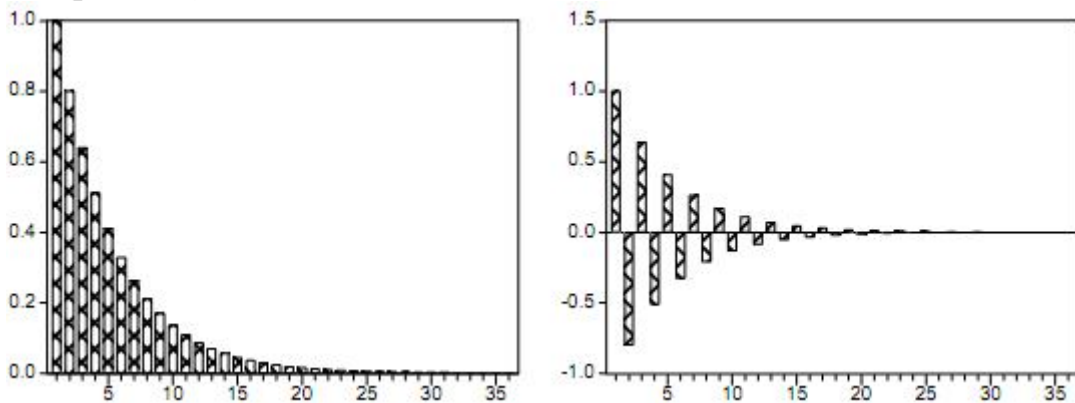


Рис. 4.3. Коррелограмма стационарного процесса AR (1)

4.4.3. Процесс скользящего среднего

Еще одной простой моделью порождения временного ряда является **процесс скользящего среднего порядка q (MA(q))**. Согласно этой модели,

$$x_t = \varepsilon_t + b_1\varepsilon_{t-1} + b_2\varepsilon_{t-2} + \dots + b_q\varepsilon_{t-q}, \quad b_q \neq 0,$$

где ε_t – процесс белого шума.

Такой процесс имеет нулевое математическое ожидание. Модель можно обобщить до процесса, имеющего ненулевое математическое ожидание μ , полагая

$$x_t - \mu = \varepsilon_t + b_1\varepsilon_{t-1} + b_2\varepsilon_{t-2} + \dots + b_q\varepsilon_{t-q}$$

или

$$x_t = \mu + \varepsilon_t + b_1 \varepsilon_{t-1} + b_2 \varepsilon_{t-2} + \dots + b_q \varepsilon_{t-q}.$$

Для процесса скользящего среднего порядка q используется обозначение **МА(q)** (скользящее среднее – moving average).

При $q = 0$ и $\mu = 0$ получаем процесс белого шума.

Если $q = 1$, то

$x_t - \mu = \varepsilon_t + b\varepsilon_{t-1}$ – скользящее среднее первого порядка. В последнем случае $D(x_t) = (1 + b^2)\sigma_\varepsilon^2$, $E((x_t - \mu)(x_{t-1} - \mu)) = b\sigma_\varepsilon^2$, $E((x_t - \mu)(x_{t-k} - \mu)) = 0, k > 1$.

Так что процесс x_t является стационарным с $E(x_t) = 0$, $D(x_t) = (1 + b^2)\sigma_\varepsilon^2$,

$$\gamma(k) = \begin{cases} (1 + b^2)\sigma_\varepsilon^2, k = 0, \\ b\sigma_\varepsilon^2, k = 1, \\ 0, k > 1. \end{cases}$$

Автокорреляции этого процесса равны

$$\gamma(k) = \begin{cases} 1, k = 0, \\ b/(1 + b^2), k = 1, \\ 0, k > 1. \end{cases}$$

Коррелированными оказываются только соседние наблюдения. Корреляция между ними положительна, если $b > 0$, и отрицательна при $b < 0$. Соответственно, процесс **МА(1)** с $b > 0$ имеет более гладкие по сравнению с белым шумом реализации, а процесс **МА(1)** с $b < 0$ имеет менее гладкие по сравнению с белым шумом реализации (рис. 4.4). Заметим, что для любого процесса **МА(1)** $\rho(1) \leq 0,5$, т. е. корреляционная связь между соседними наблюдениями невелика, тогда как у процесса **АР(1)** такая связь может быть сколь угодно сильной (при значениях a , близких к 1).

Модель **МА(q)** кратко можно записать в виде

$$x_t - \mu = b(L) \varepsilon_t,$$

где $b(L) = 1 + b_1 L + b_2 L^2 + \dots + b_q L^q$. Для нее

$$\gamma(k) = E[(x_t - \mu)(x_{t-k} - \mu)] = \begin{cases} \left(\sum_{j=0}^{q-k} b_j b_{j+k} \right) \sigma_\varepsilon^2, 0 \leq k \leq q, \\ 0, k > q, \end{cases}$$

так что $\mathbf{MA}(q)$ является стационарным процессом с нулевым математическим ожиданием, дисперсией $\sigma_x^2 = (1 + b_1^2 + \dots + b_q^2)\sigma_\varepsilon^2$ и автокорреляциями

$$\rho_k = \begin{cases} \left(\sum_{j=0}^{q-k} b_j b_{j+k} \right) / \sum_{j=0}^q b_j^2, & k = 0, 1, \dots, q, \\ 0, & k = q+1, q+2, \dots \end{cases}$$

Здесь статистическая связь между наблюдениями сохраняется в течение q единиц времени (т. е. «длительность памяти» процесса равна q). Подобного рода временные ряды соответствуют ситуации, когда некоторый экономический показатель находится в равновесии, но отклоняется от положения равновесия в силу последовательно возникающих непредсказуемых событий. Причем система такова, что влияние таких событий отмечается на протяжении некоторого временного периода. Если влияние прошлых событий ослабевает с течением времени так, что $b_j = a^j$, $0 < a < 1$, то искусственное предположение о том, что ряд ε_t начинается в «бесконечном прошлом», приводит к модели бесконечного скользящего среднего $\mathbf{MA}(\infty)$:

$$X_t = \sum_{j=0}^{\infty} a^j \varepsilon_{t-j} = \sum_{j=0}^{\infty} b_j \varepsilon_{t-j}, \quad \text{где} \quad \sum_{j=0}^{\infty} |b_j| < \infty.$$

Такое же представление допускает стационарный процесс авторегрессии первого порядка $\mathbf{AR}(1)$ $x_t = ax_{t-1} + \varepsilon_t$, $|a| < 1$, т. е. в рассматриваемом случае процесс $\mathbf{MA}(\infty)$ эквивалентен процессу $\mathbf{AR}(1)$. Вообще, всякий стационарный процесс $\mathbf{AR}(p)$ можно записать в форме процесса $\mathbf{MA}(\infty)$:

$$x_t = \mu + \frac{1}{a(L)} \varepsilon_t = \mu + \sum_{j=0}^{\infty} b_j \varepsilon_{t-j} = \mu + b(L) \varepsilon_t, \\ b(L) = \sum_{j=0}^{\infty} b_j L^j = \frac{1}{a(L)}, \quad \text{и} \quad \sum_{j=0}^{\infty} |b_j| < \infty.$$

Примеры:

а) рассмотрим процесс $\mathbf{MA}(1)$ с $b = 0,8$ и $E(x_t) = 6$, т. е.

$$x_t = 6 + \varepsilon_t + 0,8 \varepsilon_{t-1}.$$

Для него $\rho(1) = 0,8/(1 + 0,82) = 0,488$;

б) для процесса $\mathbf{MA}(1)$ с $b = -0,8$ и $E(x_t) = 6$ имеем

$$\rho(1) = -0,8/(1 + 0,82) = -0,488.$$

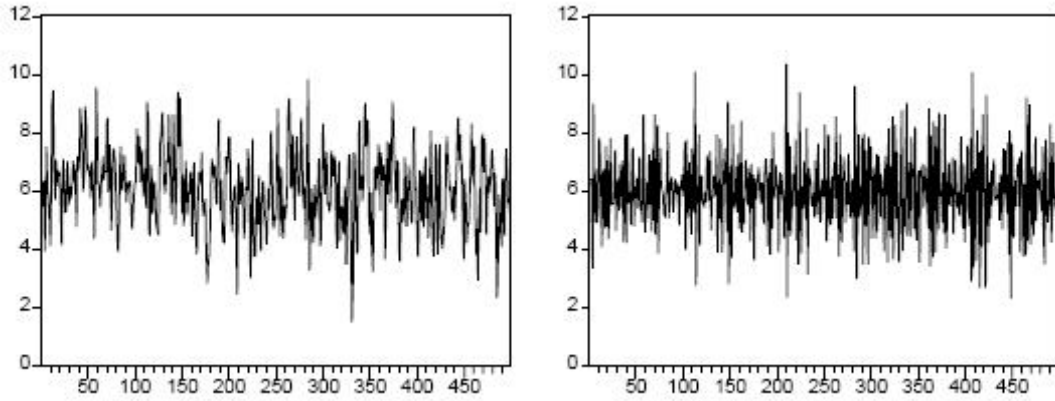


Рис. 4.4. Реализации процесса $\text{MA}(1)$

4.4.4. Смешанный процесс авторегрессии – скользящего среднего (процесс авторегрессии с остатками в виде скользящего среднего)

Процесс x_t с нулевым математическим ожиданием, принадлежащий такому классу процессов, характеризуется порядками p и q его AR и MA составляющих и обозначается как процесс $\text{ARMA}(p, q)$ (autoregressive moving average, mixed autoregressive moving average). Более точно, процесс x_t с нулевым математическим ожиданием принадлежит классу $\text{ARMA}(p, q)$, если

$$x_t = \sum_{j=1}^p a_j x_{t-j} + \sum_{j=1}^q b_j \varepsilon_{t-j}, a_p \neq 0, b_q \neq 0,$$

где ε_t – процесс белого шума и $b_0 = 1$.

В операторной форме последнее соотношение имеет вид $a(L) x_t = b(L) \varepsilon_t$, где $a(L)$ и $b(L)$ имеют тот же вид, что и в определенных ранее моделях $\text{AR}(p)$ и $\text{MA}(q)$. Если процесс имеет постоянное математическое ожидание μ , то он является процессом типа $\text{ARMA}(p, q)$, если

$$x_t - \mu = \sum_{j=1}^p a_j (x_{t-j} - \mu) + \sum_{j=1}^q b_j \varepsilon_{t-j}.$$

Отметим следующие свойства процесса $\text{ARMA}(p, q)$ с $E(x_t) = \mu$:

- процесс стационарен, если все корни уравнения $a(z) = 0$ лежат вне единичного круга $|z| > 1$;
- если процесс стационарен, то существует эквивалентный ему процесс $\text{MA}(\infty)$:

$$x_t - \mu = \sum_{j=0}^{\infty} c_j \varepsilon_{t-j}, c_0 = 1, \sum_{j=0}^{\infty} |c_j| < \infty$$

или

$$x_t - \mu = c(L)\varepsilon_t,$$

где $c(z) = \sum_{j=0}^{\infty} c_j z^j = \frac{b(z)}{a(z)}$;

• если все корни уравнения $b(z) = 0$ лежат вне единичного круга $|z| > 1$ (условие обратимости), то существует эквивалентное представление процесса x_t в виде процесса авторегрессии бесконечного порядка **AR**(∞)

$$x_t - \mu = \sum_{j=1}^{\infty} d_j (x_{t-j} - \mu) + \varepsilon_t,$$

или

$$d(L)(x_t - \mu) = \varepsilon_t,$$

где $d(z) = 1 - \sum_{j=1}^{\infty} d_j z^j = \frac{a(z)}{b(z)}$.

Отсюда вытекает, что стационарный процесс **ARMA**(p, q) всегда можно аппроксимировать процессом скользящего среднего достаточно высокого порядка, а при выполнении условия обратимости его можно также аппроксимировать процессом авторегрессии достаточно высокого порядка.

Модели **ARMA**, учитывающие наличие сезонности

Если наблюдаемый временной ряд обладает выраженной сезонностью, то модель **ARMA**, соответствующая этому ряду, должна содержать составляющие, обеспечивающие проявление сезонности в порождаемой этой моделью последовательности наблюдений.

Для квартальных данных чисто сезонными являются стационарные модели сезонной авторегрессии первого порядка (**SAR**(1)):

$$x_t = a_4 x_{t-4} + \varepsilon_t, \quad |a_4| < 1$$

и сезонного скользящего среднего первого порядка (**SMA**(1)):

$$x_t = \varepsilon_t + b_4 \varepsilon_{t-4}.$$

В первой модели $\rho(k) = a_4^{k/4}$ для $k = 4m, m = 0, 1, 2, \dots$, $\rho(k) = 0$ – для остальных $k > 0$.

Во второй модели $\rho(0) = 1, \rho(4) = b_4, \rho(k) = 0$ – для остальных $k > 0$.

Ниже приведены смоделированные реализации модели **SAR**(1) с $a_4 = 0,8$ и модели **SMA**(1) с $b_4 = 0,8$ (рис. 4.5).

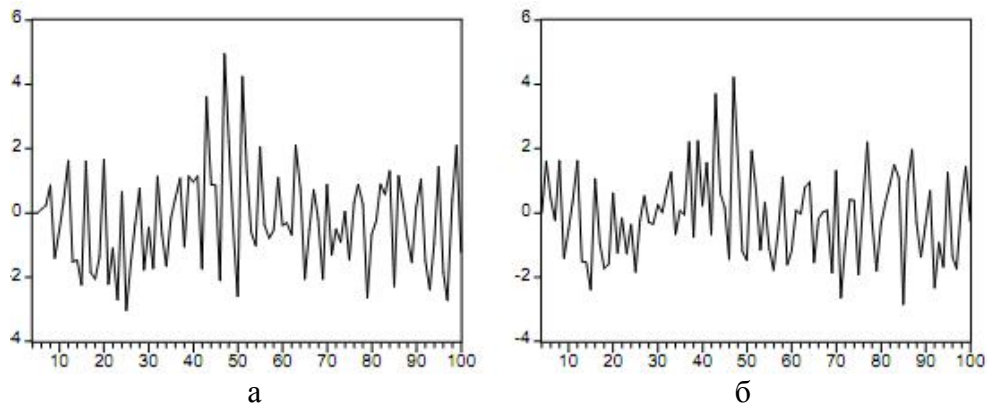


Рис. 4.5. Смоделированные реализации сезонных моделей **ARMA**:
а – **SAR(1)**; б – **SMA(1)**

Комбинации несезонных и сезонных изменений реализуются, например, в моделях **ARMA**((1, 4), 1)

$$x_t = a_1 x_{t-1} + a_4 x_{t-4} + \varepsilon_t + b_1 \varepsilon_{t-1}$$

и **ARMA**(1, (1, 4))

$$x_t = a_1 x_{t-1} + \varepsilon_t + b_1 \varepsilon_{t-1} + b_4 \varepsilon_{t-4}.$$

Следующие два графика (рис. 4.6) показывают поведение смоделированных реализаций таких рядов при $a_1 = 2/3$, $a_4 = -1/48$, $b_4 = 1/5$ у первого ряда и при $a_1 = 0,4$, $b_1 = 0,3$, $b_4 = 0,8$ у второго ряда.

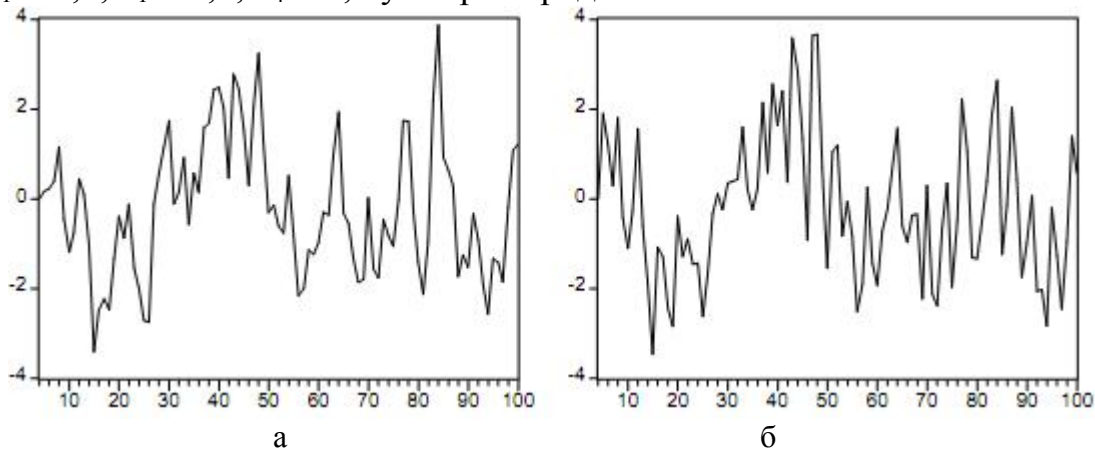


Рис. 4.6. Реализации моделей **ARMA**:
а – **ARMA**((1, 4), 1); б – **ARMA**(1, (1, 4))

Кроме рассмотренных примеров аддитивных сезонных моделей, употребляются также и мультипликативные спецификации, например,

$$(1 - a_1 L)x_t = (1 + b_1 L)(1 + b_4 L^4) \varepsilon_t,$$

$$(1 - a_1 L)(1 - a_4 L^4) x_t = (1 + b_1 L) \varepsilon_t.$$

Первая дает

$$x_t = a_1 x_{t-1} + \varepsilon_t + b_1 \varepsilon_{t-1} + b_4 \varepsilon_{t-4} + b_1 b_4 \varepsilon_{t-5},$$

а вторая –

$$x_t = a_1 x_{t-1} + a_4 x_{t-4} - a_1 a_4 x_{t-5} + \varepsilon_t + b_1 \varepsilon_{t-1}.$$

В первой модели допускается взаимодействие составляющих скользящего среднего на лагах 1 и 4 (т. е. значений ε_{t-1} и ε_{t-4}), а во второй – взаимодействие авторегрессионных составляющих на лагах 1 и 4 (т. е. значений x_{t-1} и x_{t-4}).

4.4.5 Методы построения и тестирования модели ARMA

Методология построения моделей типа **ARMA(p, q)** известна как методология Бокса – Дженкинса и предусматривает выполнение следующих трех этапов [11]: *идентификация модели, оценивание параметров, тестирование адекватности.*

На этапе идентификации производится выбор некоторой частной модели из всего класса **ARMA**, т. е. выбор значений p и q . Используемые при этом процедуры являются не вполне точными, что может при последующем анализе привести к выводу о непригодности идентифицированной модели и необходимости замены ее альтернативной моделью. На этом же этапе делаются предварительные грубые оценки коэффициентов $a_1, a_2, \dots, a_p, b_1, b_2, \dots, b_q$ идентифицированной модели.

Идентификация модели предполагает выполнение следующих условий:

1) визуальный анализ графика временного ряда с целью выявления «выбросов», «пропусков», структурных изменений, а также признаков нестационарности типа зависимости среднего значения и дисперсии временного ряда от времени, указывающих на наличие временных трендов и гетероскедастичности;

2) анализ автокорреляционных (ACF) и частных автокорреляционных (PAC) функций, позволяющий подтвердить либо опровергнуть предположение о стационарности анализируемого временного ряда, а также указать на возможные значения параметров p и q .

Оценивание параметров. Для статистического оценивания параметров модели **ARMA(p, q)** с заданными значениями p и q могут использоваться различные методы: линейный и нелинейный метод наименьших квадратов, полный и условный метод максимального правдоподобия (ММП), а также метод моментов и обобщенный метод моментов.

Тестирование адекватности основано на анализе тестовых статистик и статистической проверке гипотез относительно параметров модели. Адекватная модель должна обладать следующими свойствами.

1. Оценки параметров модели должны быть статистически значимыми.

2. Остатки построенной модели должны быть «белым шумом», т. е. быть некоррелированными. При этом сумма квадратов остатков может служить одним из критериев выбора модели.

Проверка гипотезы о том, что наблюдаемые данные являются реализацией процесса белого шума, может осуществляться с помощью:

- визуального анализа графиков остатков, а также ACF и PACF;
- асимптотического *теста значимости значений* ACF, основанного на нормальном приближении тестовой статистики: значения ACF считаются статистически значимыми на уровне значимости $\alpha = 0,05$, если выходят за границы соответствующего доверительного интервала;
- теста для проверки гипотезы об отсутствии автокорреляции значений временного ряда $\{x_t\}$ на заданном лаговом диапазоне, включающем $k > 1$ лагов, с помощью *Q-статистики Льюнга – Бокса*, рассчитываемой на основании значений ACF по формуле

$$Q = n(n+2) \sum_{i=1}^k \hat{\rho}_i^2(n-i).$$

Распределение *Q*-статистики (при условии, что верна нулевая гипотеза об отсутствии автокорреляции значений временного ряда, на заданном лаговом диапазоне) асимптотически при $n \rightarrow \infty$ приближается к распределению χ^2 с $\nu = k - (p + q)$ степенями свободы, где p и q – число **AR**- и **MA**-компонентов в построенной модели временного ряда.

Если $Q \geq \Delta(\alpha)$ (где $\Delta(\alpha)$ – критическое значение статистики, равное квантили уровня $(1 - \alpha)$ распределения χ^2 с ν степенями свободы), то нулевая гипотеза об отсутствии автокорреляции отклоняется.

3. *Распределение остатков должно быть нормальным*, т. е. остатки должны быть гауссовым «белым шумом». Отметим, что данное свойство важно при тестировании моделей по коротким временным рядам. В случае достаточно длинных временных рядов данное свойство не требуется. Для проверки гипотезы о нормальном распределении остатков могут использоваться различные тесты, например, критерий согласия Пирсона, тест Жака-Бера (*Jarque-Bera's – JB test*).

JB-тест основан на проверке статистической значимости расхождения фактических значений коэффициентов асимметрии и эксцесса, а также ожидаемых для нормального распределения нулевых значений данных характеристик.

Тестовая статистика вычисляется по формуле

$$JB = \frac{n - (p + q)}{6} (S^2 + K^2 / 4),$$

где S и K – соответственно коэффициенты асимметрии и эксцесса для ряда остатков; p и q – число AR- и MA-компонентов в построенной модели ВР.

Если гипотеза H_0 верна, то статистика JB имеет распределение χ^2 с двумя степенями свободы.

4. Модель «должна быть наиболее простой из возможных альтернативных моделей». Это требование основано на «принципе экономности» (*principle of parsimony*): из двух моделей, признанных по результатам тестирования на одном и том же наборе данных адекватными, лучшей считается модель с меньшим числом параметров, т. е. с меньшими значениями p и q . Для выбора наиболее «экономичной» модели могут использоваться AIC-статистика Акаике (*Akaike information criterion*) и SC-статистика Шварца (*Schwartz criterion*), определяемые по формулам [11]:

$$AIC = \ln\left(\frac{ESS}{n}\right) + \frac{2m}{n},$$
$$SC = \ln\left(\frac{ESS}{n}\right) + \frac{m}{n} \ln(n),$$

где ESS сумма квадратов остатков; m – **число оцениваемых** параметров, т. е. $m = p + q + \delta$ ($\delta = 1$, если используется модель со свободным членом; $\delta = 0$ в противном случае).

В соответствии с данным критерием следует выбирать модели с меньшими значениями статистик AIC и SC.

4.5. Модели нестационарных временных рядов

Признаком нестационарного стохастического процесса является нарушение одного из условий стационарности. Конкретная реализация нестационарного стохастического процесса представляет собой нестационарный временной ряд. Признаками нестационарности временного ряда могут служить наличие тенденции, систематических изменений дисперсии, периодической составляющей, систематически изменяющихся взаимозависимостей между элементами временного ряда.

Как правило, значения, характеризующие изменение экономических показателей во времени, образуют нестационарные временные ряды.

Различают два типа нестационарных временных рядов: временные ряды, нестационарные по среднему значению, и временные ряды, нестационарные по дисперсии [12].

Временной ряд $\{x_t\}$ является нестационарным по среднему значению, если его математическое ожидание зависит от времени, и нестационарным по дисперсии, если его дисперсия зависит от времени:

$$E(x_t) = \mu, D(x_t) = \sigma_t^2, t \geq 1.$$

Для описания временных рядов, нестационарных по среднему значению, используются два основных класса моделей: модели с детерминированным

трендом и модели со стохастическим трендом. Среди последних различают интегрированные временные ряды, описываемые моделью **ARIMA**, «процессы единичного корня».

Рассмотрим авторегрессионный процесс первого порядка, определяемый моделью $x_t = a_1 x_{t-1} + \varepsilon_t$.

Такой процесс является стационарным при выполнении условия $-1 < a_1 < 1$. Рассмотрим некоторые реализации этого ряда при нарушении данного условия. На рис. 4.7 приведены смоделированные реализации такого ряда при $a_1 = 0,5; 0,7; 0,9; 1$.

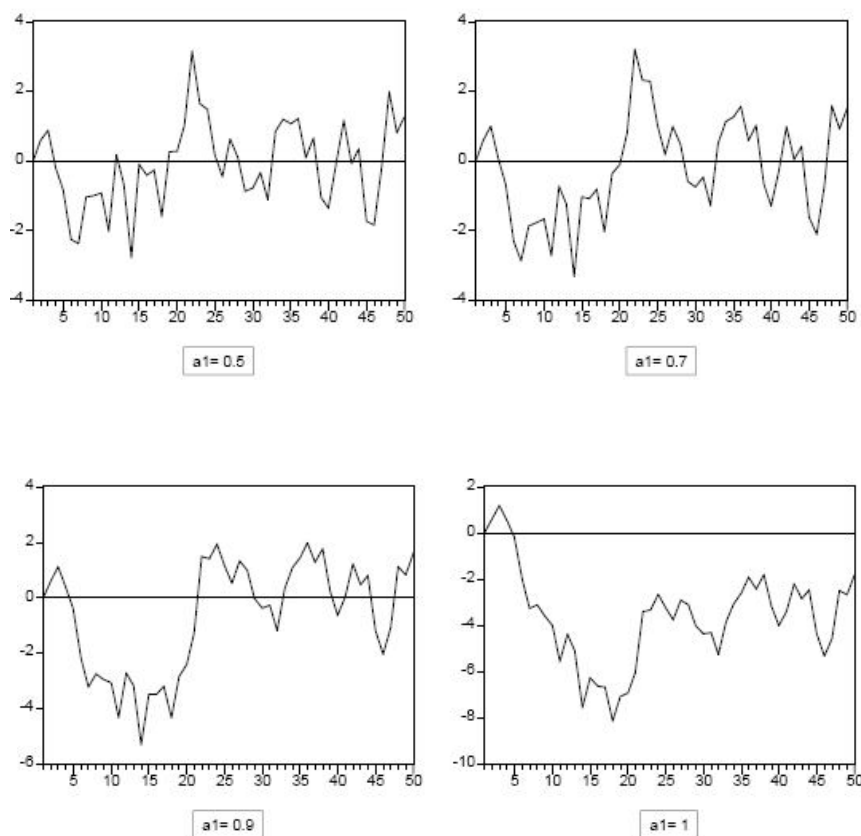


Рис. 4.7. Авторегрессионный процесс первого порядка

Во всех случаях в качестве начального значения x_0 взято $x_0 = 0$ и использовалась одна и та же последовательность значений $\varepsilon_1, \dots, \varepsilon_n$, имитирующая гауссовский «белый шум» с дисперсией, равной единице.

Однако поведение смоделированных рядов оказалось качественно различным.

При возрастании значения a_1 от $a_1 = 0$ (белый шум) до $a_1 = 1$ количество пересечений нулевого уровня уменьшается, все более длинными становятся периоды, в течение которых значения ряда находятся по одну сторону от нулевого уровня. График ряда при $a_1 = 1$ иллюстрирует характерное свойство реализаций процесса

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

состоящее в том, что такой процесс, начавшись в момент $t = 1$ со значения x_1 (в данном случае $x_1 = 0$), в дальнейшем очень редко пересекает уровень x_1 («возвращается к уровню x_1 ») и, находясь в течение длительного времени по одну сторону от этого уровня (выше или ниже), может удаляться от этого уровня на значительные расстояния.

Такой процесс называется *случайным блужданием (процесс случайного блуждания (random walk))*.

Рассмотрим процесс случайного блуждания подробнее:

$$x_t = x_{t-1} + \varepsilon_t, \quad t = 1, \dots, n,$$

стартовое значение x_0 .

Можно представить x_t в виде

$$\begin{aligned} x_t &= x_{t-1} + \varepsilon_t = (x_{t-2} + \varepsilon_{t-1}) + \varepsilon_t = x_{t-2} + \varepsilon_{t-1} + \varepsilon_t = \\ &= (x_{t-3} + \varepsilon_{t-2}) + \varepsilon_{t-1} + \varepsilon_t = x_{t-3} + \varepsilon_{t-2} + \varepsilon_{t-1} + \varepsilon_t = \dots = \\ &= x_0 + (\varepsilon_1 + \dots + \varepsilon_t), \end{aligned}$$

$$x_t = x_0 + \sum_{j=1}^t \varepsilon_j.$$

Отсюда сразу получаем:

$$E(x_t | x_0) = x_0,$$

$$D(x_t | x_0) = D(\varepsilon_1 + \dots + \varepsilon_t) = D(\varepsilon_1) + \dots + D(\varepsilon_t) = tD(\varepsilon_1) = t\sigma_\varepsilon^2.$$

Далее

$$\begin{aligned} \text{Cov}(x_t, x_{t-1} | x_0) &= E((x_t - x_0)(x_{t-1} - x_0) | x_0) = \\ &= E((\varepsilon_1 + \dots + \varepsilon_t)(\varepsilon_1 + \dots + \varepsilon_{t-1})) = (t-1)\sigma_\varepsilon^2 \end{aligned}$$

и не зависит от значения x_0 , следовательно,

$$\begin{aligned} \text{Corr}(x_t, x_{t-1}) &= \frac{(t-1)\sigma_\varepsilon^2}{\sqrt{D(x_t)}\sqrt{D(x_{t-1})}} = \frac{(t-1)\sigma_\varepsilon^2}{\sqrt{\sigma_\varepsilon^2 t}\sqrt{\sigma_\varepsilon^2 (t-1)}} = \\ &= \frac{\sqrt{t-1}}{\sqrt{t}} = \sqrt{1 - \frac{1}{t}}. \end{aligned}$$

При $x_0 = 0$ получаем $x_t = \sum_{j=1}^t \varepsilon_j, \quad t = 1, \dots, n$.

Рассматривая последний ряд, имеем:

$$E(x_t) = 0, D(x_t) = t\sigma_\varepsilon^2,$$

Таким образом, этот ряд – нестационарный.

Этот ряд является моделью **стохастического тренда**, который обнаруживается во многих экономических временных рядах и должен обязательно приниматься во внимание при построении регрессионных моделей связи между двумя или несколькими рядами, имеющими стохастический тренд.

Поясним различие между временными рядами, имеющими только детерминированный тренд, и рядами, которые (возможно, наряду с детерминированным) имеют стохастический тренд.

Для этого рассмотрим следующие две простые модели нестационарных рядов. Пусть в первой

$$x_t = \alpha + \beta t + \varepsilon_t, t = 1, \dots, n,$$

т. е. на детерминированный линейный тренд накладываются случайные ошибки в виде белого шума. Вторая модель пусть представляет **случайное блуждание со сносом**, т. е. процесс

$$x_t = \alpha + x_{t-1} + \varepsilon_t, t = 1, \dots, n, \text{ задано } x_0,$$

приращения которого имеют ненулевое математическое ожидание $E(\Delta x_t) = \alpha \neq 0$.

Процесс x_t во второй модели можно представить в виде

$$\begin{aligned} x_t &= \alpha + x_{t-1} + \varepsilon_t = \alpha + (\alpha + x_{t-2} + \varepsilon_{t-1}) + \varepsilon_t = 2\alpha + x_{t-2} + \varepsilon_{t-1} + \varepsilon_t = \\ &= 3\alpha + x_{t-3} + \varepsilon_{t-2} + \varepsilon_{t-1} + \varepsilon_t = \dots = x_0 + \alpha t + (\varepsilon_1 + \dots + \varepsilon_t), \\ x_t &= x_0 + \alpha t + \sum_{j=1}^t \varepsilon_j. \end{aligned}$$

Таким образом, ряд x_t имеет и детерминированный, и стохастический тренды. Детрендрование первого ряда приводит к ряду

$$x_t^0 = x_t - (\alpha + \beta t) = \varepsilon_t,$$

который является стационарным. Детрендрование второго ряда приводит к ряду

$$x_t^0 = x_t - (x_0 + \alpha t) = \sum_{j=1}^t \varepsilon_j,$$

который не является стационарным. Привести его к стационарному можно, если перейти от ряда уровней x_t к ряду разностей $\Delta x_t = x_t - x_{t-1}$.

Такой переход в теории временных рядов называют **дифференцированием**.

При таком переходе получаем для первого ряда

$$\Delta x_t = x_t - x_{t-1} = (\alpha + \beta t + \varepsilon_t) - (\alpha + \beta(t-1) + \varepsilon_{t-1}) = \beta + \varepsilon_t - \varepsilon_{t-1},$$

а для второго ряда $\Delta x_t = x_t - x_{t-1} = \alpha + \varepsilon_t$.

Оба продифференцированных ряда Δx_t стационарны. Первый продифференцированный ряд относится к классу **МА(1)** и имеет математическое ожидание β . Второй продифференцированный ряд относится к классу **МА(0)** и имеет математическое ожидание α . Таким образом, в отличие от детрендривания операция дифференцирования приводит к стационарному ряду в обоих случаях.

Временной ряд x_t называется *стационарным относительно детерминированного тренда $f(t)$* , если ряд $x_t - f(t)$ – стационарный. Если ряд x_t стационарен относительно некоторого детерминированного тренда, то говорят, что этот ряд принадлежит *классу рядов, стационарных относительно детерминированного тренда*, или что он является *TS-рядом* (*TS – time stationary*).

В класс TS-рядов включаются также стационарные ряды, не имеющие детерминированного тренда. Временной ряд x_t называется *интегрированным порядка k* , $k = 1, 2, \dots$, если:

- ряд x_t не является стационарным или стационарным относительно детерминированного тренда, т. е. не является *TS-рядом*;
- ряд $\Delta^k x_t$, полученный в результате k -кратного дифференцирования ряда x_t , является стационарным рядом;
- ряд $\Delta^{k-1} x_t$, полученный в результате $(k - 1)$ -кратного дифференцирования ряда x_t , не является *TS-рядом*.

Для интегрированного ряда порядка k используют обозначение **I(k)**. Если ряд X_t является интегрированным порядка k , то для краткости это обозначается как $x_t \sim \mathbf{I}(k)$. В этой системе обозначений соотношение $x_t \sim \mathbf{I}(0)$ соответствует ряду, который является стационарным и при этом не является результатом дифференцирования *TS-ряда*.

Совокупность интегрированных рядов различных порядков $k = 1, 2, \dots$ образует класс *разностно стационарных*, или **DS-рядов** (*DS – difference stationary*). Если некоторый ряд x_t принадлежит этому классу, то мы говорим о нем как о *DS-ряде*. Пусть ряд x_t – интегрированный порядка k . Подвергнем этот ряд k -кратному дифференцированию. Если в результате получается стационарный ряд типа **ARMA(p, q)**, то говорят, что исходный ряд x_t *является рядом типа ARIMA(p, k, q)*, или *k раз проинтегрированным ARMA(p, q) рядом* (*ARIMA – autoregressive integrated moving average*). Если при этом $p = 0$ или $q = 0$, то тогда употребляются и более короткие обозначения:

$$\begin{aligned} \text{ARIMA}(p, k, 0) &= \text{ARI}(p, k), \text{ARIMA}(0, k, q) = \text{IMA}(k, q), \\ \text{ARIMA}(0, k, 0) &= \text{ARI}(0, k) = \text{IMA}(k, 0). \end{aligned}$$

Возвращаясь к только что рассмотренным примерам, получаем:

$$X_t = \alpha + \beta t + \varepsilon_t \sim I(0);$$
$$X_t = \alpha + X_{t-1} + \varepsilon_t \sim I(1), X_t - \text{ряд типа } \mathbf{ARIMA(0, 1, 0)}.$$

Первый ряд является *TS*-рядом, а второй – *DS*-рядом.

При построении моделей связей между временными рядами в долгосрочной перспективе необходимо учитывать факт наличия или отсутствия у анализируемых макроэкономических рядов стохастического (недетерминированного) тренда. Иначе говоря, приходится решать вопрос об отнесении каждого из рассматриваемых рядов к классу рядов, стационарных относительно детерминированного тренда (или просто стационарных), – *TS*-ряды (trend stationary), или к классу рядов, имеющих стохастический тренд (возможно, наряду с детерминированным трендом) и приводящихся к стационарному ряду только путем однократного или *k*-кратного дифференцирования ряда – *DS*-ряды (difference stationary). Принципиальное различие между этими двумя классами рядов выражается в том, что в случае *TS*-ряда вычитание из ряда соответствующего детерминированного тренда приводит к стационарному ряду, тогда как в случае *DS*-ряда вычитание детерминированной составляющей ряда оставляет ряд нестационарным из-за наличия у него стохастического тренда.

TS-ряды имеют линию тренда в качестве некоторой «центральной линии», которой следует траектория ряда, находясь то выше, то ниже этой линии, с достаточно частой сменой положений выше – ниже. *DS*-ряды помимо детерминированного тренда (если таковой имеется), имеют еще и стохастический тренд, из-за присутствия которого траектория *DS*-ряда весьма долго пребывает по одну сторону от линии детерминированного тренда (выше или ниже этой линии), удаляясь от нее на значительные расстояния так, что по- существу в этом случае линия детерминированного тренда перестает играть роль «центральной» линии, вокруг которой колеблется траектория процесса.

В *TS*-рядах влияние предыдущих шоковых воздействий затухает с течением времени, а в *DS*-рядах такое затухание отсутствует, и каждый отдельный шок влияет с одинаковой силой на все последующие значения ряда. Поэтому наличие стохастического тренда требует проведения определенной экономической политики для возвращения макроэкономической переменной к ее долговременной перспективе, тогда как при отсутствии стохастического тренда серьезных усилий для достижения такой цели не требуется – в этом случае макроэкономическая переменная «скользит» вдоль линии тренда как направляющей, пересекая ее достаточно часто и не уклоняясь от этой линии сколь-нибудь далеко.

Для решения вопроса об отнесении исследуемого ВР к классу стационарных (относительно линейного тренда) или нестационарных процессов имеется ряд различных тестов. Однако все эти тесты страдают теми или иными недостатками. Тесты, сформулированные в виде формальных статистических критериев, как правило, имеют достаточно низкую мощность. Поэтому часто не отвергается исходная (нулевая) гипотеза, когда она в действительности не выпол-

няется. В то же время невыполнение теоретических предпосылок, на которых основывается критерий, при применении его к ВР приводит к отличию реально наблюдаемого размера критерия от заявленного уровня значимости. Из-за последнего обстоятельства теряется контроль над вероятностью ошибки первого рода, в результате чего нулевая гипотеза зачастую отвергается, когда в действительности она верна. С учетом этого исследователи при анализе рядов на принадлежность их к классу стационарных или нестационарных обычно используют не один, а несколько тестов и подкрепляют полученные выводы графическими процедурами.

При этом наиболее часто используются следующие тесты.

В тесте Дики – Фуллера [5] нулевой (альтернативной) гипотезой является тот факт, что исследуемый ВР x_t нестационарен (стационарен) и описывается одной из трех моделей авторегрессии первого порядка с поправкой на линейный тренд:

1) если ВР x_t имеет детерминированный линейный тренд, то оценивается модель

$$\Delta x_t = \phi x_{t-1} + \alpha + \beta_t + \varepsilon_t, t = 2, \dots, n;$$

2) если ВР x_t не имеет детерминированного тренда и его математическое ожидание не равно нулю, то берется модель

$$\Delta x_t = \phi x_{t-1} + \alpha + \varepsilon_t, t = 2, \dots, n;$$

3) если у ВР x_t нет детерминированного тренда и его математическое ожидание равно нулю, то выбирается модель

$$\Delta x_t = \phi x_{t-1} + \varepsilon_t, t = 2, \dots, n,$$

где ε_t – независимые случайные величины, имеющие одинаковое нормальное распределение с нулевым математическим ожиданием; ϕ , α , β – оцениваемые параметры.

С помощью МНК оцениваются параметры модели ϕ , α , β и вычисляется значение t -статистики t_ϕ для проверки нулевой гипотезы $\phi = 0$. Полученное значение сравнивается с критическим уровнем t_c . Гипотеза о нестационарности ВР отвергается, если $t_\phi < t_c$.

Мощность теста Дики – Фуллера существенно зависит от оцениваемой модели. С одной стороны, неправильный выбор оцениваемой модели может значительно отразиться на мощности этого теста. Так, если ВР порождается моделью случайного блуждания со сносом, а статистические выводы проводятся по результатам оценивания модели без включения трендовой составляющей, то мощность критерия стремится к нулю с возрастанием количества наблюдений [5]. С другой стороны, уменьшение мощности критерия может быть вызвано и избыточностью оцениваемой модели. Если же ВР описывается моделью более высокого порядка $p > 1$ и у характеристического многочлена не более одного единичного корня, для анализа данного ряда на стационарность применяется расширенный тест Дики – Фуллера (ADF-тест) [11], в котором в правые части каждой из трех рассмотренных для теста Дики – Фуллера моделей добав-

лены запаздывающие разности Δx_{t-j} , $t = 2, \dots, p - 1$. Полученные при оценивании моделей с добавленными запаздывающими разностями значения t -статистик t_ϕ для проверки нулевой гипотезы $\phi = 0$ сравниваются с теми же критическими значениями t_c , что и для теста Дики – Фуллера. Гипотеза о нестационарности ВР отвергается, если $t_\phi < t_c$. *ADF*-тест может использоваться и в том случае, когда ВР x_t описывается смешанной моделью авторегрессии и скользящего среднего.

Отметим, что тесты Дики – Фуллера нельзя применять к временным рядам, содержащим структурные и сезонные изменения, а также в случае гетероскедастичности ошибок наблюдения [12]. В указанных и в других случаях, связанных с нарушением модельных предположений, следует использовать специально разработанные тесты, например:

- тест Квятковского – Филлипса – Шмидта – Шина (KPSS-тест);
- тест Филлипса – Перрона (PP-тест) при нарушении гипотезы о некоррелированности и гомоскедастичности ошибок наблюдения в тестируемой модели;
- сезонный тест Дики – Гасжа – Фуллера (DNF-тест) при наличии сезонных эффектов в квартальных и месячных данных.

Методы построения и тестирования модели **ARIMA**.

Процесс построения модели **ARIMA** включает два этапа:

- 1) определение порядка интегрируемости нестационарного временного ряда и получение временного ряда разностей $\Delta^k x_t$;
- 2) подбор наилучшей модели для стационарного временного ряда разностей $\Delta^k x_t$ в классе моделей типа **ARMA**.

Раздел 5. Системы одновременных уравнений

При моделировании сложных экономических объектов часто приходится вводить не одно, а несколько связанных между собой уравнений. *Система одновременных уравнений* – это система взаимосвязанных регрессионных уравнений, причем для каждого момента времени среди эндогенных переменных есть такие, которые для одного регрессионного уравнения являются зависимыми переменными, а для другого – объясняющими переменными (регрессорами).

Уравнения, составляющие систему, делятся на две группы: *поведенческие* уравнения, описывающие эмпирические взаимосвязи между переменными, и уравнения *тождества*, не содержащие случайных ошибок и оцениваемых параметров (например, соотношения экономического баланса).

Системы эконометрических уравнений являются важным инструментом подготовки обоснованных решений. Они входят в состав эконометрических моделей принятия решений.

5.1. Структурная и приведенная формы уравнений

Макроэкономическая модель I является простейшей версией модели Кейнса. Модель построена на предположении, что народное хозяйство является системой закрытого типа без государственного регулирования экономики. Она состоит из двух уравнений: функции потребления и тождества, определяющего формирование национального дохода:

$$C_t = \alpha + \beta Y_t + \varepsilon_t, \quad (5.1)$$

$$Y_t = C_t + I_t, \quad (5.2)$$

где Y_t – национальный доход за период t ; C_t – личное потребление за период t ; I_t – частные инвестиции + государственные расходы + баланс внешней торговли в постоянных ценах за период t ; α, β – неизвестные параметры модели: α – свободный член функции потребления, выражающий автономное потребление, β – коэффициент регрессии, выражающий кратковременную предполагаемую склонность к потреблению; ε_t – случайная ошибка (возмущение) в функции потребления.

Уравнения (5.1), (5.2), составляющие исходную модель, называются *структурными*, или *структурной формой модели*. Первое уравнение относится к уравнению поведения, второе – к балансовому тождеству [2].

Проведем классификацию переменных в эконометрической модели: C_t, I_t – текущие эндогенные переменные, их значения определяются внутри модели; экзогенной является переменная, значение которой определяется вне модели и поэтому берется как заданное; в данном случае это переменная I_t . Наряду с экзогенными переменными в уравнение могут быть включены лаговые эндогенные. Данную модель можно представить стрелочной диаграммой (рис. 5.1).

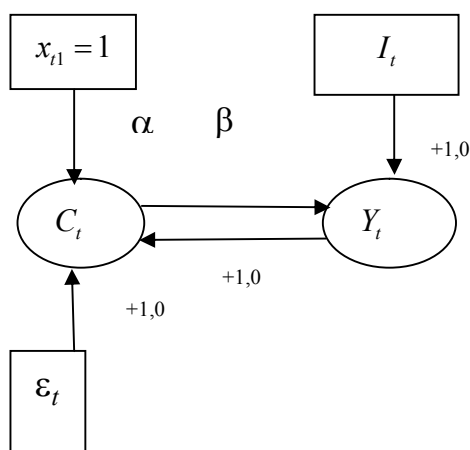


Рис. 5.1. Диаграмма структурной формы

Стрелочная диаграмма (рис. 5.1) отражает причинные связи структурной формы (5.1), (5.2). Зависимые переменные изображены окружностями, осталь-

ные – прямоугольниками. Структурные коэффициенты обозначены стрелками. Начало стрелки указывает на объясняющую переменную, конец – на объясняемую. Коэффициент переменной, которая просто суммируется, равен +1. Переменные C_t и Y_t связаны стрелками в обоих направлениях, т. е. они зависят друг от друга. Модели (5.1), (5.2) являются взаимозависимыми.

Для прогнозирования взаимосвязанных эконометрических показателей целесообразно от структурной формы переходить к приведенной форме. *Приведенные уравнения* – это уравнения, в которых эндогенные переменные выражены исключительно через предопределенные переменные и случайные составляющие:

$$C_t = \alpha + \beta(C_t + I_t) + \varepsilon_t$$

или

$$C_t = \frac{\alpha}{1-\beta} + \frac{\beta}{1-\beta} I_t + \frac{1}{1-\beta} \varepsilon_t.$$

Аналогичным образом строится приведенное уравнение для I_t .

Таким образом, получаем приведенную форму модели:

$$C_t = \frac{\alpha}{1-\beta} + \frac{\beta}{1-\beta} I_t + \frac{1}{1-\beta} \varepsilon_t, \quad (5.3)$$

$$Y_t = \frac{\alpha}{1-\beta} + \frac{1}{1-\beta} I_t + \frac{1}{1-\beta} \varepsilon_t, \quad (5.4)$$

где $\frac{\alpha}{1-\beta}$, $\frac{\beta}{1-\beta}$, $\frac{1}{1-\beta}$ – параметры ее приведенной формы.

Построим стрелочную диаграмму приведенной формы модели (рис. 5.2).

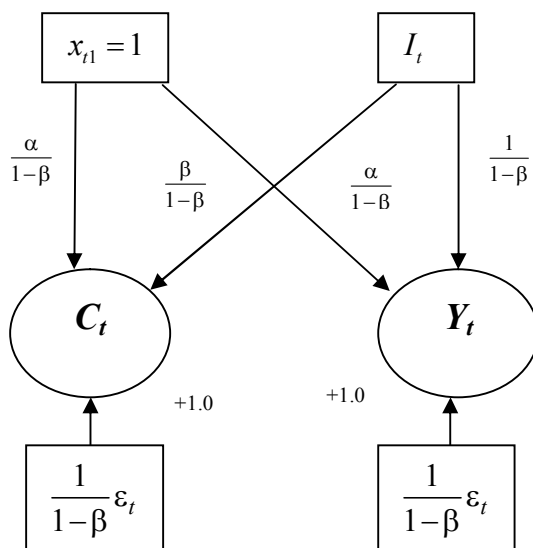


Рис. 5.2. Диаграмма приведенной формы

В правых частях уравнений приведенной формы встречаются только экзогенные переменные (к ним же относятся лаговые или запаздывающие эндогенные переменные), взаимозависимые переменные отсутствуют. На диаграмме это выражается отсутствием стрелок между зависимыми переменными. В левых частях уравнений (5.3), (5.4) стоят эндогенные переменные. Поэтому любое уравнение приведенной формы системы может использоваться в отдельности для прогнозирования зависимости переменной. Каждый коэффициент приведенной формы состоит из различных комбинаций коэффициентов структурной формы.

Структурная и приведенная формы модели дополняют друг друга в теоретико-экономическом смысле. Из модели приведенной формы видно, что функция потребления зависит от экзогенной переменной I_t . И она же позволяет выяснить, насколько сильно зависит C от I . Возмущение влияет сильнее, чем ожидалось.

В структурной форме коэффициенты носят прозрачный экономический смысл. Их легко интерпретировать в отличие от коэффициентов приведенной формы, среди которых не все допускают экономическую интерпретацию.

5.2. Мультипликаторы приведенной формы

Коэффициенты приведенной формы называются *мультипликаторами*. Все мультипликаторы одного периода составляют мультипликаторную матрицу. Запишем приведенную форму модели в матричном виде:

$$Y_t = \Pi^T X_t + V_t,$$

где

$$Y_t = \begin{bmatrix} y_{t1} \\ y_{t2} \end{bmatrix} = \begin{bmatrix} C_t \\ Y_t \end{bmatrix}, \quad X_t = \begin{bmatrix} x_{t1} \\ x_{t2} \end{bmatrix} = \begin{bmatrix} 1 \\ I_t \end{bmatrix}, \quad V_t = \begin{bmatrix} v_{t1} \\ v_{t2} \end{bmatrix} = \begin{bmatrix} \frac{\varepsilon_{t1}}{1-\beta} \\ \frac{\varepsilon_{t1}}{1-\beta} \end{bmatrix},$$

$$\Pi^T = \begin{bmatrix} \pi_{11} & \pi_{21} \\ \pi_{12} & \pi_{22} \end{bmatrix} = \begin{bmatrix} \frac{\alpha}{1-\beta} & \frac{\beta}{1-\beta} \\ \frac{\alpha}{1-\beta} & \frac{1}{1-\beta} \end{bmatrix}.$$

Чтобы сделать интерпретацию отдельных мультипликаторов более наглядной, используем оценки функции потребления, полученные обычным МНК на базе годовых данных за период 1961–1968 гг.:

$$\hat{\alpha} = 33,98 \quad \hat{\beta} = 0,64.$$

Получим оценку матрицы Π :

$$\hat{\Pi} = \begin{vmatrix} \hat{\pi}_{11} & \hat{\pi}_{12} \\ \hat{\pi}_{21} & \hat{\pi}_{22} \end{vmatrix} = \begin{vmatrix} 94,39 & 94,39 \\ 1,78 & 2,78 \end{vmatrix} = (\Pi_{ij}).$$

Для системы уравнений мультипликатор π_{ij} показывает, на сколько единиц при прочих равных условиях изменяется j -я эндогенная переменная, т. е. y_j , если в этот же период экзогенная переменная x_{ii} изменяется на 1. Таким образом, π_{ij} является мультипликатором i -й экзогенной переменной x_{ii} относительно j -й эндогенной переменной.

$$y_t = \begin{vmatrix} C_t \\ y_t \end{vmatrix} = \begin{vmatrix} \text{потребление} \\ \text{нац. доход} \end{vmatrix}, \quad x_t = \begin{vmatrix} 1 \\ I_t \end{vmatrix} = \begin{vmatrix} \text{автономная ситуация} \\ \text{инвестиции} \end{vmatrix}.$$

Пусть x_{ii} – инвестиции, y_{ij} – потребление, тогда π_{ij} – инвестиционный мультипликатор потребления.

Поскольку π_{ij} выражает наступающий без задержек полный эффект от воздействия i -й экзогенной переменной на j -ю эндогенную переменную, его также называют *мультипликатором первоначального действия*.

Для рассматриваемой модели имеем:

$\hat{\pi}_{22} = 2,78$ – мультипликатор инвестиций относительно национального дохода, т. е. если при прочих равных условиях инвестиции увеличиваются или уменьшаются на 1, то национальный доход уменьшается или увеличивается на 2,78;

$\hat{\pi}_{21} = 1,78$ – мультипликатор инвестиций относительно потребления, если при прочих равных условиях инвестиции уменьшаются на 1, то функция потребления изменяется в том же направлении на 1,78;

$\hat{\pi}_{11} = 94,39$ – мультипликатор автономного потребления, если x_{1t} изменяется на 1, то изменяется потребительское равновесие на 94,39;

$\hat{\pi}_{12} = 94,39$ – мультипликатор автономного потребления национального дохода.

Все рассматриваемые мультипликаторы выражают общий эффект воздействия объясняющих переменных на зависимые переменные в отличие от структурных коэффициентов, которые показывают эффекты частичного воздействия.

Пример. Макроэкономическая модель (Людике, 1964)

Модель Людике реалистично описывает экономику открытого типа с государственным регулированием:

$$C_t = a_0 + a_1 Y_t + a_2 C_{t-1} + \varepsilon_{t1}; \quad (5.5)$$

$$I_t = b_0 + b_1 Y_t + b_2 D_{t-1} + \varepsilon_{t2}; \quad (5.6)$$

$$J_t = c_0 + c_1 Y_t + c_2 J_{t-1} + \varepsilon_{t3}; \quad (5.7)$$

$$Y_t = C_t + I_t - J_t + G_t. \quad (5.8)$$

Эндогенными переменными модели являются: $C_t = y_{t1}$ – личное потребление; $I_t = y_{t2}$ – частные чистые инвестиции и основной капитал; $J_t = y_{t3}$ – импорт; $Y_t = y_{t4}$ – национальный доход (в текущих ценах) [2].

Предопределенные переменные: $x_{t1} = 1$, $x_{t2} = C_{t-1}$, где C_{t-1} – лаговая эндогенная переменная (личное потребление предшествующего периода), $x_{t3} = D_{t-1}$ – запаздывающие на один период доходы населения от предпринимательской деятельности, дивиденды и нераспределенная прибыль предприятия до налогообложения, $x_{t4} = J_{t-1}$ – лаговый импорт, $x_{t5} = G_t$ – государственные расходы + государственные чистые инвестиции и основной капитал + изменения в товарных запасах + экспорт + субсидии – косвенные налоги.

Все переменные в модели представлены в постоянных ценах. Уравнения (5.5) – функция потребления, (5.6) – функция инвестиций, (5.7) – функция импорта являются *уравнениями поведения*. Уравнение (5.8) представляет собой тождество, или *балансовое уравнение*.

Модель Людике в целом больше соответствует экономической теории, чем модель (5.1), (5.2), поскольку последняя пытается объяснить зависимую переменную функцией только одной объясняющей переменной.

Благодаря тому что G_t – экзогенная переменная, появляется возможность количественно измерить влияние государственных расходов на зависимые переменные модели.

Заметим, что в настоящее время в экономических и социальных исследованиях применяется намного более крупные системы уравнений (например, в Германском Федеральном Банке при прогнозировании конъюнктуры применяется более 100 уравнений).

5.3. Идентификация систем линейных одновременных уравнений

Проблема оценивания параметров уравнений структурной формы и возможность преобразования структурной формы в приведенную тесно связана с понятием идентификации модели.

Исходную систему уравнений называют *идентифицируемой* (точно определенной), если по коэффициентам приведенных уравнений можно однозначно определить значения коэффициентов структурных уравнений.

Исходную систему уравнений называют *неидентифицируемой* (недоопределенной), если по коэффициентам приведенных уравнений невозможно определить значения коэффициентов структурных уравнений.

Исходную систему уравнений называют *сверхидентифицируемой* (переопределенной), если по коэффициентам приведенных уравнений можно получить несколько вариантов значений коэффициентов структурных уравнений.

Для быстрого формального определения идентифицируемости уравнений структурной формы применяют следующие необходимые и достаточные условия.

Пусть система одновременных уравнений включает в себя N уравнений относительно N эндогенных переменных, и M – число predetermined переменных в системе. Пусть также количество эндогенных и predetermined переменных в проверяемом уравнении n и m соответственно. Тогда количество predetermined переменных, не включенных в проверяемое уравнение, но присутствующих в системе, равно $M - m$. Необходимым условием идентификации (условие порядка) уравнения является

$$M - m \geq n - 1.$$

При анализе структурного уравнения выделяют три случая:

1) если $M - m = n - 1$, уравнение точно идентифицируемо, и для его оценивания применяется косвенный МНК;

2) если $M - m > n - 1$, уравнение сверхидентифицируемо, и для оценивания используется двухшаговый МНК;

3) уравнение может быть не полностью идентифицировано, когда $M - m < n - 1$.

Если система точно идентифицируема, то и косвенный, и двухшаговый МНК дают одинаковые оценки.

Достаточное условие идентифицируемости (условие ранга): определитель матрицы, составленной из коэффициентов при переменных, отсутствующих в исследуемом уравнении, не равен нулю, и ранг этой матрицы – не менее числа эндогенных переменных системы без единицы.

5.4. Методы решения систем одновременных уравнений

Двухшаговый МНК (2МНК). Идея 2МНК (за два шага получить оценку параметров системы структурной формы) заключается в следующем:

1. Составляют приведенную форму модели и определяют численные значения параметров каждого уравнения обычным МНК.

2. Рассчитывают теоретические значения $\hat{Y}$ эндогенных переменных.

3. С помощью МНК определяют параметры структурных уравнений, заменяя в правой части Y на $\hat{Y}$.

Косвенный МНК (КМНК). Сущность его заключается в следующем:

1. Первый этап КМНК совпадает с первым этапом 2МНК.

2. На втором этапе путем линейных преобразований из оценок параметров приведенной формы получают оценки параметров уравнений систем структурной формы.

Раздел 6. Практические приложения: построение и оценивание производственных функций

6.1. Свойства производственной функции

Производственная функция (ПФ) представляет собой математическую функцию, отражающую связь между факторами и результатами производства. Производственные функции применяются в исследованиях различных уровней экономики в зависимости от характера исходных данных. В одних случаях они описывают отдельные технологические процессы, в других – отражают деятельности предприятия, отрасли или экономики страны в целом.

Производственная функция отражает устойчивую количественную связь, существующую между затратами и выпуском продукции, безотносительно к содержанию происходящих при этом реальных производственных процессов.

Рассмотрим два фактора производства – затраты капитала (K) и трудовые затраты (L). Если Q – объем выпускаемой продукции, то производственная функция может быть записана следующим образом:

$$Q = f(K, L). \quad (6.1)$$

Классическая теория производства предполагает, что предельные продукты капитала и труда положительны, а численные их значения убывают. Функция (6.1) является однозначной, непрерывной и дважды дифференцируемой. Для нее выполняются следующие соотношения:

$$\frac{\partial Q}{\partial K} = Q_K > 0; Q_{KK} < 0; \frac{\partial Q}{\partial L} = Q_L > 0; Q_{LL} < 0,$$

где Q_L и Q_K – соответственно предельные продукты капитала и труда.

Кривые, описываемые условием $f(K, L) = \text{const}$, обычно называют изоквантами. Изокванты имеют отрицательный наклон, т. е.

$$-\frac{\partial K}{\partial L} = \frac{Q_L}{Q_K} = R > 0,$$

поскольку для малых приращений вдоль каждой изокванты можно утверждать, что

$$\frac{\partial f}{\partial K} dK + \frac{\partial f}{\partial L} dL = Q_K dK + Q_L dL = 0.$$

Численное значение углового коэффициента касательной R характеризует предельную норму замещения. Предельная норма замещения представляет со-

бой отношение, в котором использование одного фактора производства может быть уменьшено, а другого – увеличено без изменения выпускаемого объема продукции. По мере роста значений одной из независимых переменных предельная норма замещения данного фактора производства уменьшается. При сохранении постоянного объема производства отношение dK/dL должно увеличиваться и d^2K/d^2L будет положительным.

Мерой скорости изменения R служит эластичность замещения факторов K и L , которая определяется как отношение изменения пропорции K/L к изменению величины R :

$$\sigma = \frac{(-K/L)d(K/L)}{dR/R},$$

причем все приращения отсчитываются вдоль кривой постоянного выпуска. С помощью показателя σ можно измерить, насколько легко заместить фактор L фактором K . При $\sigma = 0$ замена невозможна. Если $\sigma = \infty$, кривая постоянного выпуска вырождается в прямую.

Другое свойство производственной функции связано с характеристикой эффективности при изменении масштабов производства. Если объем всех затрат увеличивается в λ раз и при этом во столько же раз увеличивается объем выпускаемой продукции, то эффективность не зависит от масштабов производства, производственная функция является линейно-однородной и имеет место соотношение

$$f(\lambda K, \lambda L) = \lambda f(K, L) = \lambda Q.$$

Заметим, что если оплата каждого из используемых факторов производства совпадает с размерами предельного продукта соответствующего фактора, то на покрытие этих расходов будет направлен весь производственный продукт Q , т. е. выполняется равенство

$$Q = Q_K K + Q_L L.$$

Полезные свойства производственных функций можно получить, если предположить, что целью хозяйственной деятельности предпринимателя является максимизация прибыли.

Пусть VD обозначает размеры валового дохода; SI – совокупные издержки; π – величину чистого дохода. Тогда предприниматель выбирает такие значения K , Q , L , которые максимизируют величину π при ограничении, налагаемом производственной функцией, а именно:

$$\pi = VD - SI + \mu f(K, L) - Q, \quad (6.2)$$

где μ – неизвестный множитель Лагранжа.

Если предположить, что на рынке труда и капитала господствуют условия совершенной конкуренции, то совокупные издержки равны

$$SI = w_L L + w_K K, \quad (6.3)$$

где w_L обозначает ставку заработной платы, выплачиваемой за единицу труда; w_K – цену единицы «услуг капитала», причем w_L и w_K представляют собой постоянные величины.

На рынке готовой продукции не существует совершенной конкуренции. Пусть кривая спроса характеризуется постоянной эластичностью $\frac{1}{2}$. Тогда цена p готовой продукции определяется соотношением $p = bQ^\eta$, где b – постоянная величина. Размер валового дохода вычисляется по формуле

$$VD = pQ = bQ^{\eta+1}.$$

Если $\eta = 0$, то цена p представляет собой постоянную величину (ситуация совершенной конкуренции на рынке готовой продукции). Если $\eta = -1$, то совокупный доход VD не зависит от изменения величин p или Q и оказывается постоянным. В этом случае максимизация π в (6.2) эквивалентна минимизации величины SI в (6.3). Но SI достигает минимума при $L = 0$, $K = 0$ и $Q = 0$. Следовательно, при $\eta = -1$ объем производства должен составлять некоторую заранее заданную величину Q_0 .

Используя соотношения (6.2), (6.3), функцию прибыли можно записать как

$$\pi = bQ^{\eta+1} - w_L L - w_K K + \mu(f(K, L) - Q).$$

Необходимыми условиями максимизации прибыли первого порядка являются следующие:

$$\begin{aligned} \frac{\partial \pi}{\partial L} &= -w_L + \lambda f_L = 0, \quad f_L = \frac{\partial f}{\partial L}; \\ \frac{\partial \pi}{\partial K} &= -w_K + \lambda f_K = 0, \quad f_K = \frac{\partial f}{\partial K}; \\ \frac{\partial \pi}{\partial \mu} &= f(K, L) - Q = 0. \end{aligned}$$

Если Q не представляет собой заранее заданную величину, то добавляется еще одно соотношение:

$$\frac{\partial \pi}{\partial Q} = (\eta + 1)bQ^\eta - \mu = 0.$$

Условия максимизации прибыли второго порядка позволяют получить ряд априорных ограничений на параметры различных производственных функций. Рассмотрим условия второго порядка. Если Q – постоянная величина ($Q = Q_0$), то эти условия имеют вид

$$\begin{vmatrix} \pi_{LL} & \pi_{LK} & f_L \\ \pi_{KL} & \pi_{KK} & f_K \\ f_L & f_K & 0 \end{vmatrix}$$

или

$$D = 2f_{KL}f_Lf_K - f_{LL}f_K^2 - f_{KK}f_L^2. \quad (6.4)$$

Если Q является переменной величиной, то эти условия записываются как

$$\begin{vmatrix} \pi_{LL} & \pi_{LK} & \pi_{LK} & f_L \\ \pi_{KL} & \pi_{KK} & \pi_{KK} & f_K \\ \pi_{KK} & \pi_{LK} & \pi_{KL} & -1 \\ f_L & f_K & -1 & 0 \end{vmatrix} < 0,$$

что равносильно соотношению

$$pD(\eta + 1)Q^{-1} - \kappa(f_{KK}f_{LL} - f_{KL}^2) < 0,$$

где D определяется выражением (6.4). Для случая совершенной конкуренции ($\eta = 0$) условие второго порядка принимает вид

$$f_{KK}f_{LL} - f_{KL}^2 > 0. \quad (6.5)$$

Производственная функция Кобба – Дугласа

Производственная функция Кобба – Дугласа, включающая два фактора производства, в общем виде может быть записана следующим образом:

$$Q = AK^\alpha L^\beta, \quad (6.6)$$

где A , α и β – параметры модели.

Эта функция однозначна и непрерывна при положительных K и L , а предельные продукты капитала и труда равны соответственно

$$Q_k = \frac{\alpha Q}{K}, \quad Q_L = \frac{\beta Q}{L}. \quad (6.7)$$

Из (6.7) следует, что $\alpha > 0$ и $\beta > 0$, т. к. Q , K и L положительны. Из требований отрицательности вторых частных производных ($Q_{KK} < 0$, $Q_{LL} < 0$) производственной функции можно получить, что $\alpha < 1$ и $\beta < 1$. Нетрудно показать, что предельная норма замещения определяется формулой

$$R = \frac{Q_L}{Q_K} = \frac{\beta K}{\alpha L},$$

а эластичность замещения этой функции равна 1. Степень однородности функции (6.6) равна $\alpha + \beta$. Действительно, если затраты каждого фактора увеличить в λ раз, то новое значение Q (обозначим его через $Q^{(\lambda)}$) будет определяться следующим образом:

$$Q^{(\lambda)} = A(\lambda K)^\alpha (\lambda L)^\beta = \lambda^{\alpha+\beta} Q.$$

При $\lambda + \beta = 1$ уровень эффективности не зависит от масштабов производства. Если $\lambda + \beta < 1$, с расширением масштабов производства его эффективность снижается. При $\lambda + \beta > 1$ с расширением масштабов производства повышается и его эффективность.

В условиях совершенной конкуренции из соотношения (6.5) можно вывести неравенство

$$\frac{\alpha\beta Q^2}{K^2 L^2} (1 - \alpha - \beta) > 0,$$

которое выполняется при $\lambda + \beta < 1$. Следовательно, конечный максимум прибыли может быть достигнут только при убывании эффективности по мере увеличения масштабов производства.

В ситуации несовершенной конкуренции из условий второго порядка выводится соотношение

$$\frac{p(\eta+1)\alpha\beta Q^2}{K^2 L^2} [(\alpha + \beta)(\eta+1) - 1] < 0.$$

Поскольку величины p , α , β , Q и L положительны, необходимо одновременно выполнение условий

$$(\eta+1) > 0 \text{ и } (\lambda+\beta)(\eta+1) < 1.$$

Из этого следует, что спрос на производственную продукцию должен быть неэластичным, т. е. $-1 < \eta < 0$, и, кроме того, $\eta + 1 < \frac{1}{\alpha + \beta}$.

Производственная функция с постоянной эластичностью замещения (CES-функция)

Модель производственной функции с постоянной эластичностью замещения имеет следующий вид:

$$Q = \gamma[(1 - \delta)K^{-\rho} + \delta L^{-\rho}]^{-\gamma/\rho}, \quad (6.8)$$

где γ , δ , ρ , ν – параметры модели. Поскольку предельные продукты труда и капитала

$$Q_K = \frac{1 - \delta}{K^{1+\rho}} A Q^{1+\rho/\nu}, \quad Q_L = \frac{\delta}{L^{1+\rho}} A Q^{1+\rho/\nu} \quad (6.9)$$

положительны, то $(1 - \delta)A > 0$ и $A\delta > 0$. Следовательно, имеют место неравенства $1 > \delta > 0$ и $A\delta > 0$.

Требование отрицательности частных производных Q_{KK} и Q_{LL} и условие (6.9) приводят к неравенству $1 + \rho > 1 + \rho/\nu$, откуда либо $\rho > -1$ и $\rho/\nu > 0$, либо $\rho < -1$ и $\rho < -1$.

Параметр ρ характеризует эластичность $\sigma = \frac{1}{1 + \rho}$, которая, в свою очередь, не зависит от значений K , Q , L . Поскольку $\sigma > 0$, значение ρ не должно превышать (-1) и, следовательно, $\rho/\nu > 0$.

Параметр ν производственной функции (6.8) характеризует поведение эффективности производства. При увеличении затрат K и L в λ раз объем выпускаемой продукции изменится следующим образом:

$$Q^{(\lambda)} = \gamma[(1 - \delta)\lambda^{-\rho}K^{-\rho} + \delta\lambda^{-\rho}L^{-\rho}]^{-\gamma/\rho} - \lambda^\gamma Q.$$

При $\gamma = 1$ эффективность не зависит от изменения масштабов производства, при $\gamma > 1$ с расширением масштабов производства повышается его эффективность, при $\gamma < 1$ с расширением масштабов производства его эффективность снижается.

Для CES-функции отношение между оплатой труда и доходом на капитал зависит от величин δ , ρ и показателя капиталовооруженности труда, а именно:

$$\frac{w_L L}{w_K K} = \left(\frac{\delta}{1 - \delta} \right) \left(\frac{K}{L} \right)^\rho.$$

В ситуации совершенной конкуренции конечный максимум прибыли может быть достигнут лишь при $v < 1$. Для случая несовершенной конкуренции можно вывести условия $\eta + 1 > 0$, $v(\eta + 1) < 1$.

Окончательно априорные ограничения на параметры CES-функции могут быть записаны следующим образом:

$$\rho/v > 0; 1 > \delta > 0; \rho > -1; v > 0.$$

6.2. Оценивание параметров производственных функций. Проблемы и пути их решения

Для определения параметров и вида производственной функции необходимо провести статистические наблюдения. Как правило, пользуются двумя видами данных – временными рядами и данными одновременных наблюдений. Динамические (временные) ряды характеризуют поведение одной и той же фирмы во времени. Данные второго вида обычно относятся к одному и тому же моменту, но к различным фирмам.

Основные проблемы при использовании временных рядов возникают вследствие того, что со временем меняются не только параметры модели (относительная цена, относительное сочетание затрат), но и сама форма производственной функции (последствия технического прогресса). По мере развития науки и техники возникают новые, более эффективные технологические процессы, улучшается подготовка и повышается качество рабочей силы, меняются нормы затрат производственных факторов, параметры эффективности производственной функции и ее поведение при переходе к другим масштабам производства.

При оценивании параметров производственных функций, меняющихся со временем детерминированным или случайным образом, применяются методы идентификации регрессионных моделей с переменными параметрами. Для учета технического прогресса предлагаются два метода. Согласно первому из них технический прогресс учитывается в форме некоего временного тренда, включаемого в состав производственной функции $Q_t = f(t, K_t, L_t)$. Например, в производственной функции Кобба – Дугласа технический прогресс учитывается следующим образом:

$$Q_t = Ae^{\theta t} K_t^\alpha L_t^\beta.$$

Параметр в экспоненциальной функции показывает, что объем выпускаемой продукции ежегодно увеличивается (на θ %), независимо от изменений в затратах производственных факторов, независимо от размеров новых инвестиций.

Другой подход предполагает, что технический прогресс материализуется в инвестициях и оказывает влияние на объем выпускаемой продукции только после того, как эти капиталовложения будут осуществлены. В этом случае применяют

модели, учитывающие возрастной состав оборудования, поскольку каждой возрастной группе капитала соответствует своя производственная функция:

$$Q_t(v) = f(v, K_t(v), L_t(v)),$$

где $Q_t(v)$ – продукция, произведенная в период t на оборудовании, введенном в строй в период v ; $L_t(v)$ – труд, занятый в период t обслуживанием оборудования, введенного в строй в период v ; $K_t(v)$ – основной капитал, введенный в строй в период v и используемый в период t . Параметр v в такой производственной функции отражает состояние технического прогресса. Так, функция Кобба – Дугласа при этом подходе выглядит следующим образом: $Q_t = Ae^{\theta v} K_t^\alpha(v) L_t^\beta(v)$. Показатели совокупного объема выпускаемой продукции, общие затраты труда и капитала на момент t имеют вид

$$Q_t = \int_{-\infty}^t Q_t(v) dv, \quad L_t = \int_{-\infty}^t L_t(v) dv, \quad K_t = \int_{-\infty}^t K_t(v) dv.$$

При этом $Q_t = f(L_t, K_t)$ представляет собой агрегированную производственную функцию.

При использовании данных обследований фирм, относящихся к одному и тому же времени, относительные цены будут оставаться прежними и можно не принимать во внимание последствия технического прогресса. Однако предполагается, что поведение всех фирм может быть описано с помощью одной и той же функции, т. е. фирмы должны принадлежать одной и той же отрасли. Для того чтобы на уровне отрасли наблюдать те же соотношения между выпуском продукции и затратами, что и на уровне отдельной фирмы, производственные функции должны быть аддитивно сепарабельны, т. е. представимы в виде

$$F(K, L) = g(K) + h(L).$$

Условие аддитивной сепарабельности выполняется для рассмотренных выше видов производственных функций, если от функции Кобба – Дугласа перейти к функции

$$\log Q = \log A + \alpha \log K + \beta \log L = g(K) + h(L),$$

где $g(K) = \log A + \log K$; $h(L) = \beta \log(L)$, а вместо CES-функции работать с функциями

$$Q^{-\rho/\nu} = \gamma^{-\rho/\nu} (1 - \delta) K^{-\rho} + \gamma^{-\rho/\nu} \delta L^{-\rho} = g(K) + h(L),$$

где $g(K) = \gamma^{-\rho/\nu} (1 - \delta) K^{-\rho}$; $h(L) = \gamma^{-\rho/\nu} \delta L^{-\rho}$.

Если представить в логарифмической форме производственные функции для n фирм, то на агрегированном уровне мерой Q может служить логарифм среднего геометрического, составленного из $Q_1, \dots, Q_n$, умноженного на n :

$$\log\left(\prod_{i=1}^n Q_i\right) = \log\left(\prod_{i=1}^n A_i\right) + \log\left(\prod_{i=1}^n K_i^{\alpha_i}\right) + \log\left(\prod_{i=1}^n Z_i^{\beta_i}\right).$$

Если параметры производственных функций всех фирм совпадают ($\alpha_i = \alpha, \beta_i = \beta, i = 1, \dots, n$), то для получения агрегированных значений капитала и труда требуются средние геометрические, составленные из соответствующих показаний для отдельных фирм:

$$\log\left(\prod_{i=1}^n Q_i\right) = \log\left(\prod_{i=1}^n A_i\right) + \alpha \log\left(\prod_{i=1}^n K_i\right) + \beta \log\left(\prod_{i=1}^n Z_i\right).$$

На функции CES это утверждение не распространяется, за исключением частного случая ($\delta_1 = \delta_2 = 0$).

При работе с данными одновременных обследований основную проблему представляет интерпретация производственной функции. На основании оцененной функции нельзя получить достаточную информацию о степени замещаемости производственных затрат. Выявляемое влияние масштабов производства на уровень эффективности будет зависеть от статистической характеристики рассматриваемых отраслей.

Данные обследования могут относиться к различным отраслям различных стран. Тогда даже в случае наблюдений, относящихся к одному и тому же моменту времени, может возникнуть проблема учета научно-технического прогресса, поскольку экономика этих стран может характеризоваться неодинаковым уровнем технического прогресса. Возникают трудности с измерением Q, K и L в одних и тех же единицах, особенно в тех случаях, когда валютный курс не совсем точно отражает реальную покупательскую способность различных валют. Размеры аналогичных отраслей в разных странах также отличаются друг от друга, поэтому при статистическом оценивании может возникнуть проблема гетероскедастичности.

6.3. Способы эмпирического оценивания производственных функций разных видов

В качестве основных переменных производственной функции выступают затраты труда, затраты капитала и объем выпускаемой продукции, причем с теоретической точки зрения единицы измерений всех переменных должны быть однородны. На практике это невыполнимо.

Все рассмотренные нами виды производственных функций носили детерминированный характер, они не учитывали влияние случайных возмущений. В каждое уравнение, параметры которого предстоит оценить, необходимо ввести случайную переменную U , которая будет отражать воздействие на процесс производства всех тех факторов, которые не вошли в состав производственной функции в явном виде:

$$Q = A(K, L) + U.$$

Параметры функции Кобба – Дугласа неудобно непосредственно оценивать по методу наименьших квадратов. Как правило, прибегают к логарифмическому преобразованию:

$$\log Q = \log A + \alpha \log K + \beta \log L + U.$$

Если заранее предположить, что с применением масштабов производства уровень эффективности остается прежним ($\alpha + \beta = 1$, $\beta = 1 - \alpha$), то получаем

$$Q = AK^\alpha L^{1-\alpha} e^U = A(K/L)^\alpha \alpha e^U$$

или

$$\log(Q/L) = \log A + \alpha \log K/L + U.$$

Этот подход помогает устранить влияние мультиколлинеарности между $\log K$ и $\log L$, а также уменьшить влияние гетероскедастичности в тех случаях, когда дисперсия шума U коррелирует с L .

Не существует достаточно простого преобразования, приводящего CES-функцию к удобной для оценивания форме. Поэтому оценивание CES-функции сводится либо к проверке выполнения условий максимизации прибыли ($\omega/r = (\delta/(1 - \delta))(K/L)^{1+\rho}$), либо к аппроксимации этой функции.

Используя условие максимизации прибыли, можно для функции CES получить регрессионное уравнение относительно капитала:

$$\log\left(\frac{Q}{K}\right) = \log C' + \frac{1}{1+\rho} \log\left(\frac{r}{\rho}\right) + \frac{\rho(v-1)}{\gamma(1+\rho)} \log(Q) + U,$$

где C' – постоянная величина.

Для функции Кобба – Дугласа из того факта, что $\omega/r = \beta^{k/\alpha} \alpha$, следует

$$\log(Q/L) = -\log \alpha + \log(r/\rho) + U.$$

Указанные выражения позволяют провести соответствующие эмпирические проверки.

Для функции CES регрессионное уравнение относительно затрат труда имеет вид

$$\log\left(\frac{Q}{K}\right) = \log C + \frac{1}{1+\rho} \log\left(\frac{\omega}{\rho}\right) + \frac{\rho(v-1)}{\gamma(1+\rho)} \log(Q) + U,$$

где C – постоянная величина.

Для функции Кобба – Дугласа аналогичное уравнение запишется как

$$\log(Q/C) = \log C + \log(\omega/\rho) + U.$$

Если при оценивании окажется, что коэффициент $\log(\omega/\rho)$ близок к единице, то есть большие основания считать, что производственный процесс описывается моделью Кобба – Дугласа.

Рассмотрим один из методов аппроксимации, основанный на разложении некоторой функции в ряд Тейлора. Представим CES-функцию и случайный остаток, входящий в виде сомножителя e^u , следующим образом:

$$Q/L = e^u \gamma L^{v-1} [\delta + (1-\delta)(K/L)^{-\rho}]^{-v/\rho}.$$

Прологарифмировав это выражение, получим

$$\log(Q/L) = \log v + (v-1)\log L - (v/\rho)f(\rho) + U,$$

где $f(\rho) = \log[\delta + (1-\delta)(K/L)^{-\rho}]$.

Функцию $f(\rho)$ можно разложить в ряд Тейлора:

$$f(\rho) = f(0) + \rho f'(0) + \frac{\rho^2}{2!} f''(\rho) + \dots$$

и получить приближенное выражение для $f(\rho)$:

$$f(\rho) = -\rho(1-\delta)\log(K/L) + \frac{\rho^2}{2}\delta(1-\delta)(\log(K/L))^2.$$

Следовательно, для CES-функции имеем регрессионное уравнение

$$\log(Q/L) = \log \gamma + (v-1)\log L + v(1-\delta)\log(K/L) - 0,5v\delta\rho(1-\delta)(\log(K/L))^2 + U.$$

Можно получить аналогичное уравнение для функции Кобба – Дугласа:

$$\log(Q/L) = \log A + (\beta - 1 + \alpha)\log L + \alpha\log(K/L) + U.$$

В заключение отметим проблемы, возникающие при построении эконометрических моделей.

Последствия отсутствия в уравнении существенной независимой переменной. Если в уравнение регрессии не включена независимая переменная, оказывающая существенное влияние на результирующий признак, то в общем случае это приводит к смещению оценок коэффициентов регрессии. Смещение отсутствует только в случае, если ковариация отсутствующей переменной с переменными, включенными в модель, равна нулю. Стандартные ошибки коэффициентов регрессии становятся некорректными, что приводит к неприменимости соответствующих t -тестов. Кроме того, возможно появление автокорреляции и гетероскедастичности остатков. Признаком отсутствия значимой переменной может служить несоответствие знаков коэффициентов теоретическим предположениям. Если нет возможности включить в уравнение регрессии такую переменную, то следует использовать замещающую переменную.

Последствия включения в модель несущественной независимой переменной. Если в уравнение регрессии включена существенная независимая переменная, то в общем случае это не приводит к смещению оценок коэффициентов регрессии, но значения стандартных ошибок могут возрасти.

Последствия неправильной спецификации формы уравнения регрессии. Использование неверной формы уравнения регрессии приводит к смещенности и несостоятельности оценок параметров, низкому значению коэффициента детерминации R^2 . Возможно также появление автокорреляции и гетероскедастичности остатков.

Литература

1. Афанасьев, В. Н. Анализ временных рядов и прогнозирование : учебник / В. Н. Афанасьев, М. М. Юзбашев. – М. : Финансы и статистика, 2001. – 228 с.
2. Бородич, С. А. Эконометрика / С. А. Бородич. – Минск : Новое знание, 2005. – 416 с.
3. Винн, Р. Введение в прикладной эконометрический анализ / Р. Винн, К. Холден. – М. : Финансы и статистика, 1981. – 254 с.
4. Джонстон, Дж. Эконометрические методы / Дж. Джонстон. – М. : Статистика, 1980. – 514 с.
5. Кравцов, М. К. Эконометрический анализ временных рядов основных экономических показателей / М. К. Кравцов, А. В. Пашкевич, Н. М. Бурдыко // Белорусская экономика: анализ, прогноз, регулирование. – 2003. – № 3 (93) – С. 3–22.
6. Лукашин, Ю. П. Линейная регрессия с переменными параметрами / Ю. П. Лукашин. – М. : МГУ, 1992. – 230 с.
7. Магнус, Я. Р. Эконометрика. Начальный курс : учебник / Я. Р. Магнус, П. К. Катышев, А. А. Пересецкий. – 7-е изд., испр. – М. : Дело, 2005. – 576 с.
8. Носко, В. П. Эконометрика. Введение в регрессионный анализ временных рядов / В. П. Носко. – М. : 2002. – 273 с.
9. Салманов, О. Н. Эконометрика : учеб. пособие / О. Н. Салманов. – М. : Экономистъ, 2000. – 320 с.
10. Эконометрия / В. И. Суслов [и др.]. – Новосибирск : СО РАН, 2005. – 744 с.
11. Харин, Ю. С. Эконометрическое моделирование: учеб. пособие / Ю. С. Харин, В. И. Малюгин, А. Ю. Харин. – Минск : БГУ, 2003. – 313 с.
12. Харин, Ю. С. Математические и компьютерные основы статистического анализа данных и моделирования : учеб. пособие / Ю. С. Харин, В. И. Малюгин, М. С. Абрамович. – Минск : БГУ, 2008. – 455 с.
13. Хацкевич, Г. А. Эконометрика : учеб.-метод. комплекс для студ. эконом. спец. / Г. А. Хацкевич, А. Б. Гедранович. – Минск : Изд-во МИУ, 2005. – 252 с.
14. Шанченко, Н. И. Лекции по эконометрике : учеб. пособие / Н. И. Шанченко. – Ульяновск : УлГТУ, 2008. – 139 с.

Учебное издание

Алёхина Алина Энодиевна
Поттосина Светлана Анатольевна

ЭКОНОМЕТРИКА

УЧЕБНО-МЕТОДИЧЕСКОЕ ПОСОБИЕ

Редакторы *Т. П. Андрейченко, Е. И. Герман*

Корректоры *А. В. Бас, Е. Н. Батурчик*

Компьютерная правка и оригинал-макет *Ю. Ч. Ключкевич, А. А. Лысеня*

Подписано в печать 18.09.2013. Формат 60х84 1/16. Бумага офсетная. Гарнитура «Таймс».
Отпечатано на ризографе. Усл. печ. л 5,93. Уч.-изд. л. 5,0. Тираж 150 экз. Заказ 452.

Издатель и полиграфическое исполнение: учреждение образования
«Белорусский государственный университет информатики и радиоэлектроники»
ЛИ №02330/0494371 от 16.03.2009. ЛП №02330/0494175 от 03.04.2009.
220013, Минск, П. Бровки, 6