

13.2. Анализ временных рядов в SPSS

13.2.1. Обзор возможностей

Возможности в области анализа временных рядов. Пользователи, знакомые с ранними и неполными версиями пакета SPSS, чаще всего имеют совершенно неадекватное представление о возможностях этого пакета в области анализа временных рядов. Во-первых, ранние версии SPSS использовались, в основном, специалистами в области психологии, биологии и социологии, где задачи анализа временных рядов менее характерны. Во-вторых, современные версии пакета имеют модульную структуру (см. приложения 1 и 2), в которой анализ временных рядов выделен в отдельный модуль SPSS Trends. При отсутствии этого модуля возможности в области анализа временных рядов будут ограничены только процедурами регрессионного анализа базового модуля пакета SPSS Base.

Кратко перечислим основные процедуры пакета SPSS в области анализа временных рядов:

1. **Regression (Регрессионный анализ)** — позволяет выделять и удалять широкий набор моделей трендов;

2. **ARIMA (Модели авторегрессии–проинтегрированного скользящего среднего)** — вычисляет оценки параметров для сезонных и несезонных моделей, а также строит доверительные интервалы для прогноза. Процедура допускает пропущенные значения во временных рядах и выполняет анализ интервенций;

3. **EXSMOOTH (Экспоненциальное сглаживание)** — включает широкий круг методов экспоненциального сглаживания для сезонных и несезонных рядов с трендом;

4. **SEASON (Сезонные составляющие)** — оценивает мультипликативные или аддитивные сезонные составляющие для сезонных временных рядов;

5. **SPECTRA (Спектральный анализ)** — производит разложение временного ряда на гармонические составляющие. Вычисляет и выводит на график одномерную и двумерную периодограмму и оценку спектральной плотности. Позволяет использовать различные спектральные окна;

6. **AREG (Авторегрессионный анализ)** — оценивает регрессионную модель, когда ошибки близких по времени значений ряда коррелируют между собой;

7. **X11 ARIMA** — оценивает сезонные факторы для процессов типа авторегрессии–проинтегрированного скользящего среднего.

Пакет также выполняет широкий круг других процедур, например, генерацию временных рядов, вычисление оценок автокорреляционной и частной автокорреляционной функции, построение различных типов графиков временных рядов и т.д.

Командный макроязык и система меню. Прежде чем начать разбор примеров в пакете SPSS, сделаем одно важное замечание. Пакет SPSS обладает развитым командным макроязыком, позволяющим создавать командные файлы, полностью описывающие все этапы анализа. Только в последних Windows-версиях пакета появилась возможность проводить почти все процедуры ввода, редактирования и анализа данных в режиме меню-ориентированного интерфейса с диалоговыми окнами. Мы ограничим свой рассказ, ориентированный на начинающих пользователей пакета, только работой с этим интерфейсом. Заодно будет проиллюстрирован довольно типичный Windows-интерфейс современных статистических пакетов. Однако, работая в SPSS, следует помнить, что при решении задачи с использованием меню-ориентированного интерфейса одновременно происходит создание командного файла решаемой задачи. Одно из удобств и достоинств командного языка заключается в том, что при решении однотипных задач нет необходимости каждый раз заполнять поля ввода и настраивать режимы работы процедур. Можно просто запускать однажды сформированный командный файл. При этом можно практически ничего не знать о самом командном языке SPSS.

13.2.2. Подбор тренда и прогнозирование

Рассмотрим эти задачи на следующем примере.

Пример 13.1к. Для данных урожайности зерновых культур в СССР подобрать модель тренда с помощью процедур регрессионного анализа и построить на базе подобранной модели прогноз на несколько лет вперед.

Подготовка данных. Пусть данные таблицы 1.2 находятся в текстовом (ASCII) файле `zero.txt` в виде двух столбцов, первый из которых содержит значение года, а второй — значение урожайности. Для загрузки этих данных в пакет SPSS выберем в меню пакета пункт FILE, а в нем подпункт **Read ASCII Data....** На экран будет выведен запрос открытия файла, его вид — такой же, как в большинстве Windows-программ. Только в нижней части запроса имеется переключатель **File Format** (Формат файла), позволяющий выбирать между фиксированным (**Fixed**) и свободным (**Freefield**) форматами файла. Установим значение этого переключателя **Fixed**, выбрав фиксированный формат файла. Этот формат предполагает, что значения каждой переменной в строках файла записаны в тех же столбцах, что и в первой строке файла. Затем щелкнем мышью кнопку запроса **Define** (Определить) и перейдем к определению формата записи переменных в файле. На экране откроется диалоговое окно

File: C:\MAKAR\BOOK\DATA\ZERNO.DAT

Name:

Record:

Start Column:

End Column:

Data Type
123=123, 1.23=1.23

Value Assigned to Blanks for Numeric Variables
☒ System missing ☐ Value:

☒ Display summary table ☒ Display warning message for undefined data

Defined Variables:
1 1 - 5 date

Рис. 13.1. Окно определения переменных фиксированного формата процедуры загрузки текстовых файлов в SPSS

Define Fixed Variables (Определить переменные фиксированного формата), см. рис. 13.1.

В этом окне необходимо задать имена переменных. В нашем случае мы создали переменные `date` и `zerno`, а также указали начальную и конечную колонки столбца, в котором лежит каждая переменная, и формат этой переменной. Завершив описание переменных, щелкнем мышью кнопку окна **OK**. Произойдет загрузка данных в SPSS и они отразятся в электронной таблице пакета (рис. 13.2).

1 zerno	5.0
1945	5.6
1946	4.6
1947	7.3
1948	6.7
1949	6.9
1950	7.9

Рис. 13.2. Таблица SPSS с данными об урожайности зерновых

Комментарий. При загрузке ASCII-файлов могут возникнуть две проблемы. Первая — несовпадение разделителя целой и дробной части числа в исходном файле и в установке в Windows. (Обычно для этих целей используется либо точка либо запятая.) При этом не происходит корректной загрузки данных в SPSS. Для исправления ситуации обратитесь в пункт **Стандарты** (Windows International Settings) Панели Управления (Control Panel) и поменяйте тип десятичного

разделителя в пункте **Формат чисел**. Вторая проблема — частичное (округленное) отображение чисел в электронной таблице. Для исправления этой ситуации обратитесь в пункт меню **SPSS Data Define Variable Type** и увеличьте в выведенном окне значение поля **Decimal Places**. Это значение задает число позиций, отведенных для десятичной части числа в электронной таблице.

Построение графика. Анализ временного ряда начнем с построения графика. Возможности SPSS по построению и оформлению графиков очень широки, их описание занимает в документации более 250 страниц. Мы не будем вдаваться в детали оформления графиков, а будем излагать лишь общий порядок действий и показывать полученные результаты.

Для построения графика временного ряда в пункте меню **Graphs** (Диаграммы) можно выбрать один из двух возможных типов процедур: **Simple Scatterplot** (Простой график рассеивания) или **Sequence** (График последовательности). В этой задаче будет рассмотрена работа процедуры **Simple Scatterplot**. Ее диалоговое окно приведено на рис. 13.3.

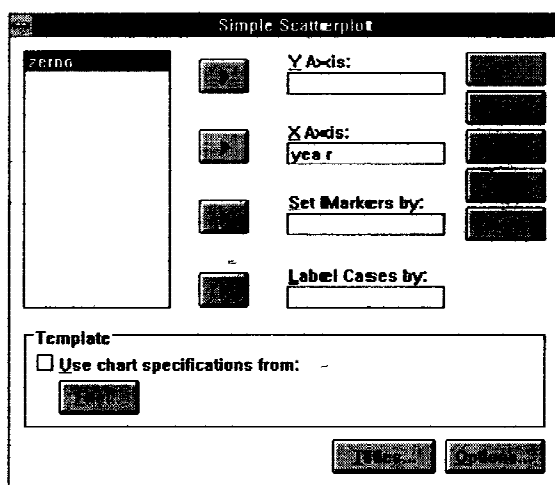


Рис. 13.3. Диалоговое окно процедуры Simple Scatterplot в SPSS

Выделяя щелчком мыши требуемые переменные и нажимая соответствующие кнопки запроса , присвоим значения переменных **date** и **zero** осям X и Y соответственно. Подобный простой способ выбора переменных используется практически во всех процедурах SPSS и в ряде других статистических пакетов под Windows (STADIA, STATGRAPHICS). На рис. 13.4 изображен полученный результат (после небольшого дополнительного оформления).



Рис. 13.4: График урожайности зерновых в SPSS

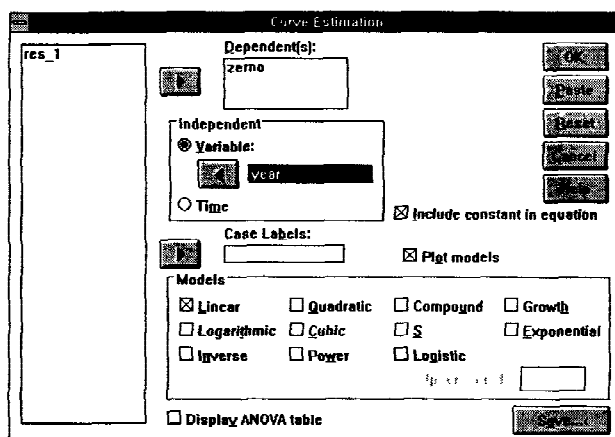


Рис. 13.5. Диалоговое окно процедуры Curve Estimation в SPSS

Выбор процедуры. На графике рис. 13.4 видно, что анализируемые данные содержат линейный тренд. Для его идентификации следует выбрать в меню пакета пункт **Statistics** (Статистика), и далее в открывшемся подменю — пункт **Regression** (Регрессия). Здесь в еще одном подменю можно выбрать один из двух методов: **Linear Regression** (линейная регрессия) или **Curve Estimation** (оценка кривой).

Процедура **Linear Regression** предоставляет широкие возможности при анализе адекватности классической модели простой и множественной линейной регрессии, включая выделение возможных «выбросов», проверку нормальности и некоррелированности остатков. А процедура **Curve Estimation** больше нацелена на выделение различных кривых трендов. Поэтому мы продолжим свой анализ с помощью процедуры **Curve Estimation**, диалоговое окно которой приведено на рис. 13.5.

- **Задание параметров процедуры.** В окне рис. 13.5 следует перенести переменную *zero* в поле *Dependent(s)* (зависимая переменная). Для этого надо выделить ее щелчком мыши и нажать кнопку . Аналогичным образом в поле *Independent* (независимая переменная) надо поместить переменную *date*. В прямоугольнике *Models* (модели) выберем линейную модель *Linear*, а также укажем включение константы в эту модель, установив флажок *Include constant in equation*. Затем нажмем кнопку окна .

Замечание. Поле независимой переменной можно оставить незаполненным, поставив переключатель типа независимой переменной в положение *Time* — в этом случае зависимая переменная будет трактоваться как временной ряд.

Модели тренда. Дадим формулы моделей тренда, приведенных в прямоугольнике *Models* на рис. 13.5. Пусть *x* — независимая переменная или время, *b_i* и *u* — константы (параметры моделей). Тогда формулы моделей можно записать так:

Linear (линейная): $y = b_0 + b_1x$;

Logarithmic (логарифмическая): $y = b_0 + b_1 \ln(x)$;

Inverse (обратная): $y = b_0 + (b_1/x)$;

Quadratic (квадратичная): $y = b_0 + b_1x + b_2x^2$;

Cubic (кубическая): $y = b_0 + b_1x + b_2x^2 + b_3x^3$;

Power (степенная): $y = b_0 \cdot (x^{b_1})$ или $\ln(y) = \ln(b_0) + b_1 \ln(x)$;

Compound (показательная): $y = b_0(b_1)^x$ или $\ln(y) = \ln(b_0) \cdot [\ln(b_1)] \cdot x$;

S (S-образная): $y = e^{(b_0+b_1/x)}$ или $\ln(y) = b_0 + b_1/x$;

Logistic (логистическая): $y = 1/(1/u + b_0 \cdot (b_1^x))$ или $\ln(1/y - 1/u) = \ln(b_0) + [\ln(b_1)] \cdot x$;

Growth (роста): $y = e^{(b_0+b_1x)}$ или $\ln(y) = b_0 + b_1x$;

Exponential (экспоненциальная): $y = b_0 \cdot (e^{b_1x})$ или $\ln(y) = \ln(b_0) + b_1 \cdot x$.

Замечание. Кнопка  диалогового окна (рис. 13.5) позволяет сохранить в виде отдельных переменных значения подобранной модели *Predicted values*, остатки *Residuals* и доверительные интервалы *Prediction intervals*, которые будут помещены в электронную таблицу пакета.

Результаты. После выполнения процедуры *Curve Estimation* в окне *Output* вывода результатов появится ряд вычисленных статистических характеристик, включая коэффициент корреляции *R*, коэффициент детерминации *R²*, таблицу анализа вариации, значения оценок коэффициентов модели и их статистические характеристики (рис. 13.6). Одновременно график ряда с подобранной кривой тренда будет помещен в окно *Chart Carousel* (рис. 13.7)

Результаты расчетов (см. рис. 13.6) показывают, что линейная модель тренда объясняет примерно 83% общей вариации данных, а полученные оценки коэффициентов модели значимо отличаются от нуля.

Dependent variable.. ZERNO Method.. LINEAR

Listwise Deletion of Missing Data

Multiple R .91521
R Square .83760
Adjusted R Square .83383
Standard Error 1.60940

Analysis of Variance:

	DF	Sum of Squares	Mean Square
Regression	1	574.46135	574.46135
Residuals	43	111.37777	2.59018

F = 221.78428 Signif F = .0000

----- Variables in the Equation -----

Variable	B	SE B	Beta	T	Sig T
YEAR	.275112	.018473	.915207	14.892	.0000
(Constant)	-528.949728	36.337744		-14.556	.0000

Рис. 13 6 Результаты работы процедуры Curve Estimation в окне выдачи результатов в SPSS

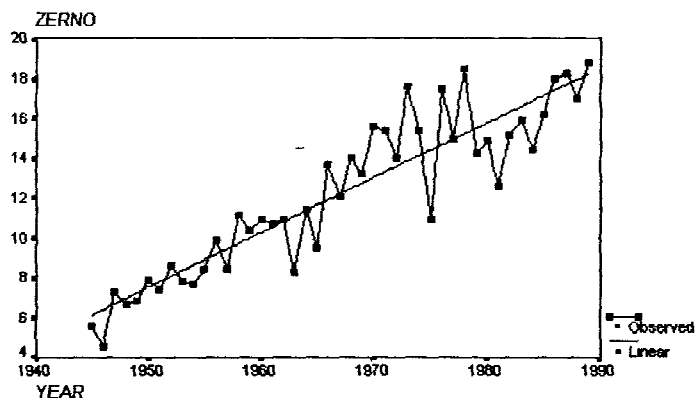


Рис. 13 7 Окно Chart Carousel с результатами работы процедуры Curve Estimation в SPSS

В частности, значение коэффициента B при переменной $year$ (то есть средний прирост урожайности за год), равен примерно 0.275 (ц/га).

Анализ остатков. Дальнейший анализ модели связан с исследованием остатков. Выясним два вопроса: можно ли считать остатки некоррелированными и насколько их распределение согласуется с нормальным.

Замечание. Учитывая небольшую длину исследуемого ряда, вряд ли можно ожидать здесь высокой точности и достоверности результатов. Однако подобный анализ позволит понять, как далеко мы могли отклониться от условий применения метода наименьших квадратов для удаления тренда, и, тем самым, насколько можно верить полученным результатам.

Проверка коррелированности остатков. Выше упоминалось, что одним из результатов работы процедуры *Curve Estimation* является создание новой переменной *err_1*, в которой хранятся остатки подобранной модели. Для выяснения коррелированности остатков вычислим оценки их автокорреляционной функции. Это можно сделать, например, вызвав процедуру *Autocorrelations* из пункта *Time Series* меню *Graphs*. Диалоговое окно процедуры приведено на рис. 13.8.

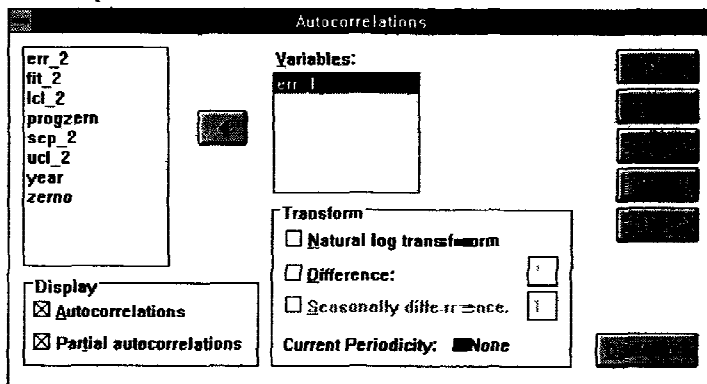


Рис 13.8. Диалоговое окно процедуры *Autocorrelations* в SPSS

В окно *Variables* (переменные) поместим переменную *err_1* со значениями остатков. С помощью кнопки *Options* зададим максимальное число лагов *Maximum Numbers of Lags* равным 10, учитывая небольшую длину изучаемого ряда. Затем нажмем кнопку *OK*. Программа выведет график автокорреляционной функции (коррелограмму), см. рис. 13.9. Из графика видно, что у нас нет основания считать остатки коррелированными, так как полученные оценки лежат внутри доверительного интервала для нулевых значений автокорреляционной функции.

Результаты расчетов процедуры помещаются в окно *Output* (см. рис. 13.10). Здесь приводятся значения оценок автокорреляционной функции, их стандартные ошибки и доверительные интервалы.

Замечание. Процедура *Autocorrelations* позволяет вычислять автокорреляционную и частную автокорреляционную функцию не только для исходного ряда, но и для прологарифмированного ряда или ряда простых и сезонных разностей. Все указанные настройки задаются в окне *Transform* (преобразования).

Еще одна проверка коррелированности остатков. Другой способ проверки некоррелированности остатков — статистика Дурбина-Уотсона. Она предназначена для проверки нулевой гипотезы H_0 о том, что все автокорреляции $r_k = 0$, против альтернативы: $r_k = r^k$ при $r \neq 0$ и $|r| < 1$. Эта альтернатива соответствует предположению, что ошибки

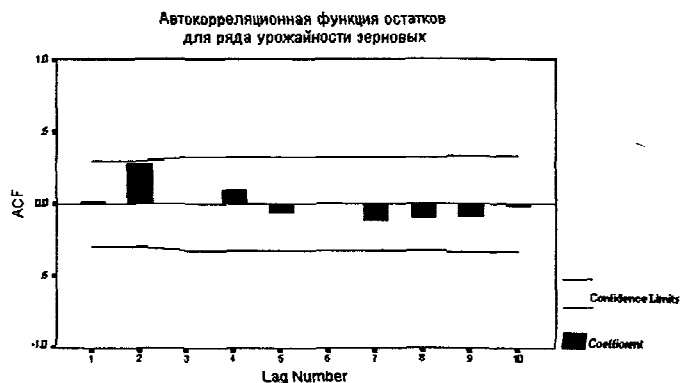


Рис 13.9 График автокорреляционной функции остатков для ряда урожайности зерновых в SPSS

Lag	Corr.	Err.	-1	-.75	-.5	-.25	0	.25	.5	.75	1	Box-Ljung	Prob.
1	.015	.149					*					.011	.916
2	.272	.149					*****					3.641	.162
3	.004	.160					*					3.642	.303
4	.096	.160					**					4.120	.390
5	-.059	.161					*					4.306	.506
6	.011	.161					*					4.312	.634
7	-.107	.162					**					4.951	.666
8	-.097	.163					**					5.494	.704
9	-.087	.164					**					5.940	.746
10	-.017	.165					*					5.958	.819

Plot Symbols: Autocorrelations * Two Standard Error Limits .
Standard errors are based on the Bartlett (MA) approximation.

Total cases: 52 Computable first lags: 44

Рис. 13.10 Результаты расчетов процедуры Autocorrelations в окне вывода результатов в SPSS

образуют процесс авторегрессии первого порядка с коэффициентом r . В SPSS эту статистику вычисляет процедура Linear Regression меню Statistics. Для остатков ряда урожайности зерновых эта статистика равна 1.96419. Таблицы процентных точек этой статистики приведены в документации модуля SPSS Trends (см. также [41]). В данном случае нулевая гипотеза должна быть отвергнута для двусторонних альтернатив, если значение статистики Дарбина-Уотсона попадает в одну из областей $(-\infty, 1.48)$ и $(2.52, \infty)$. Полученное значение статистики в эти области не попадает. Таким образом по результатам двух тестов у нас нет оснований считать остатки коррелированными.

Проверка нормальности распределения остатков. Для проверки соответствия распределения остатков нормальному распределению построим график остатков на нормальной вероятностной бумаге. Выберем в пункте меню Graphs процедуру Normal P-P (график на нор-

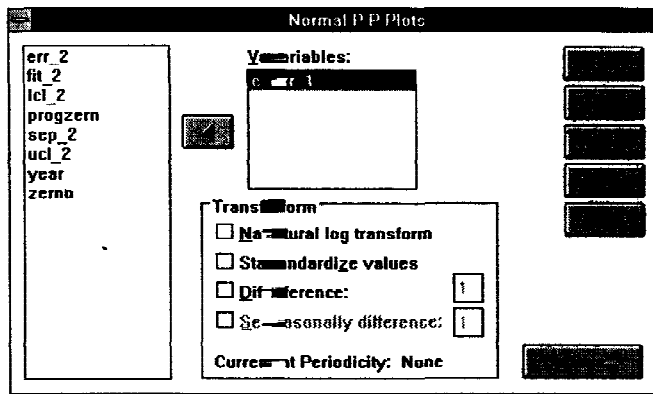


Рис. 13.11 Диалоговое окно процедуры Normal P-P в SPSS

мальной вероятностной бумаге). В диалоговом окне этой процедуры (рис. 13.11) надо указать в поле **Variables** обрабатываемую переменную.

Результаты работы процедуры приведены на рис. 13.12. Они показывают, что при хорошем согласии распределения данных с нормальным распределением «на хвостах» (зоны в районе нуля и единицы на графике), а также в центре, есть определенные отклонения в промежуточных зонах. Подобные отклонения не сильно влияют на точность полученных результатов. Поэтому результаты, полученные с помощью метода наименьших квадратов, можно считать приемлемыми.



Рис. 13.12. Результаты процедуры Normal P-P в SPSS

Замечания. 1. Процедура Normal P-P (см. рис. 13.11) предоставляет возможность проводить различные преобразования указанной переменной: переходить к логарифмической шкале, стандартизировать значения переменной, использовать простые и сезонные разностные операторы.

2. Малый объем наблюдений, естественно, не позволяет сделать обоснованных выводов о распределении остатков. Строя график остатков на нормальной вероятностной бумаге, мы прежде всего пытаемся выяснить, насколько сильно нарушается предположение о нормальности. Если эти нарушения невелики, то полученные выводы можно считать достаточно надежными. В противном случае возникает вопрос о целесообразности применения выбранного метода обработки и замене его на методы, не чувствительные к распределению данных и устойчивые к различным отклонениям.

13.2.3. Устранение сезонной компоненты

Рассмотрим эту задачу на следующем примере.

Пример 13.2к. Для данных месячных продаж шампанского построить оценки сезонной компоненты и провести сезонную коррекцию ряда.

Подготовка данных. Пусть данные о месячных продажах шампанского находятся в файле `bubbly.sav` в формате пакета SPSS в виде переменной `bubbly`. Для загрузки данных в электронную таблицу SPSS с помощью меню в пункте `File` выберите команду `Open` и в выведенном подменю укажите тип открываемого файла `Data`. В запросе открытия файла выберите файл `bubbly.sav`. Содержимое этого файла появится в электронной таблице пакета.

Для анализа сезонных временных рядов SPSS требует, кроме ввода самого ряда, дополнительного задания специальной переменной, которая указывает структуру периодичности ряда или его календарную структуру. Без этого анализ сезонных эффектов в SPSS будет недоступен. Для задания указанной переменной в меню `Data` выберите процедуру `Define Dates` (Определение календарных дат). Ее диалоговое окно приведено на рис. 13.13.

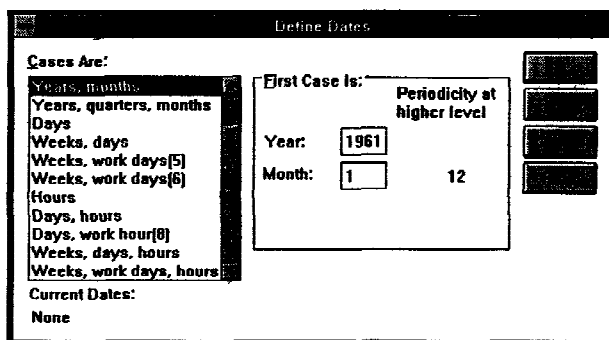


Рис. 13.13. Диалоговое окно задания типа календарных дат в SPSS

В поле `Cases Are:` (наблюдения являются) приведен обширный список различных типов календарных переменных. Из них нам следует выбрать

Years, months (годы, месяцы), так как изучаемые данные являются результатом месячных наблюдений за ряд лет. В окне *Fifth Case Is:* (первое наблюдение является) для выбранного типа переменных указывается первый год и месяц наблюдения, как это показано на рис. 13.13. В результате этой операции SPSS создаст три новых переменных *year_*, *month_* и *date_* и поместит их в электронную таблицу (см. рис. 13.14).

1 bubbly				
	2.815	1961	1	JAN 1961
	2.672	1961	2	FEB 1961
	2.755	1961	3	MAR 196
	2.721	1961	4	APR 1961
	2.946	1961	5	MAY 196
	3.036	1961	6	JUN 1961
	2.282	1961	7	JUL 1961
	2.212	1961	8	AUG 1961
	2.922	1961	9	SEP 1961
	4.301	1961	10	OCT 1961
	5.764	1961	11	NOV 1961
	7.312	1961	12	DEC 1961
	2.541	1962	1	JAN 1962
	2.475	1962	2	FEB 1962

Рис 13.14 Электронная таблица SPSS с вновь созданными календарными переменными

Выбор процедуры. Для оценки сезонных эффектов выберите в меню программы пункт *Statistics*, в появившемся подменю — пункт *Time Series* (временные ряды), и затем еще в одном появившемся подменю — пункт *Seasonal Decomposition*. Диалоговое окно выбранной нами процедуры приведено на рис. 13.15. Оно требует указания анализируемой переменной *bubbly* в окне *Variable(s)*, типа модели временного ряда: *Multiplicative* (мультипликативная) или *Additive* (аддитивная), а также указания типа взвешивания точек в скользящем среднем *Moving Average Weight*. Учитывая, что сезонная вариация данных растёт с ростом времени, выбираем мультипликативную модель временного ряда. Так как величина периода — 12 месяцев, — четная, указываем в типе взвешивания: *Endpoints weighted by .5* (веса крайних точек равны 0.5). Заполнив поля ввода, как показано на рис. 13.15, нажмем **OK** для запуска процедуры.

Результаты. Процедура сезонной декомпозиции создает четыре новые переменные:

- *saf.1* — оценки сезонных эффектов;

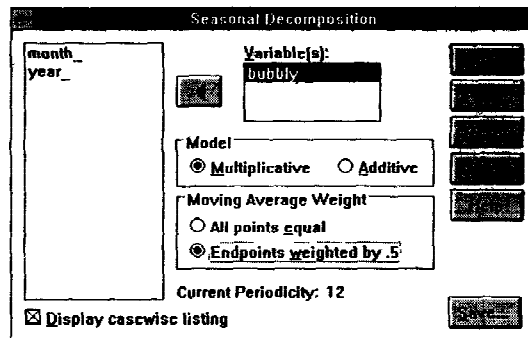


Рис. 13.15. Диалоговое окно процедуры сезонной декомпозиции в SPSS

1 bubbly							
2.815	1961	1	JAN 1961	1.01944	3.80088	.74062	3.72839
2.672	1961	2	FEB 1961	1.03829	3.77009	.70874	3.63107
2.755	1961	3	MAR 1961	.96677	3.32225	.82926	3.43644
2.721	1961	4	APR 1961	.97110	3.22230	.84443	3.31819
2.946	1961	5	MAY 1961	.97586	3.17550	.92773	3.25406
3.036	1961	6	JUN 1961	.96168	3.43123	.88481	3.56796
2.282	1961	7	JUL 1961	.81518	3.13072	.72891	3.04054
2.212	1961	8	AUG 1961	1.43449	5.98323	.36970	4.17097
2.922	1961	9	SEP 1961	.81350	3.16832	.92226	3.89470
4.301	1961	10	OCT 1961	.96689	3.55972	1.20824	3.68161
5.764	1961	11	NOV 1961	.97774	3.33071	1.73056	3.40654
7.312	1961	12	DEC 1961	1.00868	3.47404	2.10476	3.44413
2.541	1962	1	JAN 1962	.98934	3.43091	.74062	3.46787
2.475	1962	2	FEB 1962	.98246	3.49213	.70874	3.55446

Рис. 13.16. Электронная таблица SPSS с результатами работы процедуры сезонной декомпозиции

- `sas.1` — скорректированный с учетом сезонных эффектов временной ряд;
- `stc.1` — тренд и циклическая компонента скорректированного временного ряда;
- `err.1` — иррегулярная компонента скорректированного временного ряда (иными словами — случайная ошибка).

Эти переменные загружаются в электронную таблицу (рис. 13.16). Кроме того, значения этих переменных, а также вычисленные скользящие средние и промежуточные результаты расчетов сезонных индексов выводятся в окно **Output** (рис. 13.17).

Results of SEASON procedure for variable BUBBLY.
Multiplicative Model. Centered MA method. Period = 12.

DATE	BUBBLY	Moving averages	Ratios (* 100)	Seasonal factors (* 100)	Seasonally adjusted series	Smoothed trend-cycle	Irregular component
JAN 1961	2.815	.	.	74.062	3.801	3.728	1.019
FEB 1961	2.672	.	.	70.874	3.770	3.631	1.038
MAR 1961	2.755	.	.	82.926	3.322	3.436	.967
APR 1961	2.721	.	.	84.443	3.222	3.318	.971
MAY 1961	2.946	.	.	92.773	3.176	3.254	.976
JUN 1961	3.036	.	.	88.481	3.431	3.568	.962
JUL 1961	2.282	3.467	65.825	72.891	3.131	3.841	.815
AUG 1961	2.212	3.447	64.169	36.970	5.983	4.171	1.434
SEP 1961	2.922	3.450	84.685	92.226	3.168	3.895	.813
OCT 1961	4.301	3.485	123.428	120.824	3.560	3.682	.967
NOV 1961	5.764	3.542	162.737	173.056	3.331	3.407	.978
DEC 1961	7.312	3.585	203.985	210.476	3.474	3.444	1.009
JAN 1962	2.541	3.624	70.121	74.062	3.431	3.468	.989
FEB 1962	2.475	3.636	68.070	70.874	3.492	3.554	.982
MAR 1962	3.031	3.645	83.152	82.926	3.655	3.687	.991
APR 1962	3.266	3.680	88.741	84.443	3.668	3.800	1.018
MAY 1962	3.776	3.732	101.170	92.773	4.070	3.895	1.045
JUN 1962	3.230	3.821	84.541	88.481	3.650	4.003	.912

Рис. 13.17. Результаты расчетов процедуры сезонной декомпозиции

13.3. Анализ временных рядов в пакете ЭВРИСТА

13.3.1. Общие сведения о пакете

Основные возможности. Кратко перечислим основные функциональные возможности Windows-версии пакета Эвриста в области анализа временных рядов и в смежных областях. Меню статистических методов пакета включает следующие процедуры:

1. **Предварительный анализ** — вычисляет основные описательные статистики ряда, осуществляет различные преобразования ряда (удаление среднего, преобразование Бокса-Кокса, сезонные и несезонные разности), проверяет гипотезы о случайности и нормальности выборки (критерии повторных точек, хи-квадрат, Колмогорова-Смирнова) и др.;

2. **Анализ тренда** — оценивает тренд методом простой линейной или нелинейной регрессии, методом скользящего среднего, а также удаляет тренд из временного ряда и предоставляет другие полезные сервисные возможности;

3. **Прогнозирование** — строит прогнозы с помощью сглаживания методом Брауна и сезонного сглаживания Хольта-Уинтерса, а также строит прогнозы на базе подобранных моделей простой и полиномиальной регрессии и моделей типа авторегрессии-скользящего среднего;