

МИНИСТЕРСТВО ОБРАЗОВАНИЯ И НАУКИ РОССИЙСКОЙ ФЕДЕРАЦИИ
Федеральное государственное автономное образовательное учреждение
высшего образования
«Севастопольский государственный университет»

МЕТОДИЧЕСКИЕ УКАЗАНИЯ
к проведению лабораторных занятий по
дисциплине «Эконометрия»
студентов экономических специальностей
всех форм обучения

Севастополь

2019

УДК 658

Методические указания к проведению лабораторных занятий по дисциплине «Эконометрия» для студентов экономических специальностей всех форм обучения // сост. Т.А. Кокодей, И.А. Гребешкова – 2019 – 116с.

Целью методических указаний является систематизация теоретического материала и практических рекомендаций для выполнения лабораторных занятий по дисциплине «Эконометрия» для студентов всех форм обучения экономических специальностей

Утверждена на заседании кафедры Менеджмента и бизнес-аналитики от «30» июня 2019., протокол №10

СОДЕРЖАНИЕ

Лабораторная работа №1. Введение в пакет программ GRETL 1.7.1.....	4
Лабораторная работа №2. Линейный регрессионный анализ взаимосвязи статистических данных в среде GRETL 1.7.1.....	17
Лабораторная работа №3. Применение GRETL 1.7.1. при построении и анализе регрессионных моделей с гетероскедастичной случайной составляющей.....	32
Лабораторная работа №4. Реализация метода главных компонент в среде GRETL 1.7.1.....	52
Лабораторная работа №5. Анализ временных рядов в среде Gretl 1.7.1.....	68
Лабораторная работа №6 Анализ систем одновременных эконометрических уравнений в среде Gretl 1.7.1.....	97

ЛАБОРАТОРНАЯ РАБОТА №1

ВВЕДЕНИЕ В ПАКЕТ ПРОГРАММ GRETL 1.7.1

1. ЦЕЛЬ РАБОТЫ

Целью данной работы является ознакомление с функциональными возможностями программного продукта Gretl 1.7.1.

2. ТЕОРЕТИЧЕСКИЙ РАЗДЕЛ

2.1. Общие сведения о пакете GRETL

Пакет программ GRETL (GNU Regression Econometrics and Time Series Library) представляет собой инструментарий для практической реализации сложных вычислительных процедур эконометрического моделирования. В 2002 году его автор проф. Аллен Котрелл (США) включил GRETL в проект www.sourceforge.net, делая его общедоступным, бесплатным продуктом с возможностью дальнейшей доработки открытых кодов (Open Source – свободным программным обеспечением). Таким образом, данный пакет программ, статистические данные для обработки, учебное пособие и исходный код всех выпущенных версий доступны на Интернет-сайтах <http://gretl.sourceforge.net> или <http://www.kufel.torun.pl>.

Возможности программы:

1. Основные описательные статистики (среднее арифметическое, медиана, минимальное и максимальное значения, среднеквадратическое отклонение, коэффициент изменчивости (вариации), коэффициент асимметрии, коэффициент эксцесса).
2. Проверка нормальности распределения, распределение частот случайной величины, распределение плотности вероятностей, определение коэффициентов корреляции и т.д.
3. Предусматривает непосредственный доступ к статистическим таблицам. Пакет Gretl содержит встроенные статистические таблицы для следующих распределений: нормального, t-распределения Стюдента, F-распределения Фишера, хи-квадрат, Пуассона, биномиального и распределения Дарбина-Уотсона. Существует возможность вычисления критических значений, p-value.
4. Анализ временных рядов (набор методов оценивания обобщённым МНК, модели ARMAX и GARCH, система уравнений авторегрессии (VAR), проверка коинтеграции; построение линии тренда, коррелограммы, периодограммы; проверка единичных корней, моделирование типа ARIMA, а также процедуры десезонализации X-12-ARIMA и TRAMO).
5. Регрессионный анализ (одношаговый метод наименьших квадратов (1МНК), взвешенный МНК, двухшаговый МНК - оценка систем одновременных уравнений, методы оценивания логитовых, пробитовых и тобитовых моделей и нелинейных моделей, и т.д.)
6. Метод главных компонент.

7. Экспорт и импорт Gretl- Microsoft Excel и текстовые редакторы (Notepad и т.д).

8. Построение графиков и др.

Запуск программы осуществляется через **Пуск-Программы-Gretl-Gretl** или двойным щелчком мыши по иконке Gretl на рабочем столе.

2.2. Стартовый экран Gretl

Стартовый экран пакета программ GRETЛ (рисунок 1) подразделяется на три части:

- **Меню**, из которого реализуется набор функций. Меню функций состоит из следующих разделов: **File** (файл), **Tools** (инструменты), **Data** (данные), **View** (вид), **Add** (добавить), **Sample** (выборка), **Variable** (переменная), **Model** (модель), **Help** (помощь). Каждый раздел содержит группу программных функций.

- **Список переменных (процессов)**, который содержит перечень названий и описаний переменных открытого набора данных.

- **Набор иконок** (пронумерованный от 1 до 10), обеспечивает быстрый доступ

к выбранным программным функциям. Набор иконок №1-10, рисунок 1. обеспечивает быстрый доступ к некоторым программным функциям:

1. Открывает окно системного калькулятора.
2. Открывает новое окно для скриптов GRETЛ.
3. Открывает окно инструкций GRETЛ.
4. Открывает окно иконок.
5. Обращается к сайту пакета программ GRETЛ.
6. Открывает окно «Руководство» в pdf формате.
7. Открывает окно помощи.
8. Открывает окно определения графика разброса точек.
9. Открывает окно спецификации модели для оценивания с применением МНК.
10. Открывает окно с примерами – базы фактических данных.

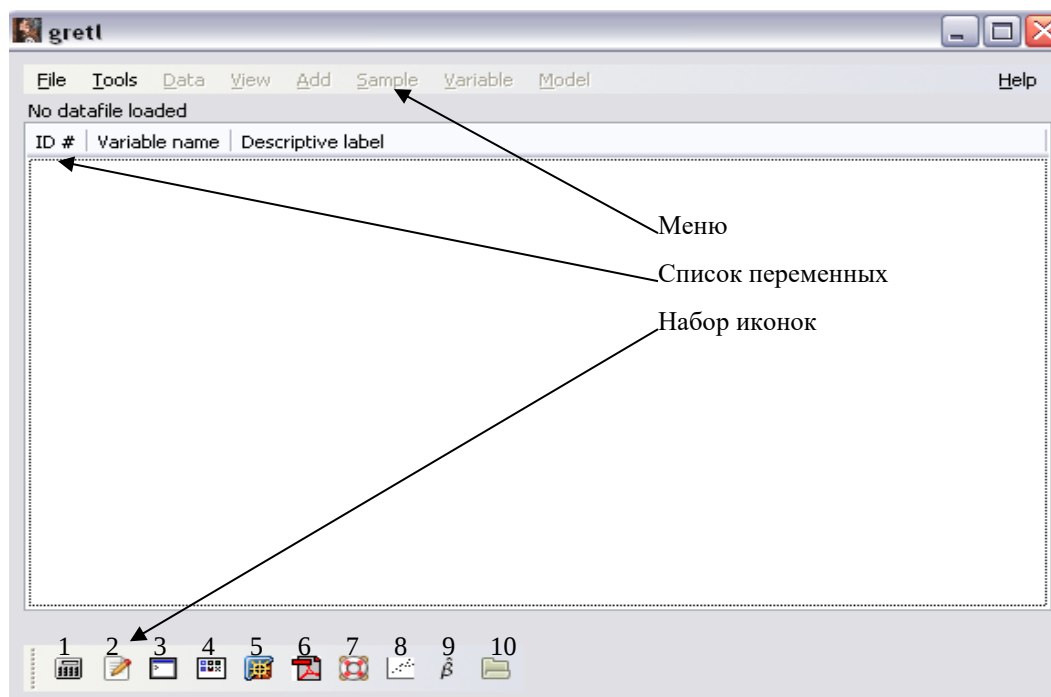


Рисунок 1 - Стартовый экран GRETl

2.3. Построение набора статистических данных

В начале работы с пакетом GRETl необходимо создать или открыть набор статистических данных. Меню File содержит команды по созданию, открытию и сохранению файлов с данными и командами GRETl.

Каждый набор данных должен иметь один из трёх типов:

1. **Срезы данных (cross-sectional)** – это неупорядоченный набор данных, например, данные, относящиеся к одному моменту времени и дающие нечто вроде поперечного среза: данные опроса семей об их уровнях дохода на определённый момент времени или данные о курсе доллара в различных городах на какую-то фиксированную дату (таблица 1);

2. **Временные ряды (time series)** – это наблюдения некоторых показателей с фиксацией периодичности (год, месяц), например, курс доллара за несколько дней (таблица 2);

3. **Панельные данные (panel)** – срезово-временные - это наблюдения за одной и той же группой объектов (фирмы, индивиды и т.п.), проведённые через определённые промежутки времени, т.е. это набор срезов (данные среза в динамике). Это могут быть данные ежегодных опросов выбранной группы семей об их уровнях дохода или ежеквартальный набор сведений (о прибыли, доходе и т.д.) об избранной группе фирм (таблица 3).

Таблица 1 - Данные типа срез данных (cross-sectional)

Курс доллара на 9.02.2008 по регионам	
Севастополь	1
Ялта	6
Одесса	7
Киев	23
Львов	45

Таблица 2 - Данные типа временной ряд (time series)

Курс доллара в г. Киев	
07.02.2008	4,9
08.02.2008	5
09.02.2008	5,06

Таблица 3 - Панельные данные (panel)

Панельные данные			
	07.02.2008	08.02.2008	09.02.2008
Киев	4,9	5	5,06
Севастополь	5	5,1	5,08
Одесса	4,8	5,02	5,03

2.3.1. Ручной ввод информации с клавиатуры

Для **создания нового набора данных** необходимо выбрать пункт **New Data Set** в меню **File** (рисунок 1), указав число наблюдений создаваемого ряда (number of observations), один из перечисленных выше типов данных, и название первой создаваемой переменной набора, рисунок 2.

Сохранение данных в формате Gretl *.gdt осуществляется при помощи **File\Save Data**, а закрытие набора данных при помощи **File\Clear Dataset**.

Пример 1. В меню **File** выберем пункт **New Data Set** и введём непосредственно с клавиатуры:

Шаг 1. Число наблюдений (number of observations): «5» и нажмём кнопку «ок» (рисунок 2).

Шаг 2. Тип данных (structure of dataset): выберем срез данных (cross-sectional) и кнопку «forward».

Шаг 3. Нажмём кнопки «ок» для подтверждения структуры данных.

Шаг 4. Нажмём кнопку «yes» для подтверждения ввода данных в Gretl.

Шаг 5. Введём название переменной «А» в открывшемся окне и нажмём кнопку «ок».

Шаг 6. В окне редактирования данных введём с клавиатуры пять значений наблюдений 1, 6, 7, 23, 45 (рисунок 3), нажмём кнопки Apply и Close, чтобы записать данную переменную в список стартового экрана (в создаваемый набор данных).

При помощи **File\Save Data** сохраним данный набор данных как файл example1.gdt. Закроем созданный набор данных при помощи **File\Clear Dataset**.

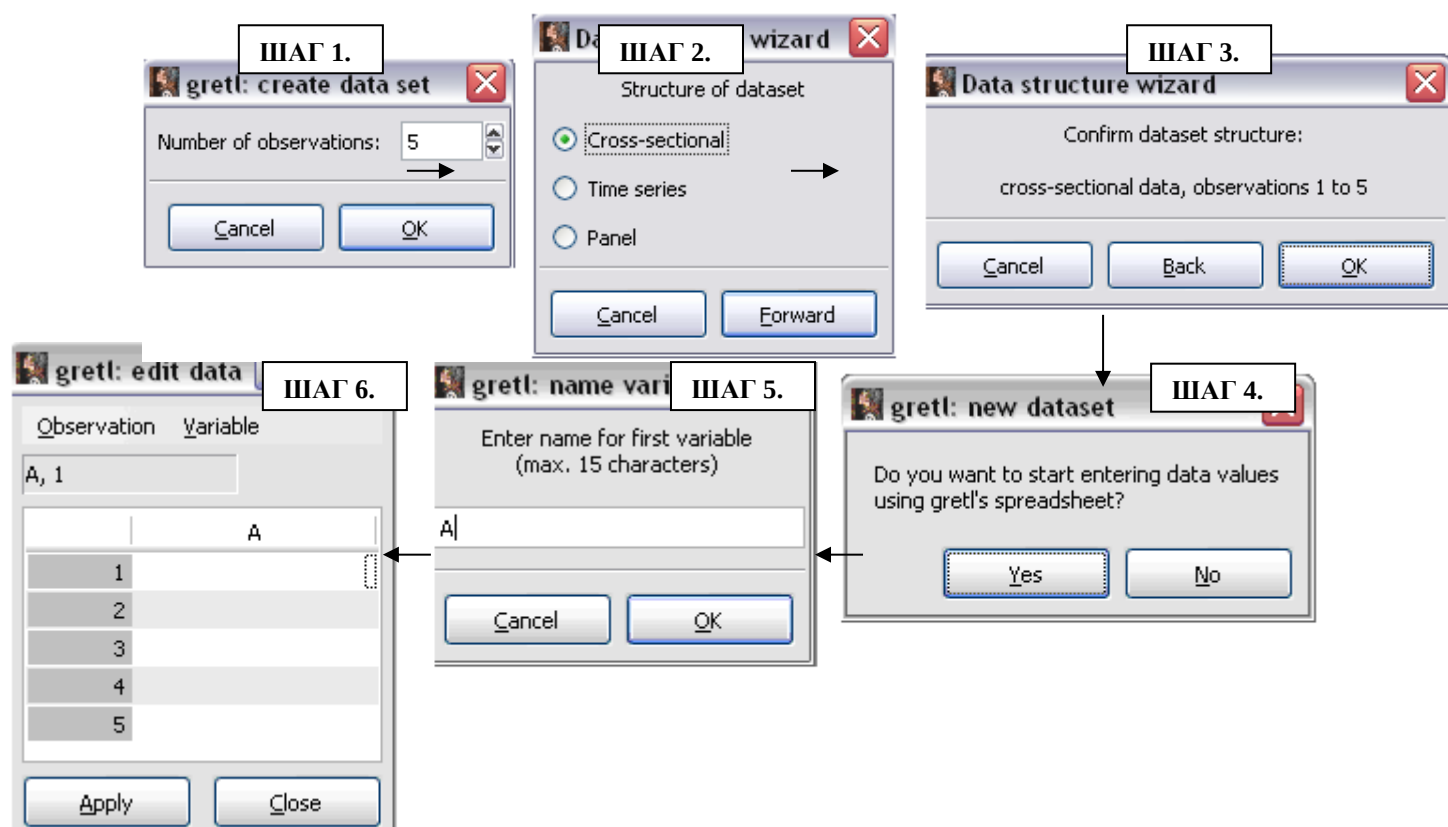


Рисунок 2 - Этапы построения набора данных «срез данных»

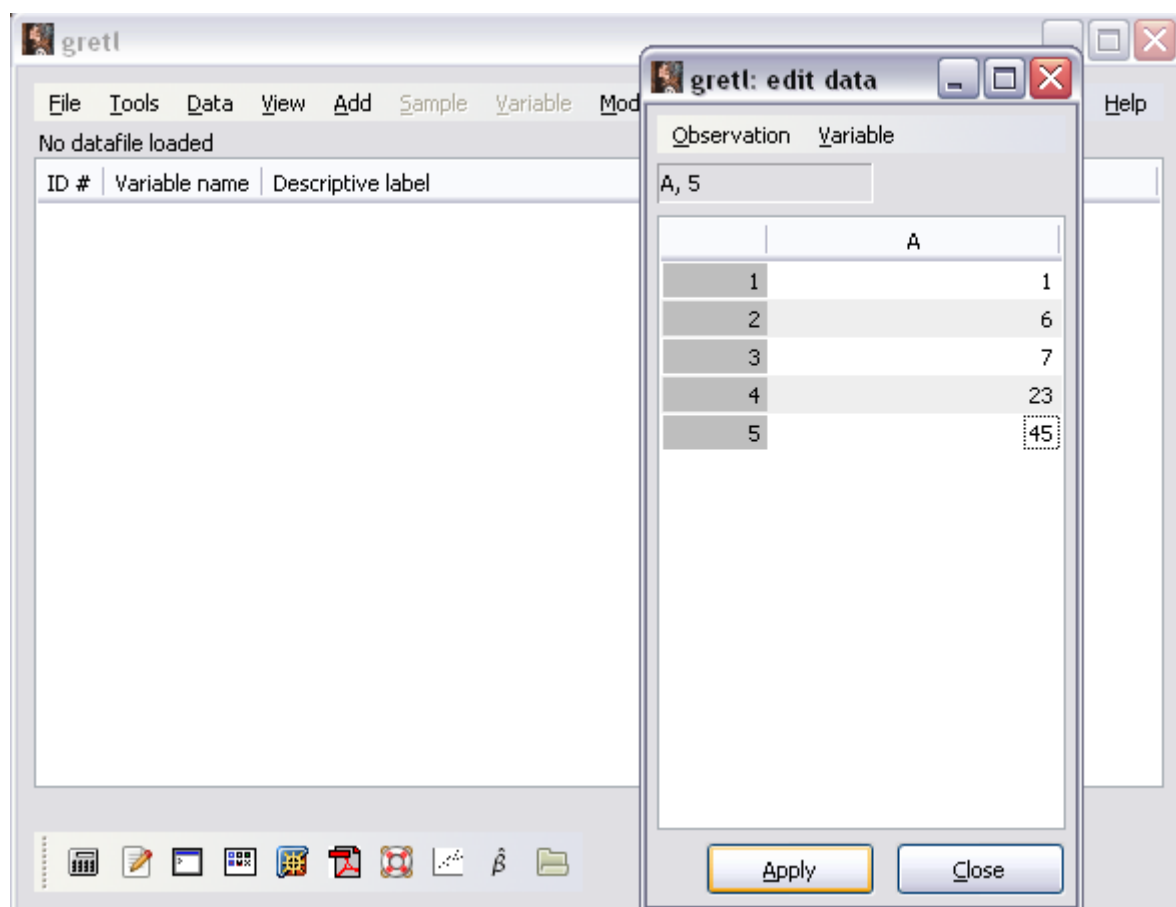


Рисунок 3 - Ввод с клавиатуры данных типа «срез данных»

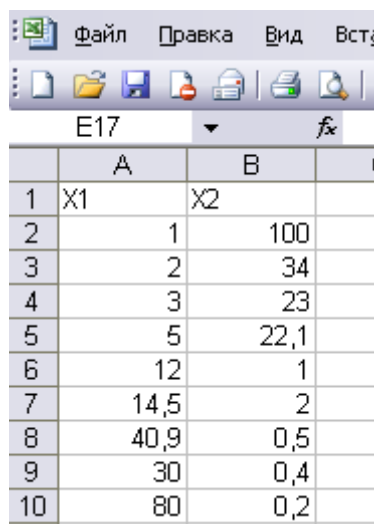
2.3.2. Импорт данных

Также возможен импорт заранее подготовленной таблицы EXCEL при помощи команды **File\Open Data\Import\Excel**, а также некоторых других типов файлов.

Пример 2.

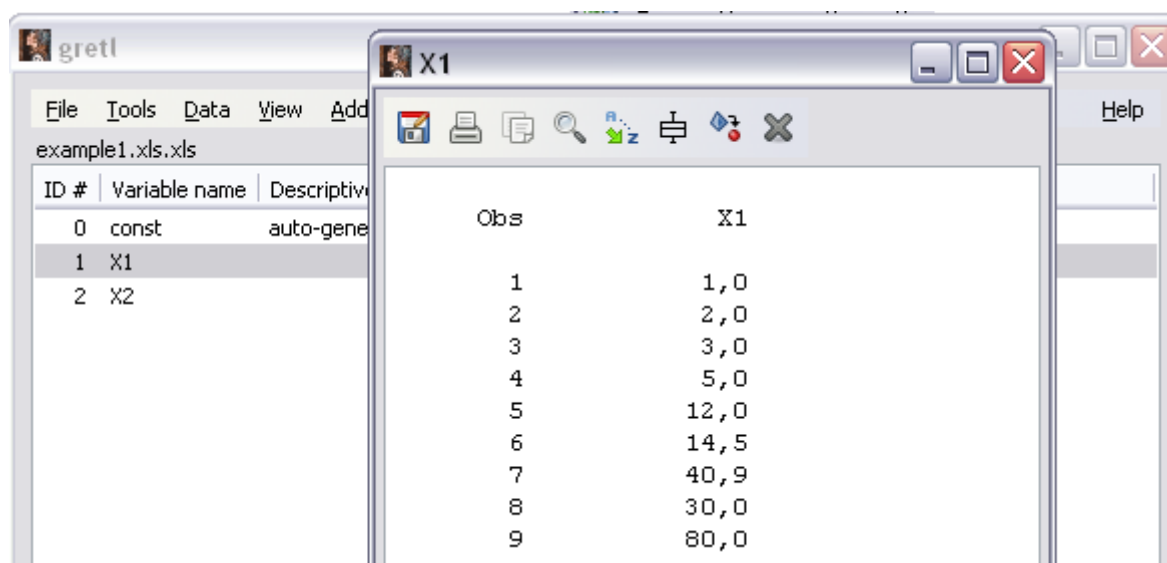
Осуществим импорт данных из файла Excel в пакет GRETЛ:

1. Создадим и поместим на рабочий стол файл example2.xls (рисунок 4).
2. Обратимся к команде Gretl - File\Open Data\Import\Excel, укажем номер строки и столбца начала таблицы Excel и выберем тип данных, например, срез данных (cross- sectional).
3. Дважды щёлкнем левой кнопкой мыши по названию появившихся в списке переменных стартового экрана Gretl переменных X1 и X2 для просмотра их значений (рисунок 5).
4. Сохраним набор данных File\Save Data как файл example2.gdt на рабочем столе.



	A	B
1	X1	X2
2	1	100
3	2	34
4	3	23
5	5	22,1
6	12	1
7	14,5	2
8	40,9	0,5
9	30	0,4
10	80	0,2

Рисунок 4 - Файл с данными example2.xls



ID #	Variable name	Descriptive
0	const	auto-generated
1	X1	
2	X2	

Obs	X1
1	1,0
2	2,0
3	3,0
4	5,0
5	12,0
6	14,5
7	40,9
8	30,0
9	80,0

Рисунок 5 - Вывод значений импортированной переменной X1

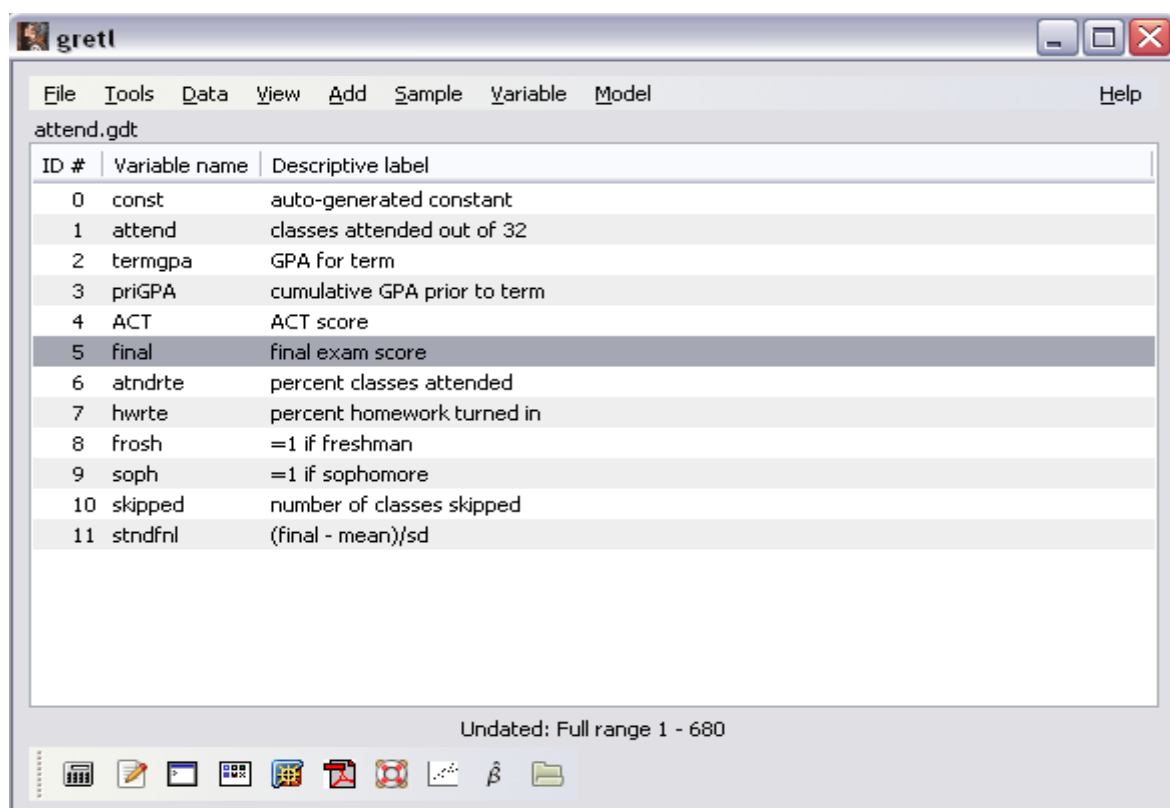
2.4. Открытие встроенного или ранее созданного набора данных

В GRETЛ существуют встроенные примеры наборов данных (*.gdt), созданные разработчиками и преподавателями. Для их открытия необходимо обратиться к команде **File\Open Data\Sample File** и выбрать (на соответствующей закладке, например Wooldridge) имя открываемого файла, например, attend.gdt двойным щелчком мыши (рисунок 6).

Для просмотра значений отдельной переменной обратимся к команде: **Variable\Display Values** (или дважды щёлчком мышью по названию переменной).

Для просмотра всего набора данных обратимся к команде: **View\Icon View\Data Set**

Открытый набор данных, рисунок 6, состоит из 11 переменных, атрибуты которых -номера (ID#), названия (Variable name) и описания\формулы (Descriptive label)- представлены в списке переменных стартового экрана.



ID #	Variable name	Descriptive label
0	const	auto-generated constant
1	attend	classes attended out of 32
2	termgpa	GPA for term
3	priGPA	cumulative GPA prior to term
4	ACT	ACT score
5	final	final exam score
6	atndrte	percent classes attended
7	hwrtte	percent homework turned in
8	frosh	=1 if freshman
9	soph	=1 if sophomore
10	skipped	number of classes skipped
11	stndfml	(final - mean)/sd

Рисунок 6 - Набор данных Class Attendance Rates and Grades (Посещаемость занятий и оценки) с закладки Wooldridge

Аналогичным образом можно открыть созданные ранее (в Примере 1 и Примере 2) на рабочем столе файлы example1.gdt и example2.gdt при помощи **File\Open Data\User File**, выбрав их в открывшемся окне.

Просмотреть общую информацию о наборе данных можно выбрав **Data\Print Description**

2.5. Редактирование набора статистических данных

Возможны следующие операции над переменными открытого набора данных:

1. Добавление переменной вручную: **Variable\Define new variable** или из Excel файла: **File\Append Data\Excel**
2. Удаление переменной: нажатие кнопки **DEL** на клавиатуре.
3. Редактирование значений переменной: **Data>Edit Values**.
4. Добавление наблюдений: **Data>Add Observations**.
5. Изменение атрибутов переменной: **Variable>Edit Attributes**.
6. Просмотр и редактирование всего набора данных: **View\Icon View\Data Set**.
7. Удаление наблюдений с пропущенными значениями: **Sample\Drop all obs with missing values**.

Функция **Variable\Define New Variable** позволяет добавить ещё одну переменную, а функция **Data>Add Observations** – добавить определённое число наблюдений к существующему выбранному ряду. Двойной щелчок мыши по названию переменной позволяет просмотреть ряд её значений, а функция **Data>Edit Values** – редактировать данные значения. Удаление переменной из списка осуществляется нажатием кнопки **del** на клавиатуре. Чтобы изменить атрибуты переменной необходимо щелчком мыши выбрать её название в списке и вызвать функцию **Variable>Edit Attributes**, затем в открывшемся диалоговом окне ввести новое имя переменной (**Name of variable**) и её формулу или текстовое описание (**Description**).

Быстрый доступ к данным функциям возможен из контекстного меню, вызванного нажатием правой кнопкой мыши на выбранной переменной в списке стартового экрана.

Пример 3. Редактирование набора данных example1.gdt:

1. Откроем ранее созданный (Примере1) набор данных example1.gdt: **File\Open Data\User File** (рисунок 7)
2. Изменим название переменной на X1 и введём её описание «Объём продаж»: **Variable>Edit Attributes** (рисунок 7)

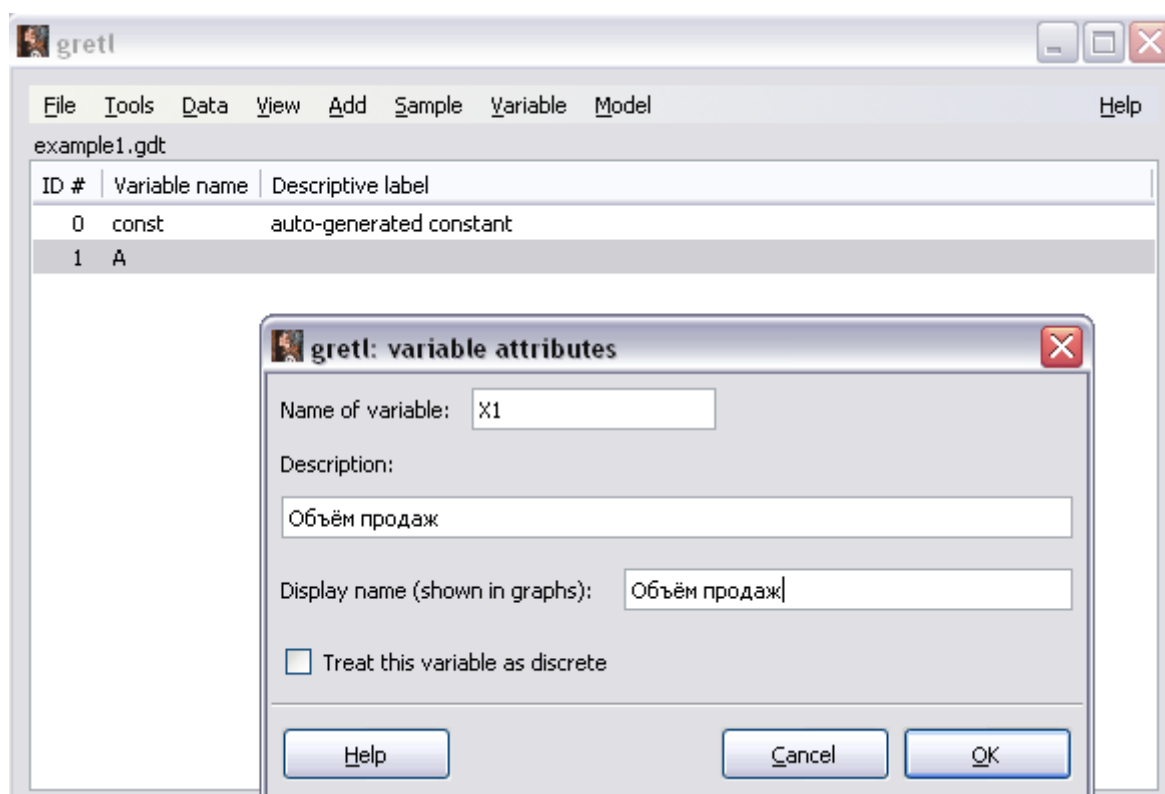


Рисунок 7 - Изменение атрибутов переменной: ввод названия и описания

3. Добавим в набор новую переменную Y. Для этого выберем команду **Define new variable** в меню **Variable** и в открывшемся окне редактирования введём её значения:

-4,2 23,8 34,2 748,992 1615 (рисунок 8)

4. Щелчком мыши выберем переменную X1 из списка (рисунок 8) и добавим одно наблюдение “100” к ряду её значений: **Data\Add Observations** (введём 1) и **Data>Edit Values** (в открывшемся окне “Edit data” введём шестое наблюдение “100”).

5. В этом же окне «Edit data» вручную изменим значение четвёртого наблюдения на 30,7 (рисунок 8), нажмём кнопки «apply» и «close» для завершения редактирования значений переменной X1.

6. Аналогичным образом выберем переменную Y, обратимся к команде **Data>Edit Values** и введём ещё одно значение ряда 7995.

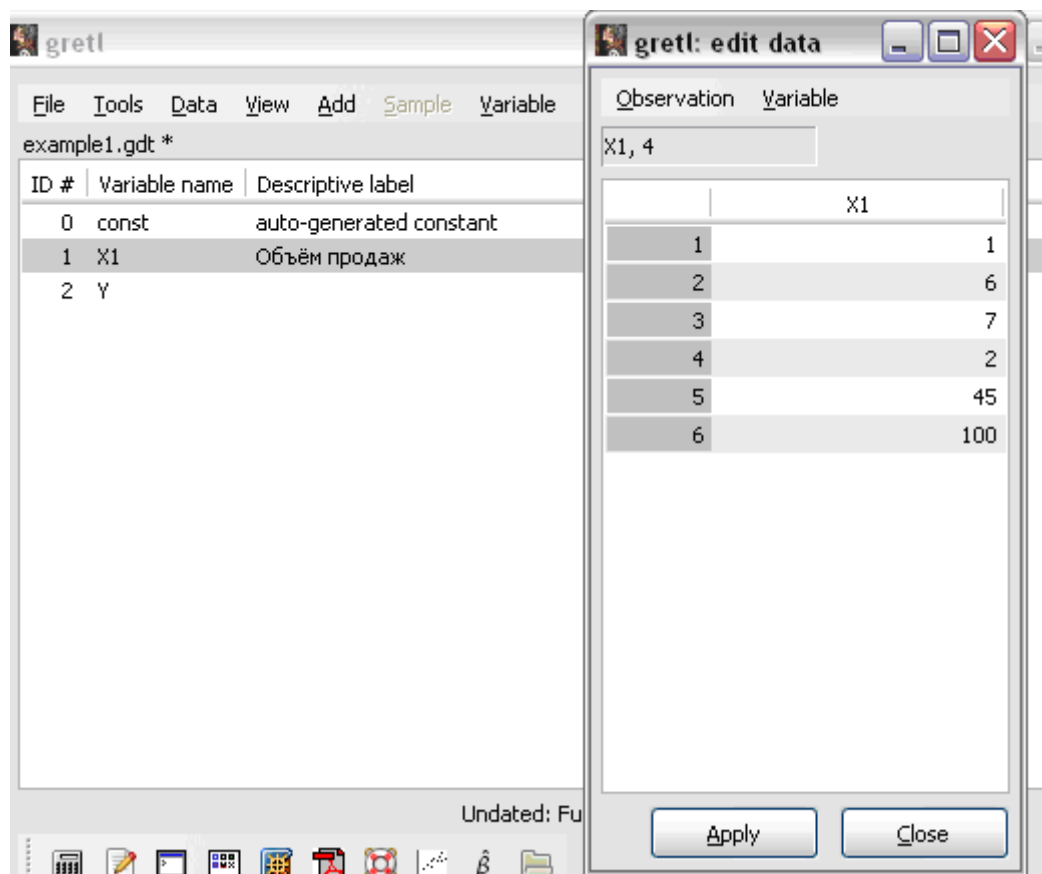
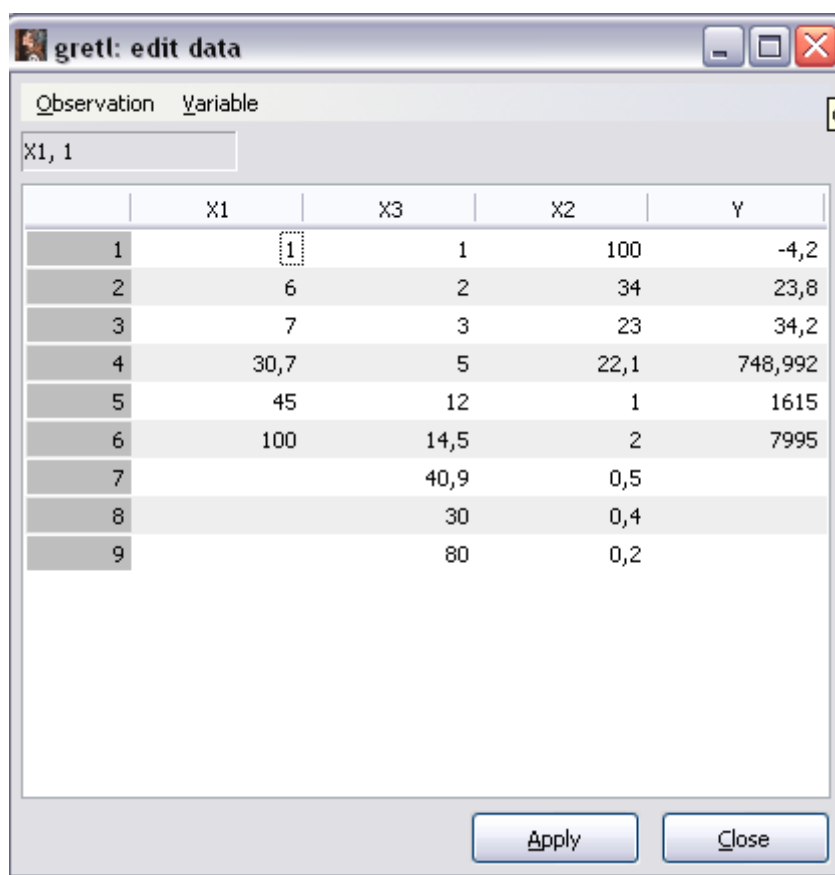


Рисунок 8 - Редактирование значений переменной X1

8. В ранее созданном (в Примере 2.) на рабочем столе файле example2.xls изменим название переменной X1 на X3. Затем добавим к редактируемому в данном примере набору данных example1.gdt переменные из файла example2.xls: **File\Append Data\Excel** (выбрать название файла). В результате в список переменных стартового экрана будут добавлены переменные X2 и X3.

9. Откроем весь набор данных для редактирования: **View\Icon View\Data Set** (рисунок 9).

10. Удалим наблюдения №7-9 с пропущенными значениями переменных X1 и Y, сократив выборку до шести первых наблюдений: **Sample\Drop all obs with missing values**. Нажмём кнопки «apply» и «close». Сохраним набор данных **File\Save Data**, ответив «по» на вопрос о восстановлении первоначального размера выборки.



gretl: edit data

Observation Variable

X1, 1

	X1	X3	X2	Y
1	1	1	100	-4,2
2	6	2	34	23,8
3	7	3	23	34,2
4	30,7	5	22,1	748,992
5	45	12	1	1615
6	100	14,5	2	7995
7		40,9	0,5	
8		30	0,4	
9		80	0,2	

Apply Close

Рисунок 9 - Окно редактирования набора данных

2.6. Экспорт данных

Экспорт данных в Excel осуществляется с использованием команды **File\Export Data\SCV**. Экспортируем в Excel полученный в Примере 3 набор данных:

1. Откроем в Gretl файл example1.gdt (File\Open Data\User File).
2. Обратимся к команде File\Export Data\SCV. В открывшемся окне поставим флажки разделителей **semicolon** и **comma (,)** для интерпретации информации в Excel как количественных данных в отдельных столбцах, рисунок 10.
3. В открывшемся окне Save Data при помощи кнопки Select перенесём все переменные из списка в левой части окна в правую часть и нажмём кнопку ОК.
4. Введём имя файла example1.csv и нажмём кнопку save. Данный файл появится на рабочем столе.

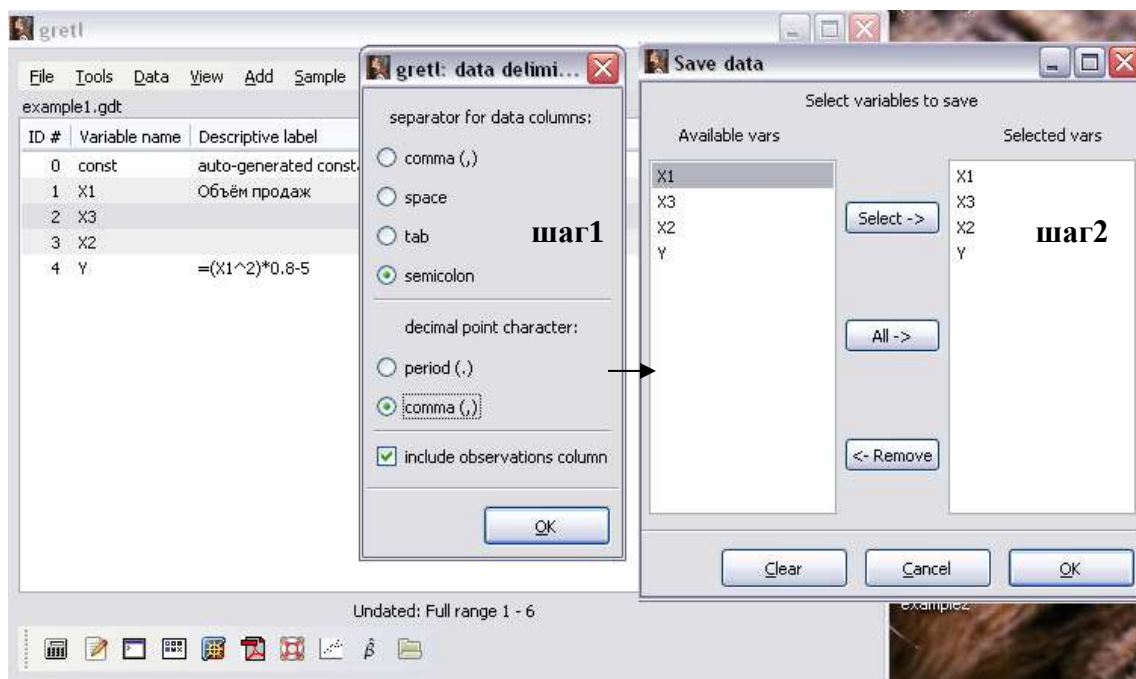


Рисунок 10 - Экспорт данных из среды GRETL в таблицу Excel

3. ПОРЯДОК ВЫПОЛНЕНИЯ ЛАБОРАТОРНОЙ РАБОТЫ

1. Выполнить примеры №1-3 данных методических указаний.
2. Выполнить упражнения №1-3 согласно варианту.
3. Подготовить отчёт по выполненной работе в электронном виде (MS-Word).

Упражнение 1. Создать набор данных из двух переменных X1 и X2 (тип данных – cross-sectional) по исходной информации, представленной в таблице 4 согласно варианту. Сохранить в файл Ex1.gdt на рабочем столе и закрыть созданный набор данных.

Таблица 4 - Значения переменной X₁ и X₂ по вариантам №№1-10

№1		№2		№3		№4		№5	
3	1	3,5	2	3,6	1	10	11	0	10
5,3	5,3	5,3	5,3	5,3	5,3	5,3	5,3	5,3	6,3
7	7	9	7,5	7,1	7	7,9	41,5	7,1	7
10,3	10,3	10,3	10,3	10,3	10,3	10,3	10,3	10,3	10,3
40,2	40	30,2	40,2	30,2	40,2	40,2	40,2	40,2	40,2
80	80	80	80	80	81	80	80	80	87
100	101	200	100	100	101	100	100	100	106
120	120	130	120	129	121	120	128	120	125
№6		№7		№8		№9		№10	
3	1	3,5	2	3,6	1	10	11	0	10
5,3	5,3	5,3	5,3	5,3	5,3	5,3	5,3	5,3	6,3
7	7	9	7,5	7,1	7	7,9	41,5	7,1	7
10,3	10,3	10,3	10,3	10,3	10,3	10,3	10,3	10,3	10,3
40,2	40	30,2	40,2	30,2	40,2	40,2	40,2	40,2	40,2
80	80	80	80	80	81	80	80	80	87
100	101	200	100	100	101	100	100	100	106
120	120	130	120	129	121	120	128	120	125

Упражнение 2.

Импортировать файл example2.xls (созданный в вышеописанном Примере 2) в среду Gretl как временной ряд с поквартальными данными, начиная с даты X.0X.2008г., где X – номер варианта.

Упражнение 3.

1. Открыть встроенный набор данных **File\Open data\Sample File\Ramanathan**файл в зависимости от варианта (рисунок 11).
2. Просмотреть общую текстовую информацию о наборе данных **Data\Print Description**.
3. Изменить первые три значения одной из переменных.
4. Сократить выборку на последние пять наблюдений.
5. Экспортировать данные в таблицу Excel.

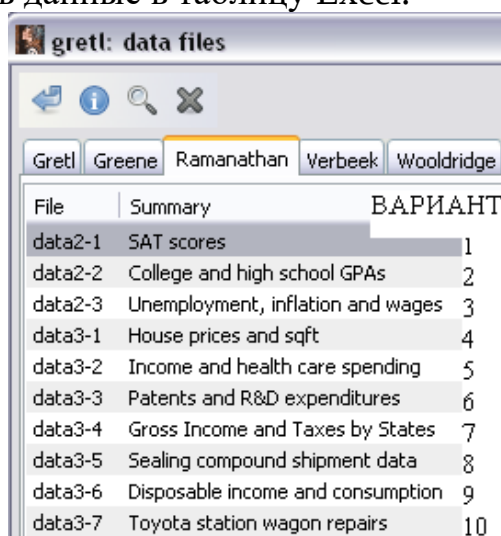


Рисунок 11 - Варианты заданий для выполнения упражнения 3

4. СОДЕРЖАНИЕ ОТЧЕТА О ВЫПОЛНЕНИИ ЛАБОРАТОРНОЙ РАБОТЫ

- 1) Название и цель работы.
- 2) Постановка задачи.
- 3) Этапы выполнения задачи в Gretl.
- 4) Выводы.

ЛАБОРАТОРНАЯ РАБОТА №2 ЛИНЕЙНЫЙ РЕГРЕССИОННЫЙ АНАЛИЗ ВЗАИМОСВЯЗИ СТАТИСТИЧЕСКИХ ДАННЫХ В СРЕДЕ GRETЛ 1.7.1.

1. ЦЕЛЬ РАБОТЫ

Целью данной работы является получение практических навыков регрессионного анализа в системе Gretl для автоматизированного поиска ранее неизвестных закономерностей в имеющихся в распоряжении менеджера данных с последующим использованием полученной информации для подготовки управленческих решений.

2. ТЕОРЕТИЧЕСКИЕ СВЕДЕНИЯ О ЛИНЕЙНОМ РЕГРЕССИОННОМ АНАЛИЗЕ

Целью регрессионного анализа является оценка функциональной зависимости $y = \hat{f}(x_1, x_2, \dots, x_n) + u$ результативного признака (y) от факторных ($x_1, x_2, \dots, x_n$). Формулы (1) и (2) представляют собой линейные модели парной и множественной регрессии соответственно.

$$y = f(x) + u = \alpha_0 + \alpha_1 \cdot x + u, \quad (1)$$

$$y = \alpha_0 + \alpha_1 \cdot x_1 + \dots + \alpha_n x_n + u, \quad (2)$$

где y — фактическое значение результативного признака;

x_i - признак-фактор;

α_i — параметр регрессионной модели;

u — случайная ошибка (остаток), характеризующая отклонения реального значения результативного признака от теоретического. Она включает влияние не учтенных в модели факторов, случайных ошибок и особенностей измерения.

Оценивание параметров линейной модели основан на обычном или одношаговом методе наименьших квадратов (МНК или OLS – Ordinary Least Squares).

Этот метод позволяет получить такие оценки параметров, при которых сумма квадратов отклонений фактических значений результативного признака (y) от расчетных (теоретических) $\tilde{y}_x$ минимальна, формула (3).

$$\sum_i (y_i - \tilde{y}_{x_i})^2 \rightarrow \min, \quad (3)$$

Статистическое моделирование связи методом линейного регрессионного анализа осуществляется в 3 этапа:

а) Оценка параметров линейной регрессионной модели методом МНК

Вектор оценок параметров модели (2) определяется выражением (4).

$$\alpha = [X^T \cdot X]^{-1} \cdot X^T \cdot Y \quad (4)$$

б) Проверка адекватности регрессионной модели (проверки значимости индивидуальных оценок коэффициентов модели с помощью t-

критерия Стьюдента и оценка значимости уравнения регрессии в целом с помощью F-критерия Фишера)

На первом шаге проверки адекватности (качества) модели оценивается существенность влияния каждой объясняющей переменной x_i , на зависимую переменную y , для этого необходимо оценить значимость полученных параметров α_i , используя t- критерий Стьюдента, формула (5). Значимость параметра определяется путём проверки нулевой гипотезы о равенстве его нулю (для выбранного уровня значимости).

$$t_{p_{\alpha_i}} = \frac{|\alpha_i|}{\sqrt{\sigma_{\alpha_i}^2}}, \quad (5)$$

где α_i - оценка i -го коэффициента модели, COEFFICIENT;

$\sigma_{\alpha_i}^2$ - оценка дисперсии параметра α_i , $\sqrt{\sigma_{\alpha_i}^2} = \text{STDERROR}$.

На втором шаге проверки адекватности модели оценивается её значимость (пригодность) в целом, используя показатели: F-критерий Фишера, формула (6), коэффициент детерминации R^2 , формула (7), (Unadjusted R^2 и Adjusted R^2), сумма квадратов остатков RSS Sum of squared residuals), стандартная ошибка регрессии (Standard error of residuals), информационные критерии (Akaike information criterion, Schwarz Bayesian criterion, Hannan-Quinn criterion).

Значимость регрессии проверяется путём проверки нулевой гипотезы о равенстве нулю всех параметров модели (для выбранного уровня значимости).

$$F_p = \frac{R^2}{1-R^2} \cdot \frac{n-k}{k-1}, \quad (6)$$

где R^2 - коэффициент детерминации - часть вариации (дисперсии) зависимой переменной y , которая объясняется уравнением регрессии, UNADJUSTED R^2 .

$$R^2 = \frac{\hat{\alpha}^T \cdot X^T \cdot Y - n \cdot \bar{Y}^2}{Y^T \cdot Y - n \cdot \bar{Y}^2}, \quad (7)$$

n - число наблюдений;

k – число коэффициентов факторов.

При анализе адекватности уравнения регрессии исследуемому процессу возможны следующие варианты:

- Построенная модель на основе ее проверки по F-критерию Фишера в целом адекватна, и все коэффициенты регрессии значимы. Такая модель может быть использована для принятия решений к осуществлению прогнозов.
- Модель по F-критерию Фишера адекватна, но часть коэффициентов регрессии незначима. В этом случае модель пригодна для принятия некоторых решений, но не для производства прогнозов.
- Модель по F-критерию Фишера адекватна, но все коэффициенты регрессии

незначимы. Поэтому модель полностью считается неадекватной. На ее основе не

принимаются решения и не осуществляются прогнозы.

с) Анализ выполнения предпосылок 1МНК (условий Гаусса—Маркова).

Регрессионный анализ линейных функций, основанный на обычном или одношаговом методе наименьших квадратов (1МНК) должен удовлетворять четырем условиям Гаусса—Маркова:

1. Математическое ожидание случайной составляющей, $M(u_i)$ в любом наблюдении должно быть равно нулю.

2. Дисперсия случайной составляющей σ_u^2 должна быть постоянна для всех наблюдений. Дисперсия случайной составляющей σ_u^2 должна быть постоянна для всех наблюдений. Если это условие не соблюдается, то имеет место гетероскедастичность. Наличие гетероскедастичности можно определить при помощи теста Уайта, который позволяет проверить значимость регрессии квадратов остатков относительно комплекса переменных модели и их квадратов. При этом формулируется нулевая гипотеза о гомоскедастичности остатков (равенстве нулю всех коэффициентов модели).

3. Отсутствие систематической связи между значениями случайной составляющей u_i в любых двух наблюдениях. Отсутствие автокорреляции остатков.

Автокорреляция остатков означает наличие корреляции между остатками текущих и предыдущих (последующих) наблюдений. Оценить эту зависимость можно вычислив коэффициент корреляции между этими остатками по формуле (8).

$$r_{u_i u_j} = \frac{\overline{u_i u_j} - \bar{u}_i \cdot \bar{u}_j}{\sigma_{u_i} \sigma_{u_j}}. \quad (8)$$

4. Случайный характер остатков. Случайная составляющая должна быть распределена независимо от переменных u_i и x_i .

3. ОПИСАНИЕ СРЕДСТВ СИСТЕМЫ GRETL ДЛЯ ВЫПОЛНЕНИЯ РЕГРЕССИОННОГО АНАЛИЗА

3.1. Оценка параметров линейной регрессионной модели методом 1МНК (OLS) и проверка адекватности модели

В пакете программ GRETL параметры эконометрической модели можно оценить с применением метода наименьших квадратов, в частности, одношагового 1МНК (Ordinary Least Squares - OLS) для оценки линейных регрессионных моделей для срезов данных (cross-sectional data type).

Окно спецификации эконометрической модели, оцениваемой с применением 1МНК, вызывается функцией меню **Model\Ordinary Least**

Squares... В этом окне Y (Dependent variable) выбирается при помощи кнопки «Choose», а объясняющие переменные (Independent variable) – при помощи кнопки “Add”.

Пример 1.

Откроем встроенный набор данных attend.gdt на закладке Wooldridge (**File\Open Data\Sample File\attend.gdt**) и обратимся к функции **Model\Ordinary Least Squares**, чтобы построить линейную регрессионную модель, отражающую зависимость переменной Final (оценка за итоговый экзамен) от переменных attend (число посещённых занятий), termGPA (средний балл за семестр), priGPA (средний балл на начало семестра), ACT (оценка по вступительному тесту ACT), hwrte (процент сданных домашних работ). В открывшемся окне спецификации модели выберем соответствующие зависимую и независимые переменные при помощи кнопок «Choose» и “Add”, затем нажмём кнопку ОК (рисунок 1).

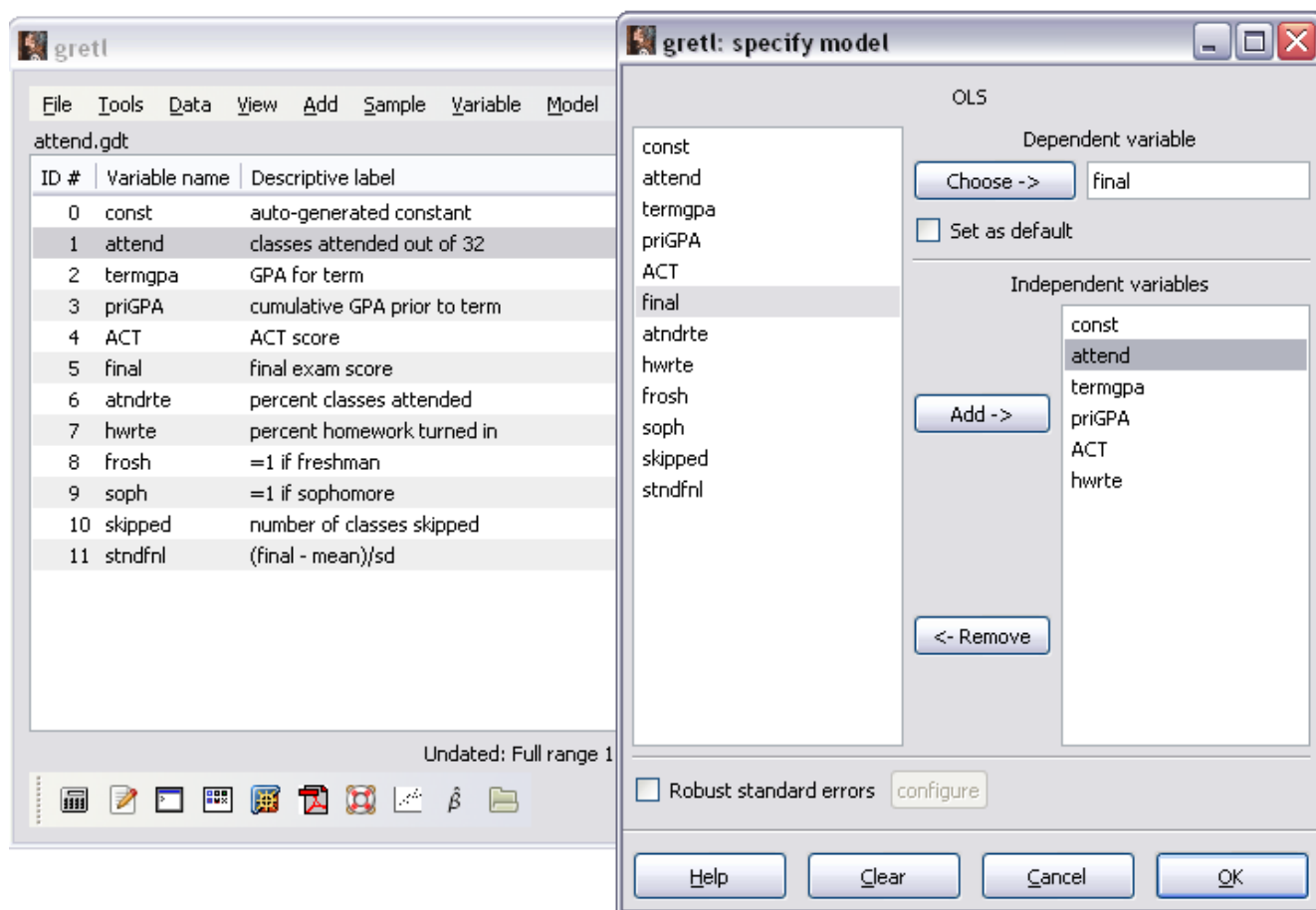


Рисунок 1 - Окно спецификации эконометрической модели, оцениваемой с применением МНК

Полученные результаты представлены на рисунке 2. По данным наблюдений attend.gdt была составлена модель (9).

$$\text{final} = 13,784 - 0,069 \cdot \text{attend} + 3,54 \cdot \text{termgpa} - 0,066 \cdot \text{priGPA} + 0,272 \cdot \text{Act} - 0,015 \cdot \text{hwrte} + u. \quad (9)$$

Рассматривая значения параметров модели α_i для данной отдельной выборки, можно отметить существенность зависимости переменной final от переменных termgra (сильная положительная связь, 3,54) и АСТ (положительная связь 0,272), остальные коэффициенты имеют значения близкие к 0 и не оказывают существенного влияния на результирующий признак. Необходимо установить насколько вероятно, что зависимость, подобная найденной, подтвердится на данных другой выборки, извлеченной из той же самой генеральной совокупности, т.е. можно ли свойства данной выборки перенести на всю генеральную совокупность.

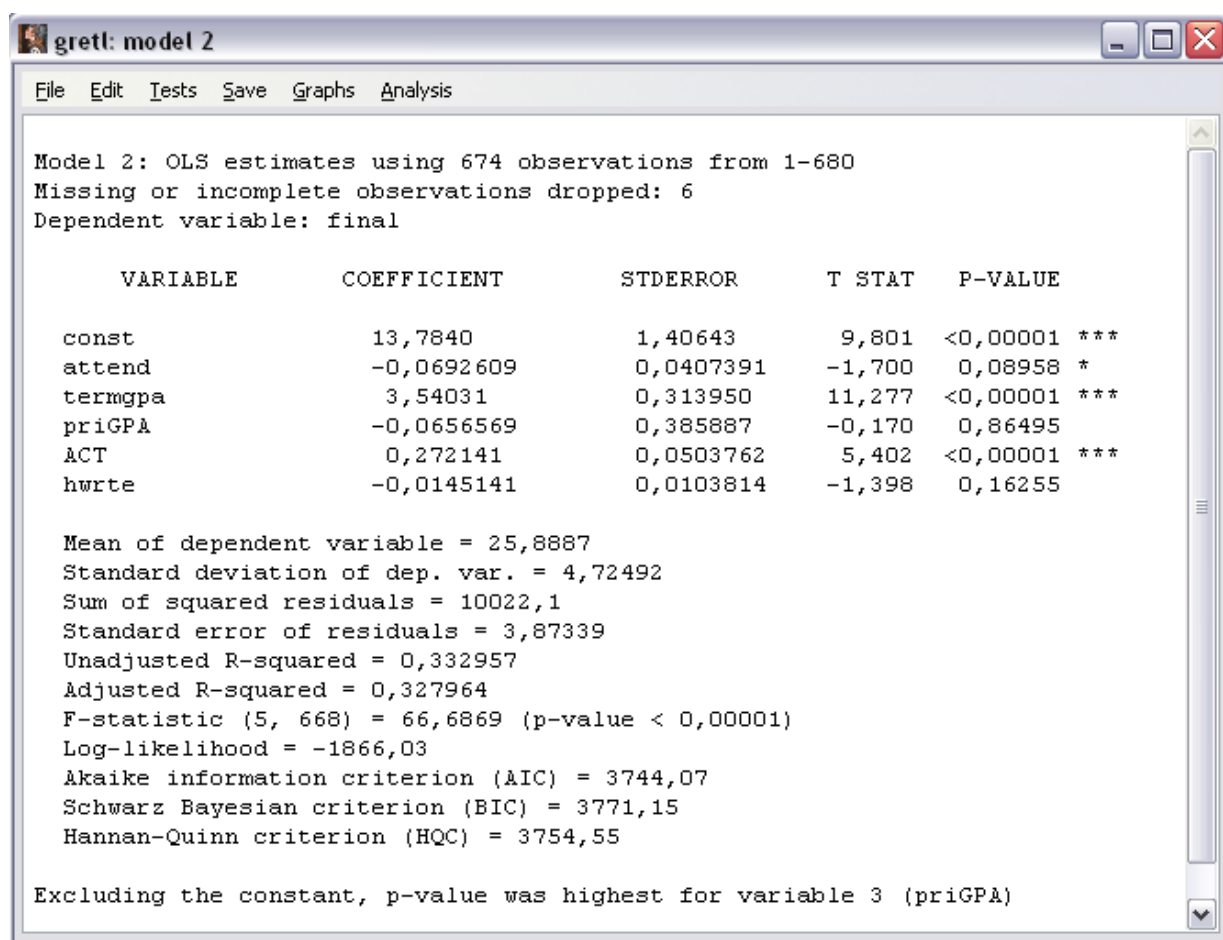


Рисунок 2 - Окно результатов моделирования с применением 1МНК

Рассмотрим сущность показателей (таблица 1), используемых в таблице регрессии окна результатов моделирования (рисунок 2):

Таблица 1- Показатели таблицы регрессии

Variable -	независимая переменная, существенность влияния которой необходимо оценить.
Coefficient-	коэффициент (параметр) модели, α_i , значимость которого необходимо оценить.
Std. Error-	стандартная ошибка параметра модели, является оценкой среднеквадратичного отклонения параметра регрессии от его истинного значения, даёт общую оценку степени точности параметра.
T-STATISTIC-	расчётные значения t- критерия Стьюдента, отношение значения параметра α_i к STD.ERROR, формула (5).
P-value-	показывает вероятность того, что соответствующее значение критерия для генеральной совокупности может оказаться больше, чем расчётное значение по рассматриваемой выборке. Если p-value не превышает уровень значимости, то коэффициент является значимым и принимается альтернативная гипотеза.
Mean of dependent variable-	среднее значение зависимой переменной (y).
Standard deviation of dep. var.-	стандартное (среднеквадратическое) отклонение зависимой переменной (y) – корень квадратный из дисперсии, мера разброса данных.
Sum of squared residuals -	сумма квадратов остатков ($RSS = Y^T \cdot Y - \alpha^T \cdot X^T \cdot Y$), измеряет необъяснённую часть вариации зависимой переменной, используется как основная минимизируемая величина в МНК.
Standard error of residuals-	стандартная ошибка регрессии (среднеквадратическое отклонение ошибки), оценивает степень соответствия модели эмпирическим данным и качество оценивания; измеряет величину квадрата ошибки, приходящейся на одну степень свободы модели $(RSS^2/(n-k))^{1/2}$.
Unadjusted R ² -	нескорректированный коэффициент детерминации R^2 - показывает долю объяснённой (уравнением регрессии) дисперсии зависимой переменной y, формула (7).
Adjusted R ² -	скорректированный коэффициент детерминации, используемый при необходимости учёта количества наблюдений и оцениваемых параметров, чтобы обеспечить сопоставимость различных моделей.
F-statistic-	расчетное значение F-критерия Фишера, формула (6), отношение объяснённой суммы квадратов (в расчёте на одну независимую переменную) к остаточной сумме квадратов (в расчёте на одну степень свободы)
Log-likelihood -	логарифмическая функция правдоподобия. Функция правдоподобия – это плотность распределения y.
Akaike information criterion-	информационный критерий Акайке, анализирует правильность спецификации модели. Позволяет выбирать наилучшую модель из множества различных спецификаций.

Продолжение таблицы 1

Schwarz Bayesian criterion-	информационный Байесовский критерий Шварца, анализирует правильность спецификации модели, позволяет выбирать наилучшую модель из множества различных спецификаций.
Hannan-Quinn criterion-	информационный критерий Хеннана – Куинна, анализирует правильность спецификации модели
OLS	1МНК
Model 2: estimates using 674 observations from 1-680	модель2: использует для оценки 674 наблюдения из 680
Missing or incomplete observations dropped	число пропущенных наблюдений
Dependent variable	зависимая переменная (y)

Пример проверки адекватности регрессионной модели

1. Шаг Оценим существенность влияния каждой объясняющей переменной (attend, termGPA, priGPA, ACT, hwtte согласно приведённому выше примеру, формула (9), на зависимую переменную final, для этого необходимо оценить значимость полученных параметров α_i (рисунок 2), используя t-критерий Стьюдента.

Сформулируем нулевую гипотезу о не значимости коэффициента ($\alpha_i = 0$, и лишь в силу случайных обстоятельств оказался равным проверяемой величине) и альтернативную – о значимости ($\alpha_i \neq 0$), а также выберем уровень значимости (1%, 5%, или 10% - максимально допустимая вероятность ошибочного принятия альтернативной гипотезы).

В оцениваемой модели (формула 9) (рисунок 2) существенные параметры при уровне значимости 1% обозначены ***, 5% - **, 10% - *. Обозначение звёздочками облегчает быстрое оценивание значимости параметров, в рассматриваемом примере существенными являются только константа const и коэффициенты при переменных termGPA и ACT (во всех трёх случаях вероятность ошибки при принятии гипотезы об их значимости P-VALUE=0,001%).

В последнем столбце представляется эмпирический уровень значимости P-VALUE (вероятность допустить ошибку при принятии альтернативной гипотезы, т.е. вероятность того, что значение t-критерия для генеральной совокупности превысит его расчётное значение по данной выборке), который позволяет проверить гипотезы о значимости каждого коэффициента и осуществить отбор существенных (P-value меньше выбранного уровня значимости) и наиболее слабых переменных модели (P-value больше выбранного уровня значимости). В рассмотренном примере самой слабой является переменная PriGPA - вероятность ошибки при принятии гипотезы о её значимости 86,5%, на которую также указывает сообщение в последней строке окна.

Значения столбца T- Stat (рисунок 2), представляющие собой отношение соответствующих величин в столбцах COEFFICIENT и STDERROR, показывают расчётные значения t- критерий Стьюдента. Согласно методу отбора объясняющих переменных *a posteriori* предполагается исключение переменных с минимальными (по модулю) значением t- критерия, в рассматриваемом случае - переменных attend, priGPA, и hwrtte.

2 Шаг. Оценим значимость (пригодность) модели (формула 9) в целом, используя показатели: F-критерий Фишера, коэффициент детерминации R^2 (Unadjusted R^2 и Adjusted R^2), сумма квадратов остатков (RSS, Sum of squared residuals), стандартная ошибка регрессии (Standard error of residuals), информационные критерии (Akaike information criterion, Schwarz Bayesian criterion, Hannan-Quinn criterion).

В рассматриваемом примере F-критерий Фишера F-statistic (5, 668) = 66,6869 для p-value < 0,00001. Поскольку p-value меньше выбранного уровня значимости (p=1%) принимается решение о принятии альтернативной гипотезы, т.е. об адекватности модели в целом. Однако $R^2 = 33,3\%$, что свидетельствует о невысоком уровне объяснения моделью фактических данных, однако согласно F-тесту, он может быть признан достаточно существенным.

Т.о. в результате анализа рассматриваемой модели на адекватность можно сделать вывод: модель по F-критерию Фишера адекватна, но три коэффициента регрессии (при переменных attend, priGPA, и hwrtte) незначимы. В этом случае модель пригодна для принятия некоторых решений относительно зависимости переменной final от переменных termgpa и АСТ, но не для производства прогнозов.

Пример 2. Согласно вышеизложенным рекомендациям исключим из полученной в Примере 1. модели, формула (9), переменные attend, priGPA, и hwrtte и повторим рассмотренную последовательность действий для получения линейной регрессионной модели, устанавливающей зависимость переменной FINAL от АСТ и termgpa. Получим скорректированную модель $final = 10,8 + 0,339 \text{ АСТ} + 2,87 \text{ termgpa} + u$ (рисунок 3), в которой все переменные существенны и модель в целом пригодна для практического использования (согласно рассмотренным выше критериям) для принятия решений и составления прогнозов.

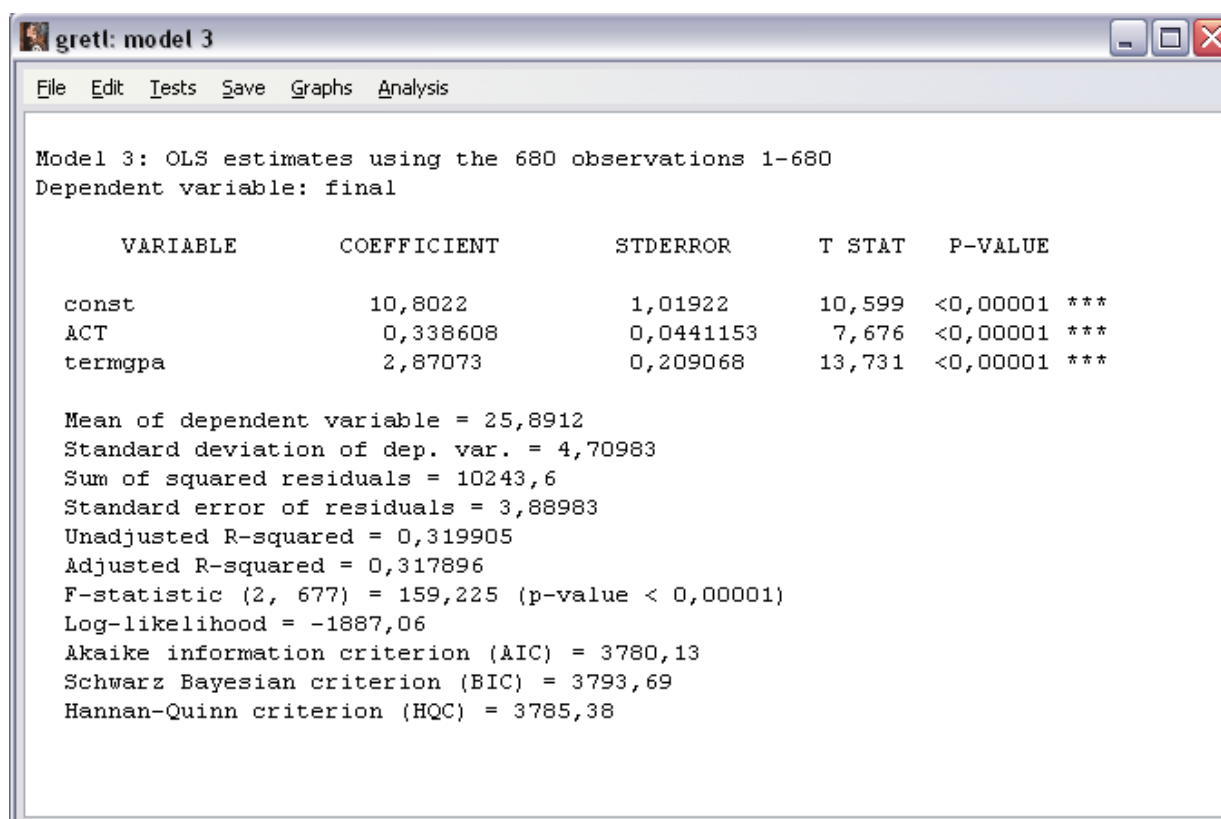


Рисунок 3 - Окно результатов моделирования с применением 1МНК, скорректированная модель

Сохраним значения остатков данной модели как отдельную переменную RESIDUALS набора attend.gdt при помощи функции **Save\Residuals** окна результатов моделирования (рисунок 3). После нажатия кнопки ОК диалогового окна, данная переменная добавится в список переменных рассматриваемого набора данных attend.gdt. Аналогичным образом сохраним модельные значения результативного признака (final) как FITSfinal при помощи функции **Save\Fitted Values**.

Пример построения графика регрессионной модели

Для графического отражения фактических и модельных данных рассмотренного Примера 2. необходимо обратиться к команде **Graphs\Fitted, Actual Plot\ Against ACT and Termgpa** окна результатов моделирования (рисунок 3).

Получим графическое изображение фактических и модельных данных (рисунок 4). Левой кнопкой мыши возможно вращать данное изображение для удобства его просмотра.

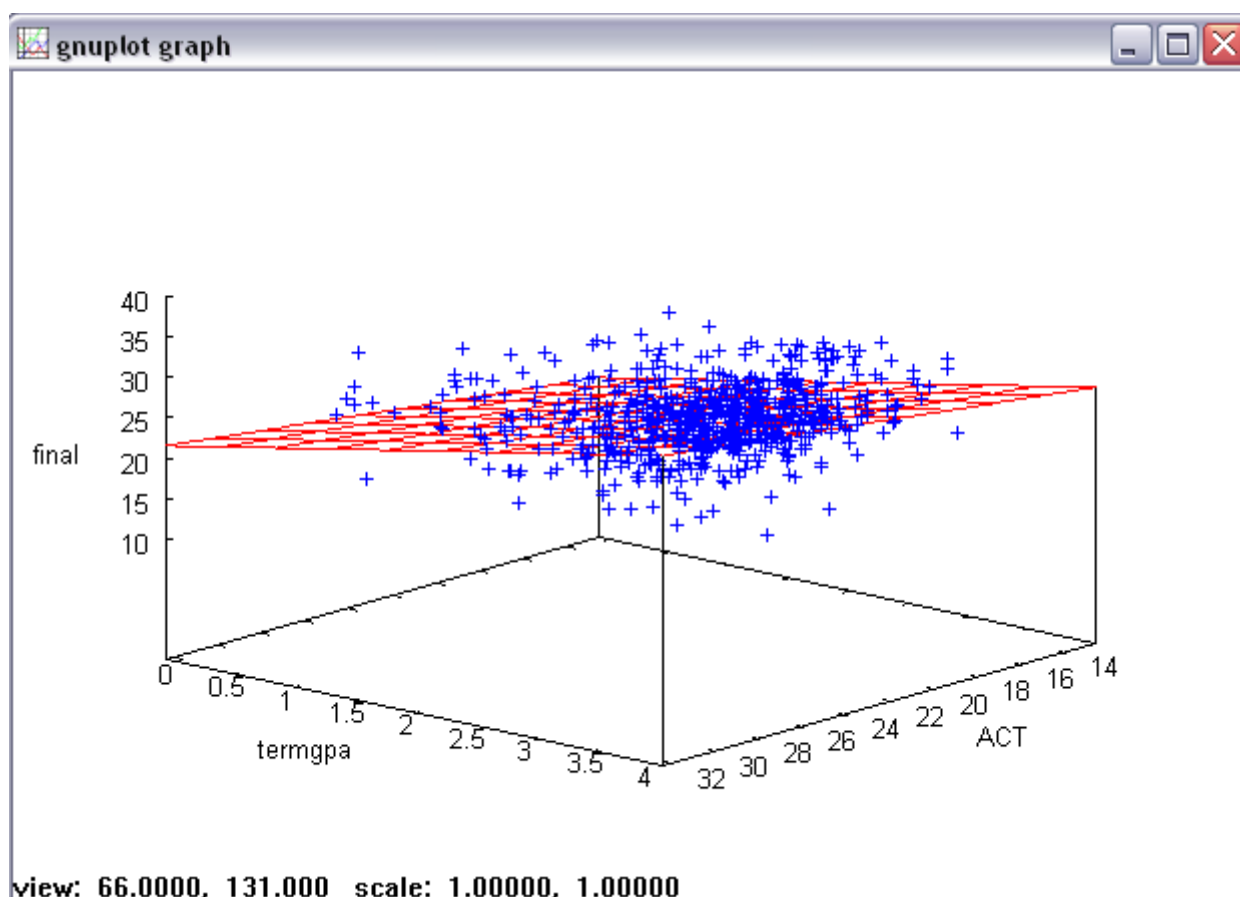


Рисунок 4 - Исходные данные и график функции
 $final = 10,8 + 0,339ACT + 2,87termgpa + u$

3.2. Анализ выполнения предпосылок 1МНК

Проверим условия Гаусса—Маркова при помощи инструментария GRETЛ для данных примера 2.:

1. Нулевая средняя величина (математическое ожидание) остатков, $M(u_i) = 0$.

Для проверки данного утверждения выберем щелчком мыши ранее созданную переменную RESIDUALS в списке переменных стартового экрана и обратимся к функции **View\ Summary Statistics** (рисунок 5), в открывшемся окне среднее значение остатков (mean) равна 0.

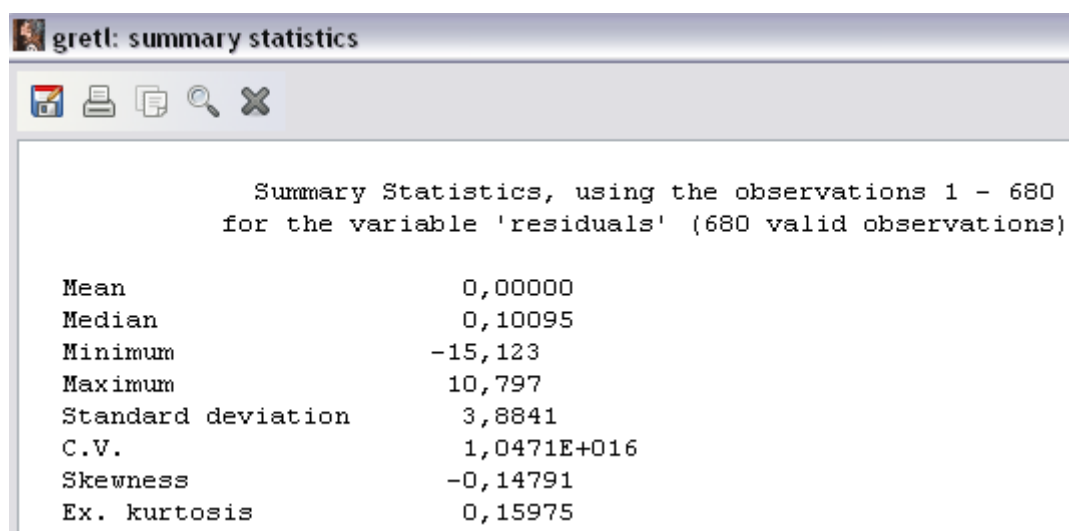


Рисунок 5 - Общая статистика для переменной RESIDUALS

2. Проверка условия гомоскедастичности остатков:

Проверку можно выполнить в окне текущей модели (рисунок 3), для чего в меню следует выбрать Tests\heteroskedasticity. Окно результатов в этом случае имеет вид, представленный на рисунке 6. Значение P-value = 0,734603 больше уровня значимости 0,01 свидетельствует о том, что нулевую гипотезу следует принять и условие гомоскедастичности остатков выполняется.

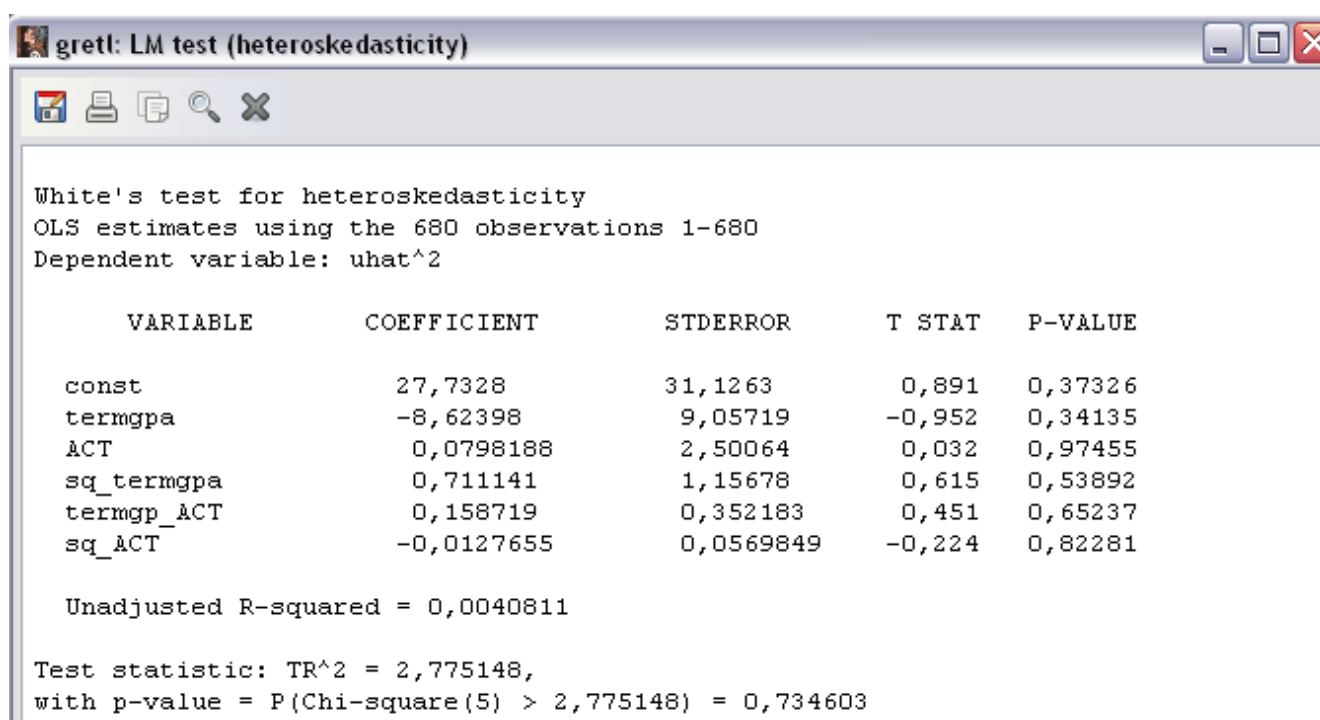


Рисунок 6 - Тест Уайта на гетероскедастичность остатков

3. Отсутствие систематической связи между значениями случайной составляющей u_i в любых двух наблюдениях (отсутствие автокорреляции остатков).

Определим наличие автокорреляции остатков рассматриваемой модели.

Экспортируем ряд значений созданной в Примере 2. переменной Residuals (остатки модели) в файл Residuals.csv (File\Export data\ CSV..., поставив флажок comma (,) в разделе decimal point character). Создадим в файле Residuals.csv новую переменную Residuals1, которая отличается на один лаг от переменной Residuals (длина рядов сокращается на одно наблюдение), затем сохраним файл в формате Residuals.xls. Создадим новый набор данных в Gretl (File\New dataset) и импортируем в него данные из файла Residuals.xls (File\Open Data\Import\Excel), ответив «по» на вопрос о смене типа данных.

Рассчитаем коэффициент корреляции между данными переменными, обратившись к функции **View\Correlation matrix**, выбрав переменные Residuals и Residuals1. Получим коэффициент -0,1488, свидетельствующий о несущественной корреляции (корреляция считается сильной, если ее коэффициент выше $|0,6|$).

4. Случайная составляющая должна быть распределена независимо от переменных x и y (случайный характер остатков).

Для проверки строится график зависимости остатков u_i от теоретических значений результативного признака y и x .

Способом, аналогичным описанному выше, построим парную регрессию ошибки RESIDUALS от модельных значений результативного признака FitsFINAL (рисунок 7). В результате получим нулевое значение коэффициента и единичное значение p -value, а также расположение остатков на графике в виде горизонтальной полосы, что свидетельствует об отсутствии данной зависимости и о случайном характере остатков.

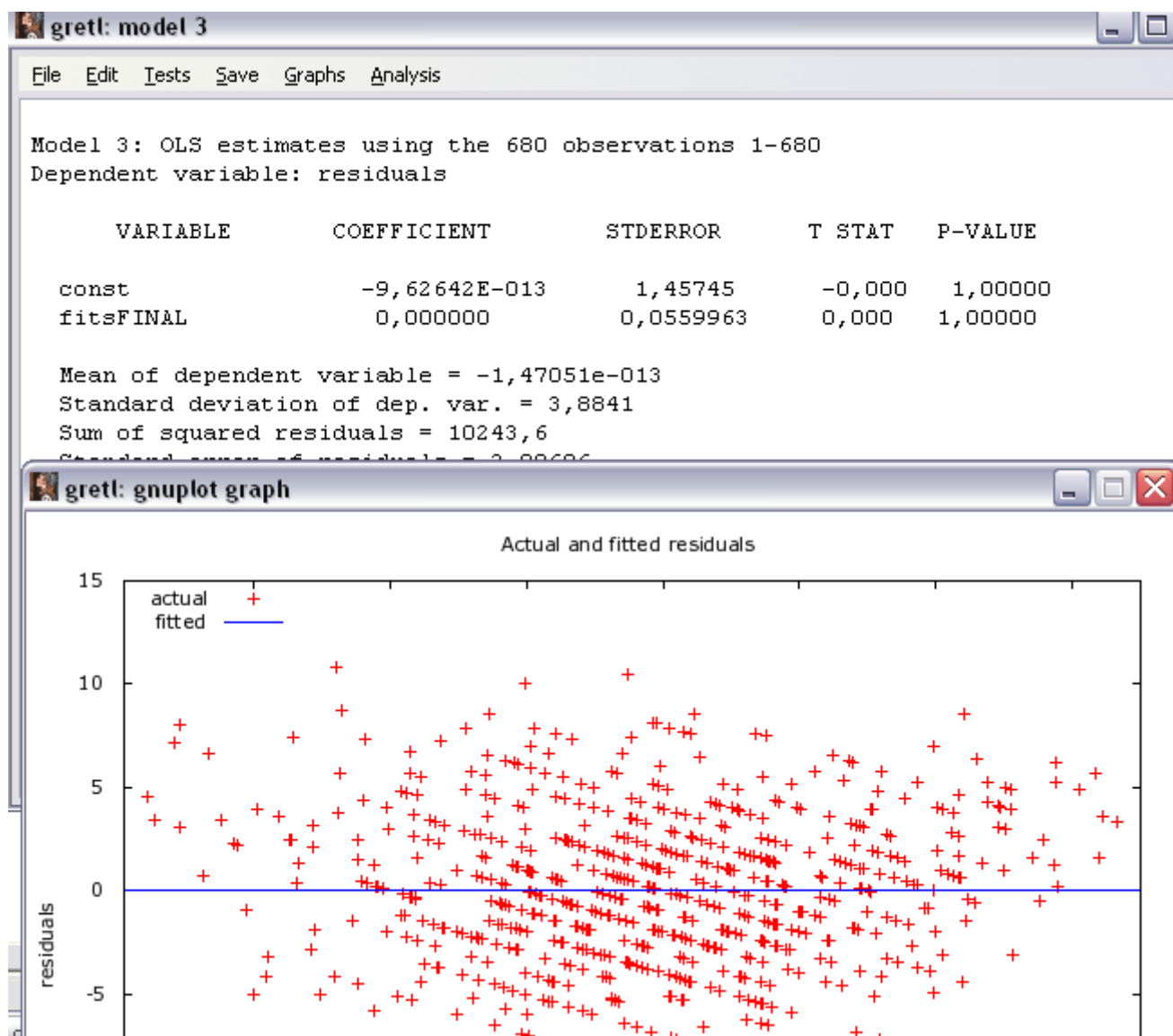


Рисунок 7 - Проверка случайного характера остатков

Проверку зависимости остатков от переменных termgpa и АСТ можно осуществить из окна модели $\text{final} = 10,8 + 0,339\text{АСТ} + 2,87\text{termgpa} + u$ (рисунок 3), построив соответствующие графики Graphs\Residual Plot\Against termgpa (Against АСТ) (рисунок 8, 9). На полученных графиках остатки также расположены в виде горизонтальных полос, что свидетельствует об отсутствии соответствующих зависимостей.

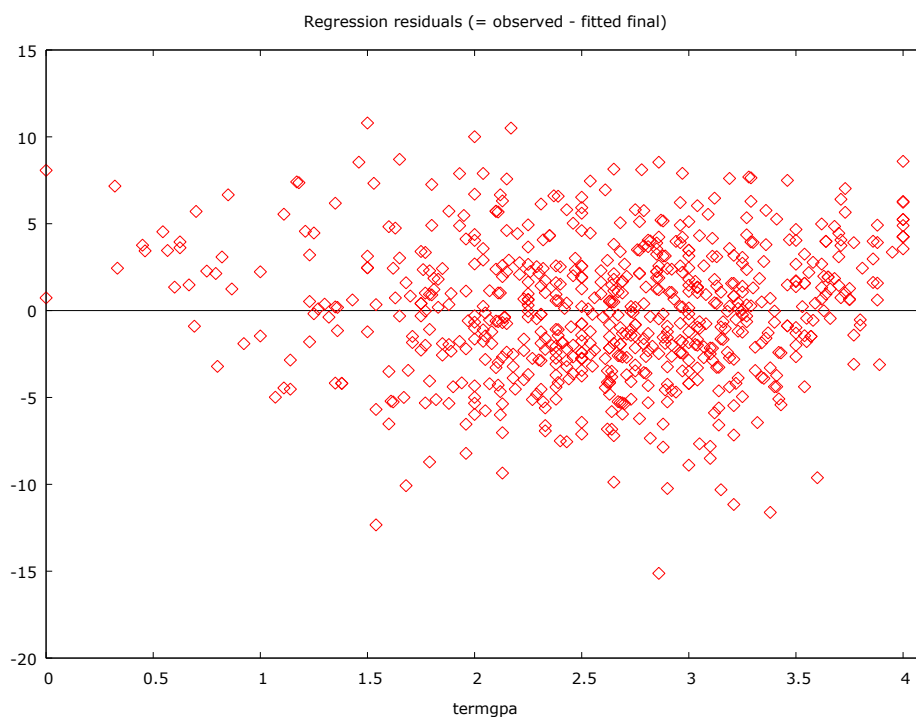


Рисунок 8 - График зависимости остатков от переменной termgra

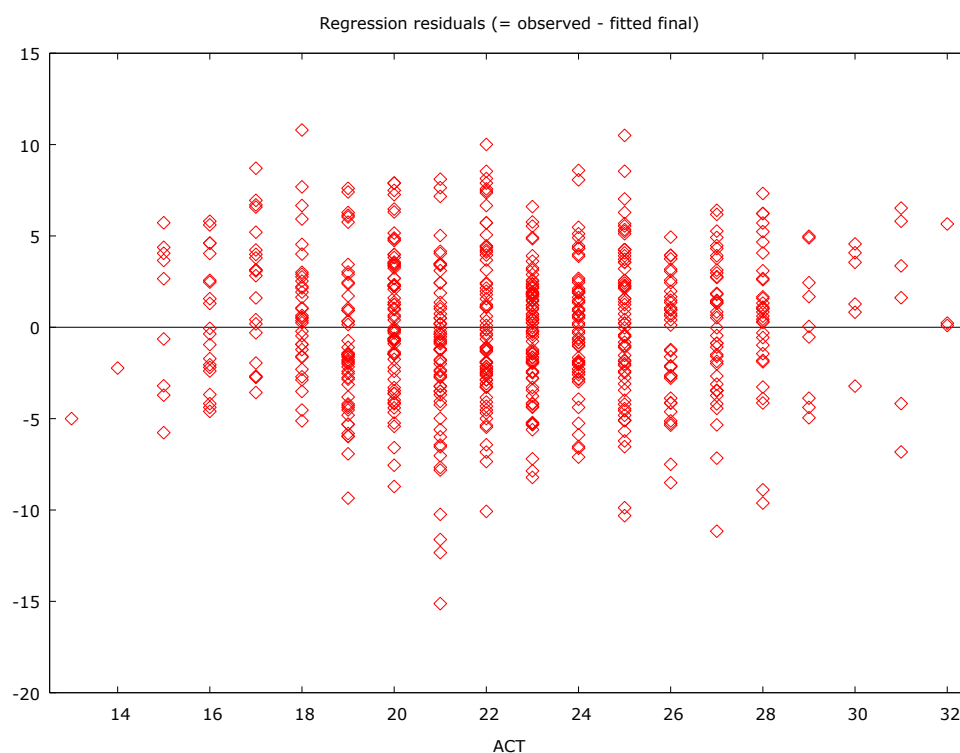


Рисунок 9 - График зависимости остатков от переменной АСТ

Из вышесказанного можно установить, что выполняются все предпосылки для применения МНК для определения параметров рассматриваемой модели (полученной в примере 2.). Построенная модель $final = 10,8 + 0,339АСТ + 2,87termgra + u$ на основе ее проверки по F-критерию Фишера в целом адекватна, и все коэффициенты регрессии значимы (в результате проверки по t-критерию

Стьюдента). Такая модель может быть использована для принятия решений и осуществления прогнозов.

4. ПОРЯДОК ВЫПОЛНЕНИЯ ЛАБОРАТОРНОЙ РАБОТЫ

1. Открыть набор данных **File\Open Data\ Sample File** (закладка **Wooldridge**) в соответствии с номером варианта (рисунок 10). Для каждого варианта также указаны номера переменных ID# для y , x_1 , x_2 , информацию о которых можно просмотреть, обратившись к команде **Data\Print Description**.

2. Провести оценку параметров линейной регрессионной модели $y = \alpha_0 + \alpha_1 \cdot x_1 + \alpha_2 x_2 + u$ методом 1МНК.

3. Оценить адекватность регрессионной модели в целом и значимость её отдельных параметров

4. Проверить были ли все предпосылки к тому, чтобы применять 1МНК и линейное уравнение регрессии к исходным данным.

№ Вар.	Файл	ID# для	Y	X1	X2
1	401k Pension plan generosity and participation	3	4	6	
2	401ksubs Wealth, income and pension plans	2	5	6	
3	fish Daily data from Fulton fish market	4	2	3	
4	affairs Incidence of extra-marital affairs	11	4	7	
5	consump Consumption, income and interest rate	9	7	2	
6	earns Earnings, productivity and hours (macro)	1	2	3	
7	ceosal2 Firm performance and CEO salary	1	6	3	
8	cps91 Wages and their determinants, 1991 (large)	3	1	4	
9	beautv Physical attractiveness and earnings	1	2	5	
10	cps78_85 Wages and demographics, 1978 and 1985	8	1	9	

Рисунок 10 - Варианты заданий: название файла на закладке Wooldridge, номера варианта и переменных.

5. СОДЕРЖАНИЕ ОТЧЕТА О ВЫПОЛНЕНИИ ЛАБОРАТОРНОЙ РАБОТЫ

- 1) Название и цель работы.
- 2) Постановка задачи.
- 3) Этапы выполнения задачи в Gretl.
- 4) Выводы.

ЛАБОРАТОРНАЯ РАБОТА №3

ПРИМЕНЕНИЕ GRETЛ 1.7.1. ПРИ ПОСТРОЕНИИ И АНАЛИЗЕ РЕГРЕССИОННЫХ МОДЕЛЕЙ С ГЕТЕРОСКЕДАСТИЧНОЙ СЛУЧАЙНОЙ СОСТАВЛЯЮЩЕЙ

1. ЦЕЛЬ РАБОТЫ

Цель работы заключается в освоении инструментария системы Gretl в области построения и анализа регрессионных моделей с гетероскедастичной случайной составляющей для выявления и последующего применения ранее неизвестных закономерностей в имеющихся данных в процессе подготовки и принятия решений менеджерами компаний.

2. ТЕОРЕТИЧЕСКИЕ СВЕДЕНИЯ

Для регрессионной модели, формула (1), построенной по фактическим данным типа срез данных (cross-sectional) дисперсия случайных отклонений (ошибок) часто представляет собой переменную величину $\sigma^2_{u_i} \neq const$.

$$y = \tilde{y}_x + u = \alpha_0 + \alpha_1 \cdot x_1 + \dots + \alpha_n x_n + u, \quad (1)$$

где y — фактическое значение результативного признака;

$\tilde{y}_x$ - модельное значение результативного признака;

α_i – параметр регрессионной модели;

x_i - признак-фактор;

u — случайная ошибка.

Данная ситуация представляет собой проблему гетероскедастичности («неодинакового разброса») - нарушения, возникающего при невыполнении одного из классических предположений линейного регрессионного анализа о постоянстве дисперсий случайных отклонений (гомоскедастичности или «одинакового разброса» $\sigma_{u_i} = const$), при этом остальные условия Гаусса-Маркова выполняются:

- математическое ожидание случайной составляющей, $M(u_i) = 0$
- отсутствие автокорреляции остатков (взаимосвязи u_i и u_{i-1}).
- случайный характер остатков – их независимость от y_i и x_i .

Случаи гетероскедастичности и гомоскедастичности показаны на рисунках 1 и 2 соответственно.

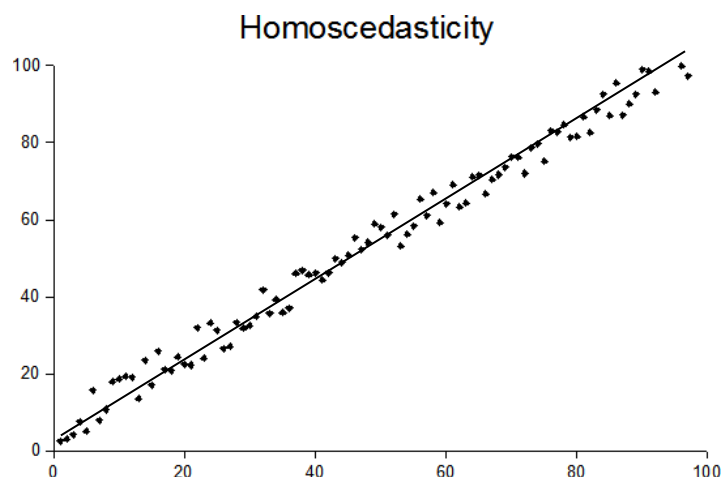


Рисунок 1 - Иллюстрация случайных данных и модели с гомоскедастичностью

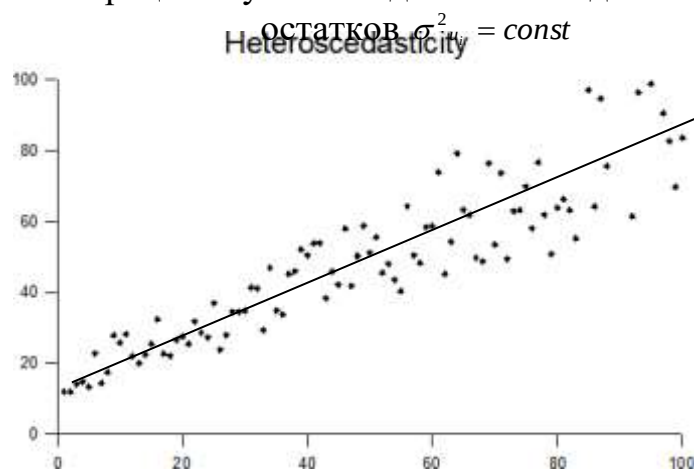


Рисунок 2 - Иллюстрация случайных данных и модели с гетероскедастичностью
остатков $\sigma^2_{u_i} \neq const$

Гетероскедастичность приводит к тому, что при применении обычного метода наименьших квадратов (МНК) полученные параметры $\alpha_0 \dots \alpha_n$ модели, формула (1), больше не представляют собой наиболее эффективные оценки или не являются оценками с минимальной дисперсией.

Наблюдение, для которого теоретическое распределение ошибки имеет малое стандартное отклонение, будет обычно находится близко к линии регрессии и, следовательно, может стать хорошим ориентиром, указывающим на место этой линии. В противоположность этому наблюдение, где теоретическое распределение имеет большое стандартное отклонение, не сможет в той же мере помочь в определении местоположения линии регрессии. Обычный МНК не делает различия между качеством наблюдений, придавая одинаковые "веса" каждому из них независимо от того, является ли наблюдение существенным или несущественным для определения местоположения этой линии. Следовательно, обычным МНК мы получим неэффективные оценки коэффициентов.

Также результаты t- и F- тестов будут ненадёжными, т.к. мы получим неверные оценки стандартных ошибок параметров $\alpha_0 \dots \alpha_i$ (STDERROR), т.к. они вычисляются на основе предположения о том, что остатки модели

гомоскедастичны, что скажется на правильности расчёта t- и F- статистик и приведёт к принятию ошибочных гипотез.

Поскольку в данном случае использование обычного метода наименьших квадратов (МНК) неэффективно, необходимо сделать поправку на гетероскедастичность, применив взвешенный метод наименьших квадратов ВМНК (WLS) для её устранения.

Построение гетероскедастичной регрессионной модели состоит из двух этапов:

1. Обнаружение гетероскедастичности случайной составляющей,
2. Оценивание модели с использованием взвешенного метода наименьших квадратов (WLS).

На первом этапе в случае однофакторной регрессии $y_i = f(x_i) + u_i = \alpha_0 + \alpha_1 \cdot x_i + u_i$ изначально проводится графический анализ остатков – строится и анализируется зависимость квадратов ошибок от x_i или от теоретического значения $\tilde{y}_{x_i}$, или строится $x_i - y_i$ диаграмма рассеяния. При множественной регрессии графический анализ также возможен для каждой из объясняющих переменных. Рост дисперсии с ростом одного из факторов свидетельствует о гетероскедастичности.

Затем проводится один из формальных тестов на гетероскедастичность: тест ранговой корреляции Спирмена, тест Парка (The Park test), тест Голдфелда-Квандта (Goldfeld-Quandt test), тест Бреуша-Погана, тест Глейзера, или тест Уайта (White's test) и осуществляется интерпретация результатов теста.

В каждом тесте пытаются опровергнуть гипотезу о гомоскедастичности, если это удаётся, то можно сделать вывод, что в модели наблюдается гетероскедастичность.

Рассмотрим алгоритм теста Уайта на гетероскедастичность, не требующего нормальности распределения остатков. Алгоритм состоит из следующих шагов:

- Получение остатков оцененной регрессионной модели $u = y - \alpha_0 - \alpha_1 \cdot x_1 - \dots - \alpha_n x_n$
- Оценивание вспомогательного уравнения регрессии квадратов остатков относительно комплекса переменных модели, их произведений и их квадратов $u^2 = \beta_0 + \beta_1 \cdot x_1 + \dots + \beta_n x_n + \beta_{n+1} \cdot x_1^2 + \dots + \beta_{2n} x_n^2 + \beta_{2n+1} \cdot x_1 \cdot x_2 + \dots + \beta_p \cdot x_n \cdot x_{n-1} + \varepsilon$
- Проверка общей значимости уравнения с помощью критерия χ^2 . Тестовой статистикой является величина $n \cdot R^2$ (n - число наблюдений, R^2 - коэффициент детерминации). Число степеней свободы равно числу регрессоров вспомогательного уравнения. Если $n \cdot R^2 \succ \chi_{\varepsilon \varepsilon \varepsilon}^2$, то нулевая гипотеза гомоскедастичности $H_0: \beta_0 = \beta_1 = \dots = \beta_i = 0$ отвергается.

На втором этапе для оценки моделей с гетероскедастичностью используется взвешенный метод наименьших квадратов (ВМНК - WLS). Метод ВМНК, как и

1МНК, применим к однофакторной и множественной линейной регрессии и использует то же правило минимизации суммы квадратов остатков, RSS, но вместо одинаковых весов для каждого наблюдения им приписываются значения, обратные соответствующим дисперсиям ошибки, что отражено формулой (2).

$$\sum_{i=1}^n w_i (y_i - \tilde{y}_{x_i})^2 = \sum_{i=1}^n \frac{1}{\sigma_{u_i}^2} (y_i - \tilde{y}_{x_i})^2 \rightarrow \min, \quad (2)$$

где $y_i, \tilde{y}_i$ - фактическое и модельное i -е значение зависимой переменной;

w_i - вес i -го наблюдения, $w_i = \frac{1}{\sigma_{u_i}^2}$;

$\sigma_{u_i}^2$ - дисперсия i -й случайной составляющей.

Тогда коэффициенты линейной регрессии находятся по формуле (3).

$$\hat{\alpha} = [X^T \cdot W \cdot X]^{-1} \cdot X^T \cdot W \cdot Y, \quad (3)$$

где $W = \text{diag}\{w_1, \dots, w_n\}$ - диагональная матрица весов;

n - число наблюдений.

Дисперсия ошибки $\sigma_{u_i}^2$ чаще всего неизвестна, но возможно существование некоторого соотношения между дисперсией ошибки и значением какой-либо объясняющей переменной в регрессионной модели $y_i = \alpha_0 + \alpha_1 x_{1i} + \dots + \alpha_p x_{pi} + u_i$, например, $\sigma_{u_i}^2 = c \cdot x_{1i}^2$, где c - ненулевая константа и x_{1i} - значение объясняющей переменной x_1 в i -ом наблюдении. В случае подобного соотношения $\sigma_{u_i}^2$ можно считать известными, т.к. постоянная величина « c » не влияет на взвешенную процедуру. Тогда значения весов

$$w_i = \frac{1}{c \cdot x_{1i}^2}.$$

3. ОПИСАНИЕ СРЕДСТВ СИСТЕМЫ GRETL ДЛЯ ВЫПОЛНЕНИЯ РЕГРЕССИОННОГО АНАЛИЗА ПРИ НАЛИЧИИ ГЕТЕРОСКЕДАСТИЧНОСТИ

3.1. Пример обнаружения гетероскедастичности в Gretl

Шаг 1. Выберем команду меню **File\Open Data\Sample File...** и на закладке **Verbeek** двойным щелчком мыши откроем встроенный разработчиками набор данных **bwages.gdt** (рисунок 3).

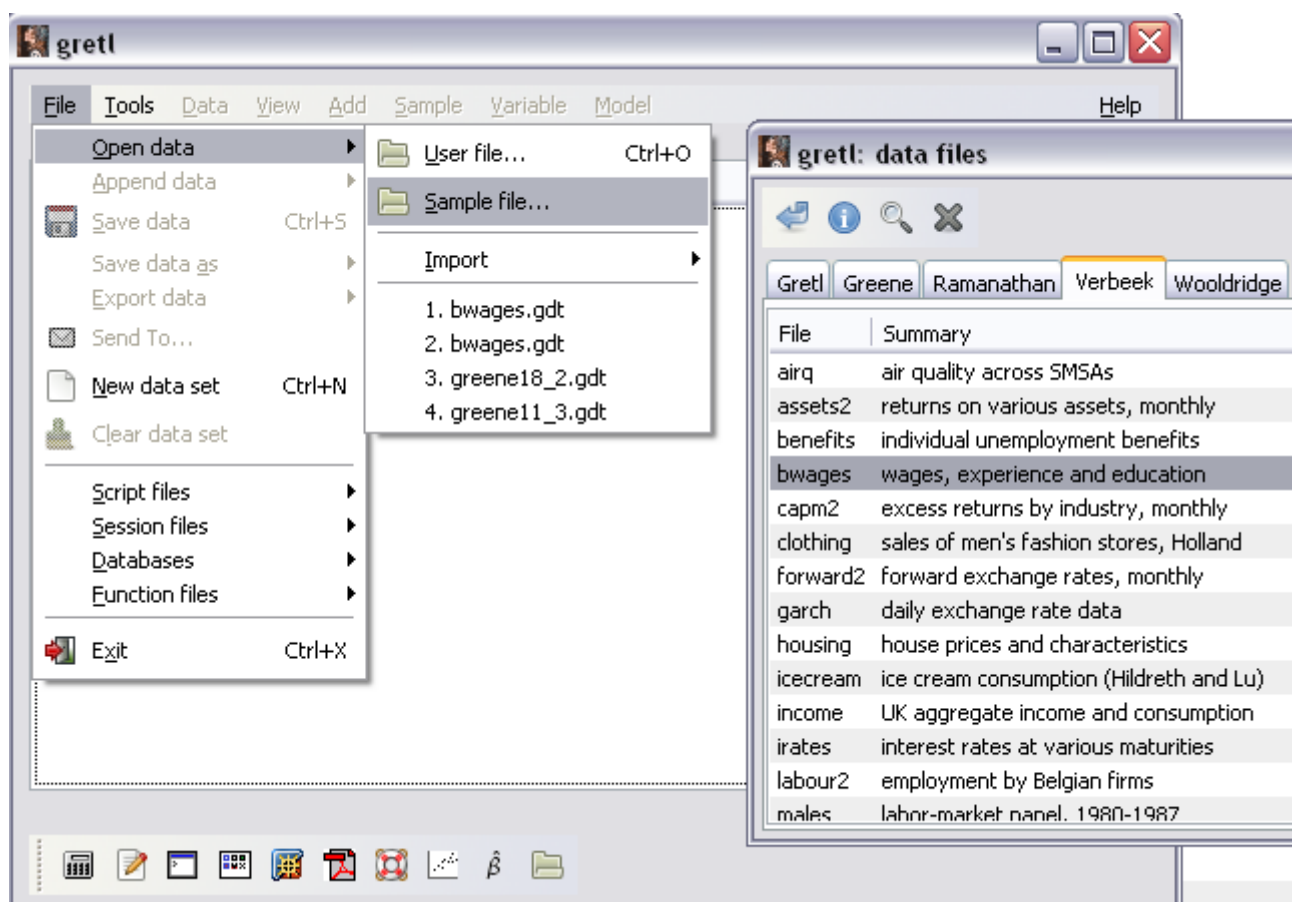


Рисунок 3 – Открытие встроенного набора данных bwages.gdt

Шаг 2. Просмотрим текстовую информацию о наборе данных, выбрав команду **Data\Print Description**. В открывшемся окне (рисунок 4) появится информация о наборе данных **bwages.gdt** в целом и о каждой переменной.

Файл содержит 1472 наблюдения, относящихся к 1994, группе бельгийских семей Европейского Сообщества. Тип данных **undated** (срез данных для фиксированного момента времени – cross-sectional). Переменные:
wage – заработная плата в час, до выплаты налогов (Евро),
educ – уровень образования от 1 (низкий) до 5 (высокий),
exper – опыт работы (лет)
male – фиктивная переменная, принимает значение «1», если мужчина и «0», если женщина. Также представлены логарифмы переменных wage, educ, exper. Требуется оценить влияние уровня образования (educ) в бельгийских семьях на величину заработной платы (wage), используя фактические данные 1994 года по 1472 семьям.

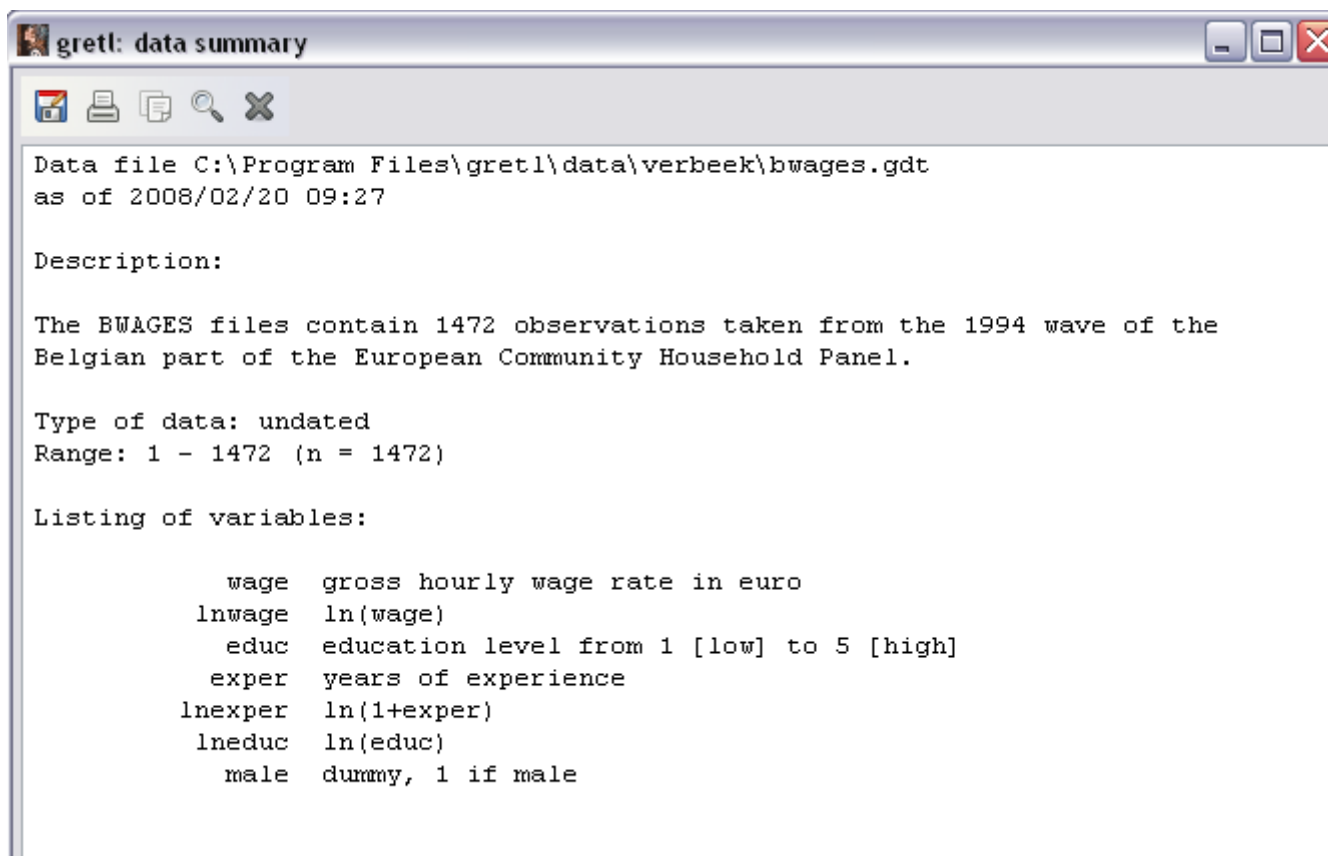


Рисунок 4 – Текстовое описание набора данных bwages.gdt

Шаг 3. Проведём предварительный графический анализ данных, построив диаграммы рассеяния пар переменных educ -wage, exper-wage, male -wage. Выберем команду меню **View\Multiple graphs Vars\ X-Y scatters...**(рисунок 5). В открывшемся окне выберем переменную wage (по оси OY), а переменные educ, exper, male (по оси OX) нажатием кнопок «Choose» и «Add» соответственно.

Как показывает рисунок 6, дисперсия wage (которую будем рассматривать как зависимую переменную Y) увеличивается при росте каждого из факторов **educ**, **exper**, **male** (которые будем рассматривать как независимые переменные, X_i), что свидетельствует о гетероскедастичности.

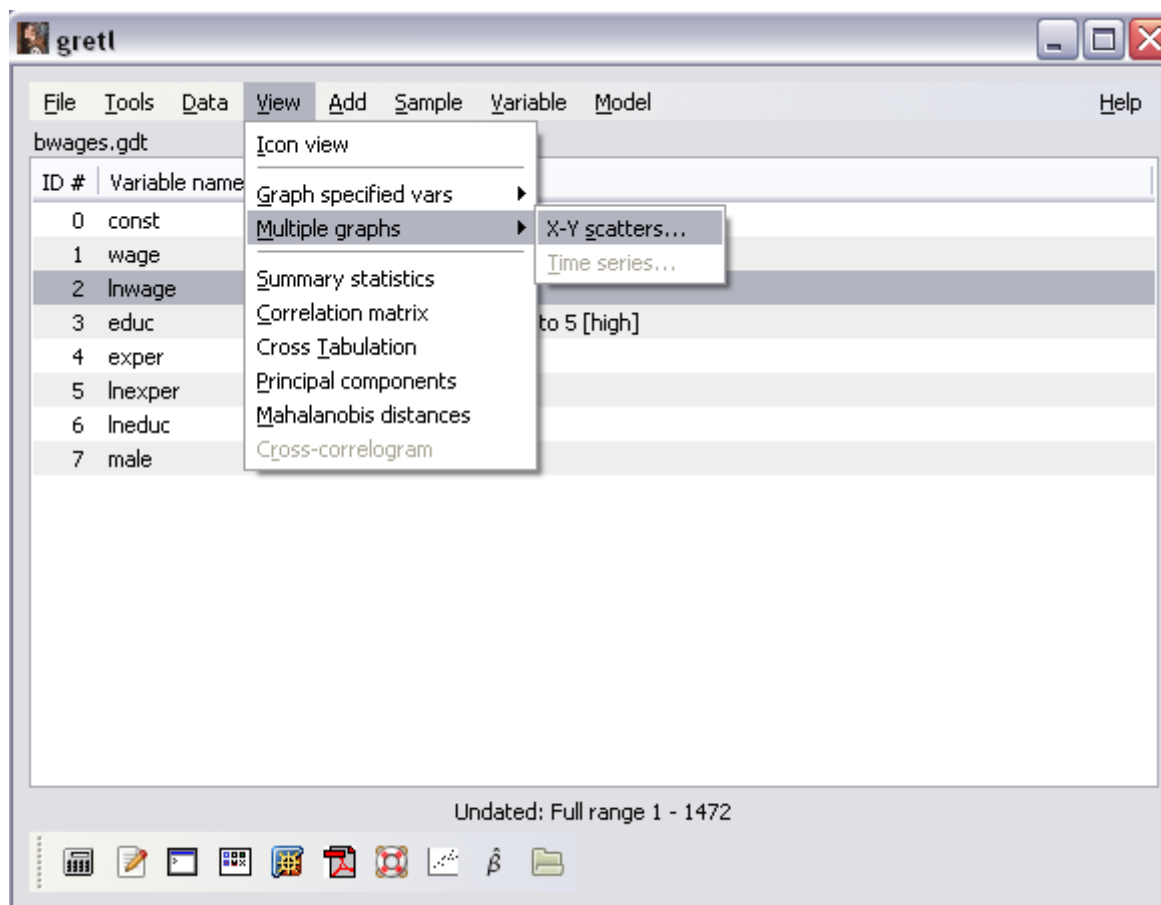


Рисунок 5 – Построение диаграмм рассеяния educ-wage, exper-wage, male-wage

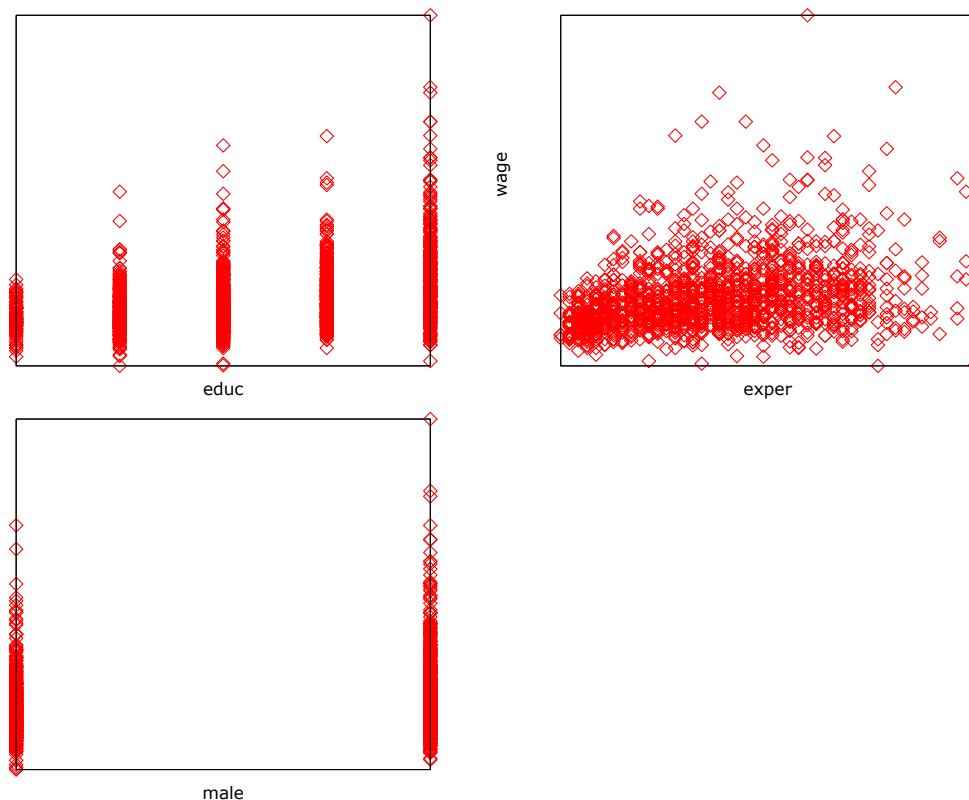


Рисунок 6 – Диаграммы рассеяния пар переменных educ -wage, exper-wage, male –wage

Шаг 4. Проведём оценивание модели $wage = \alpha_0 + \alpha_1 \cdot educ + u$ обычным методом наименьших квадратов 1МНК (OLS). Для этого выберем команду меню **Model\Ordinary Least Squares...** и в открывшемся окне спецификации модели выберем зависимую переменную *wage* при помощи кнопки «Choose» и независимую переменную *educ* при помощи кнопки «Add». После нажатия кнопки ОК появится окно с результатами моделирования (рисунок 7).

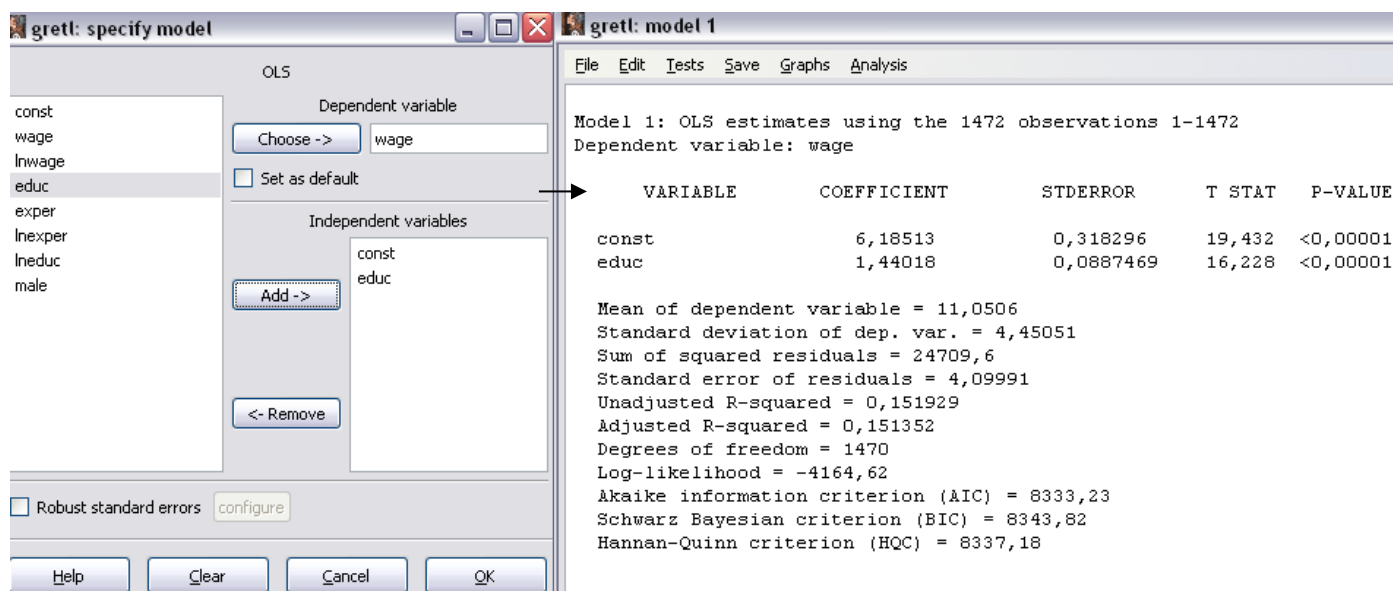


Рисунок 7 – Построение модели $wage = \alpha_0 + \alpha_1 \cdot educ + u$ методом 1МНК (OLS)

По результатам 1МНК в полученной модели $wage = 6,18513 + 1,44018 \cdot educ$ параметры α_0, α_1 являются значимыми при уровнях значимости 5% и 1%, поскольку $P\text{-VALUE} = 0,001\% < 1\% < 5\%$ (принимая альтернативную гипотезу $H_1: \alpha_1 \neq 0$).

Хотя по данным 1МНК можно установить существенность влияния переменной *educ* (образование) на переменную *wage* (з.пл.), необходимо сделать поправку на гетероскедастичность в связи с ошибочностью расчёта величин COEFFICIENT, STDERROR и T-STAT и ненадёжностью результатов данного оценивания -возможностью принятия неправильной гипотезы.

Сохраним величины квадратов остатков в отдельную переменную *usq1* набора данных, обратившись к команде **Save\Squared Residuals** меню окна результатов моделирования (рисунок 7) и нажав кнопку ОК.

Шаг 5. Подтвердим наличие гетероскедастичности, используя формальный тест Уайта. Обратимся к команде **Tests\Heteroskedasticity** меню окна результатов моделирования (рисунок 8).

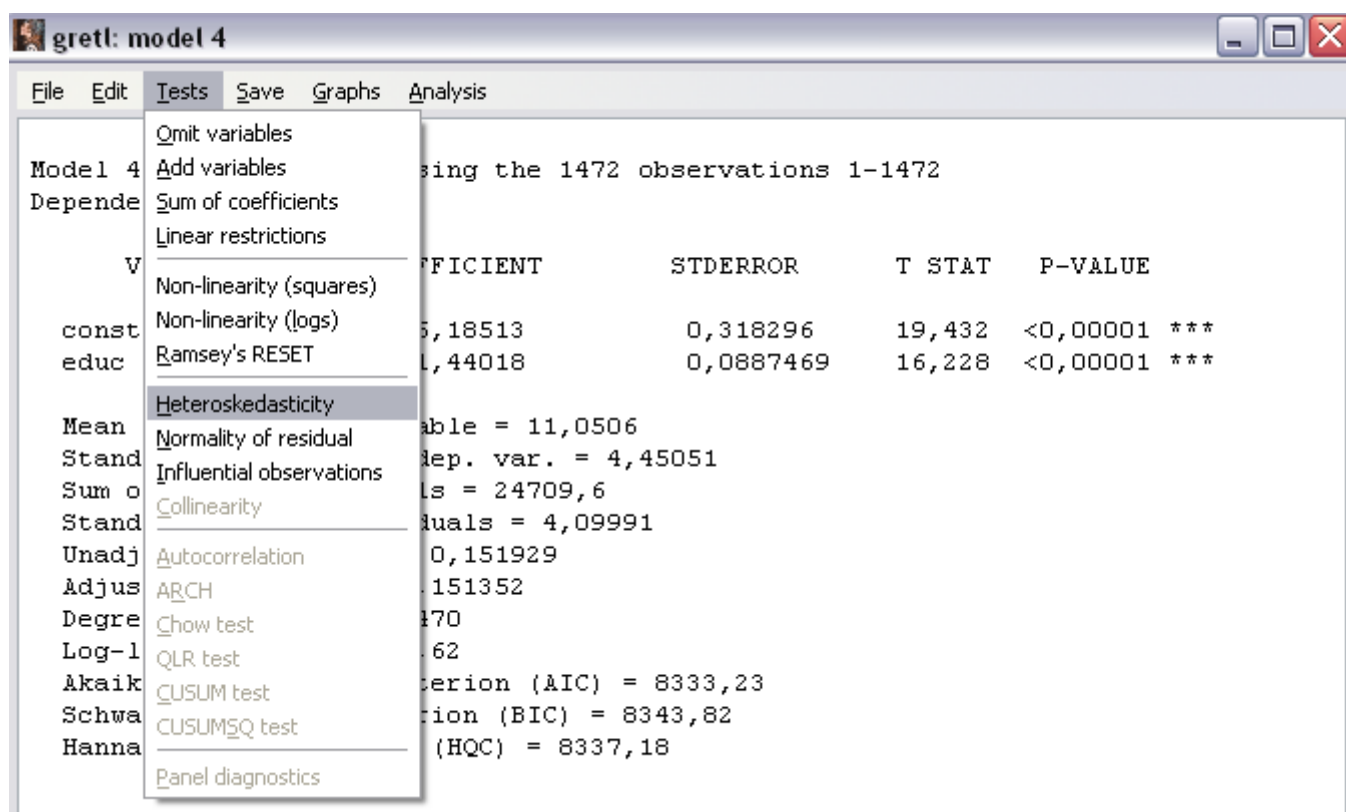


Рисунок 8 – Проведение теста Уайта на гетероскедастичность

В результате теста получим вспомогательную модель регрессии остатков относительно переменной *educ* и её квадрата $u^2 = 16,1382 - 10,3147 \cdot educ + 2,7594 \cdot educ^2$ (рисунок 9).

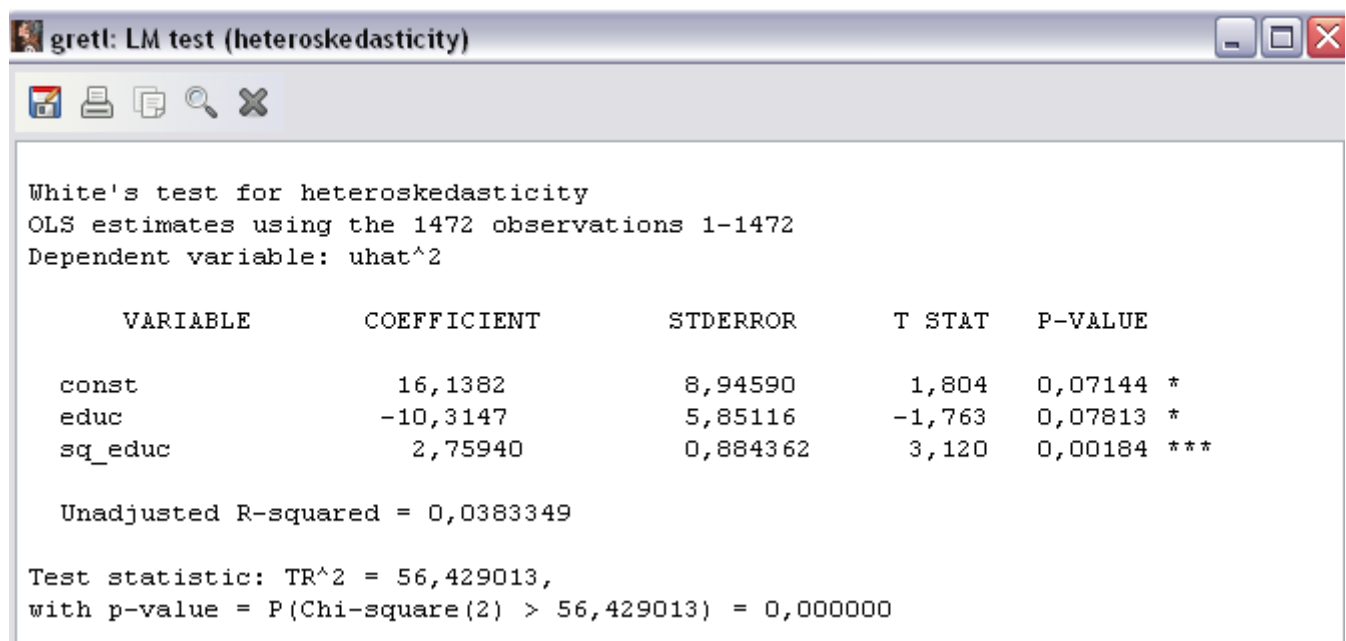


Рисунок 9 – Окно результатов теста Уайта на гетероскедастичность

Проверим общую значимость (адекватность в целом) данной модели, используя критерий χ^2 .

В окне результатов теста (рисунок 9) значение p-value для статистики теста $n \cdot R^2$ (TR^2) составило 0,000%, что меньше уровней значимости 1% и 5% и свидетельствует о наличии гетероскедастичности (адекватности вспомогательной модели, неравенстве нулю всех её параметров).

Т.о. отвергаем гипотезу H_0 о гомоскедастичности (равенстве нулю всех её параметров), поскольку расчётное значение статистики $n \cdot R^2 = TR^2 = 56,429$ больше χ^2 критического ($n \cdot R^2 > \chi_{\text{крит}}^2$).

Найдём χ^2 критическое для уровня значимости 1% , обратившись к команде основного меню **Tools\Statistical Tables** и введя в открывшемся окне на закладке **Chi-Square** число степеней свободы (2, равное числу регрессоров ($educ, educ^2$) вспомогательной модели) и уровень значимости 0.001 (рисунок 10). Получим $\chi_{\text{крит}}^2 = 13,8155$ (Critical value).

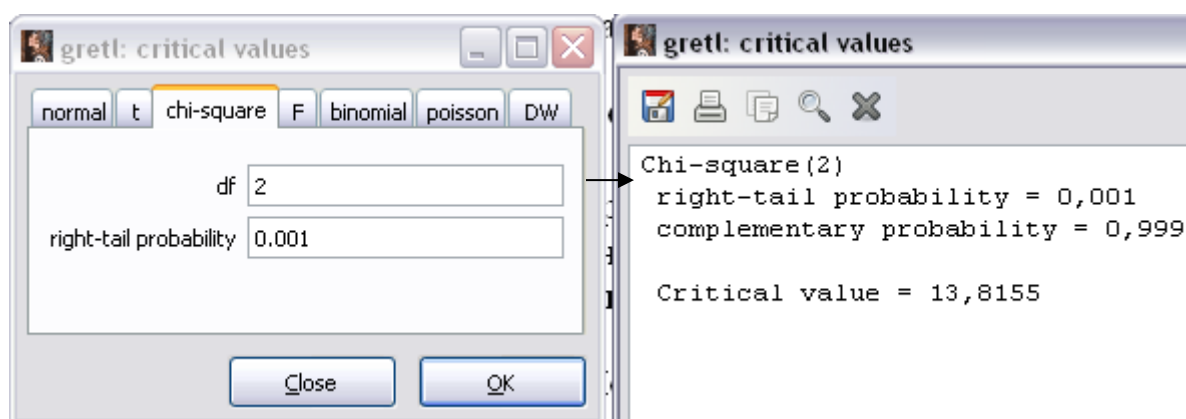


Рисунок 10 – Нахождение критического значения χ^2 для $df=2$, $p=1\%$

Построим график данного χ^2 распределения, выбрав команду основного меню **Tools\Distribution Graphs**, закладку Chi-square и $df=2$ в открывшемся окне (рисунок 11).

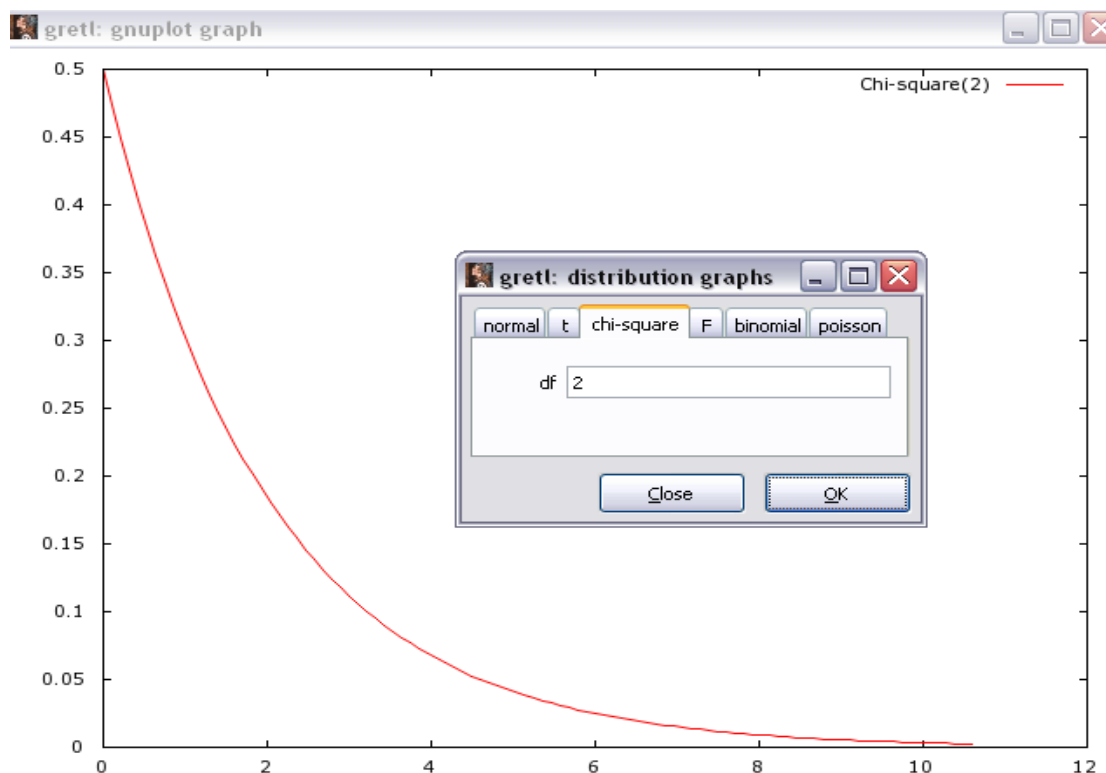


Рисунок 11 - График χ^2 распределения, df=2

3.2. Оценивание гетероскедастичной модели с использованием взвешенного метода наименьших квадратов ВМНК (WLS)

Существует два способа реализации ВМНК (WLS) в пакете Gretl:

1. При выборе из главного меню команды **Model\Other Linear Models\Weighted Least Squares....**(рисунок 12) открывается окно спецификации модели, которое предусматривает выбор из списка переменных открытого набора переменных - веса w . Данная переменная добавляется к набору путём ввода данных вручную или определяется путём ввода соответствующей формулы. Данный способ используется если дисперсия ошибки $\sigma_{u_i}^2$ или известна или неизвестна, но существует соотношение между ней и одной из объясняющих переменных (x_i), например, $\sigma_{u_i}^2 = c \cdot x_{ii}^2$.
2. При выборе из главного меню команды **Model\Other Linear Models\Heteroskedasticity Corrected...** (рисунок 12) веса наблюдений (w_i) определяются Gretl по формуле (4) и не вводятся пользователем. В качестве условия корректности применения рассматриваемого метода оценивания рассматривается получение нормального распределения остатков. Форму распределения можно проверить при помощи команды меню окна результатов моделирования **Tests\normality of residual.....**

$$w_i = \frac{1}{\sqrt{e^{\hat{y}_i}}}, \quad (4)$$

где $\tilde{y}_i$ - модельные значения зависимой переменной, полученные при использовании 1МНК; $e=2,718..$; w_i -вес i-го наблюдения.

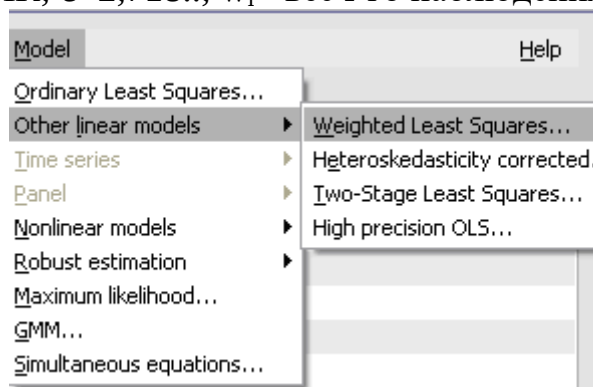


Рисунок 12 – Способы реализации ВМНК (WLS) в пакете Gretl

Обратимся к вышеприведённому примеру набора данных bwages.gdt для иллюстрации обоих способов оценивания регрессионной модели $wage = \alpha_0 + \alpha_1 \cdot educ + u$ методом ВМНК.

Этап 1. Введём новую переменную – вес w , которая будет определяться по формуле $w_i = \frac{1}{\sigma_{u_i}^2} = \frac{1}{3 \cdot educ_i^2}$.

Для этого выберем команду Define new variable из меню Variable (рисунок 13).

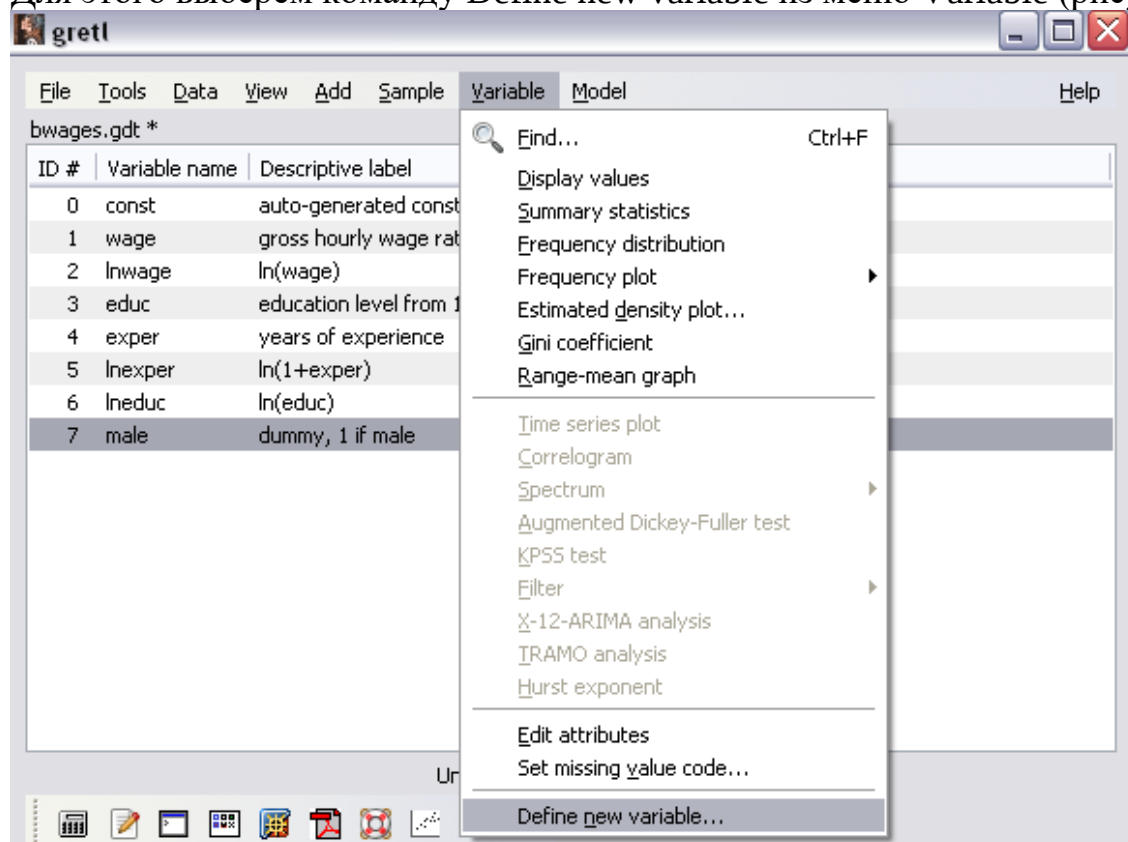


Рисунок 13 – Создание новой переменной

В открывшемся окне “add var” введём имя переменной *w* и нажмём ОК, в появившемся окне «edit data» не вводим значения переменной, а нажимаем кнопку «close» для создания переменной с пустыми значениями (рисунок 14).

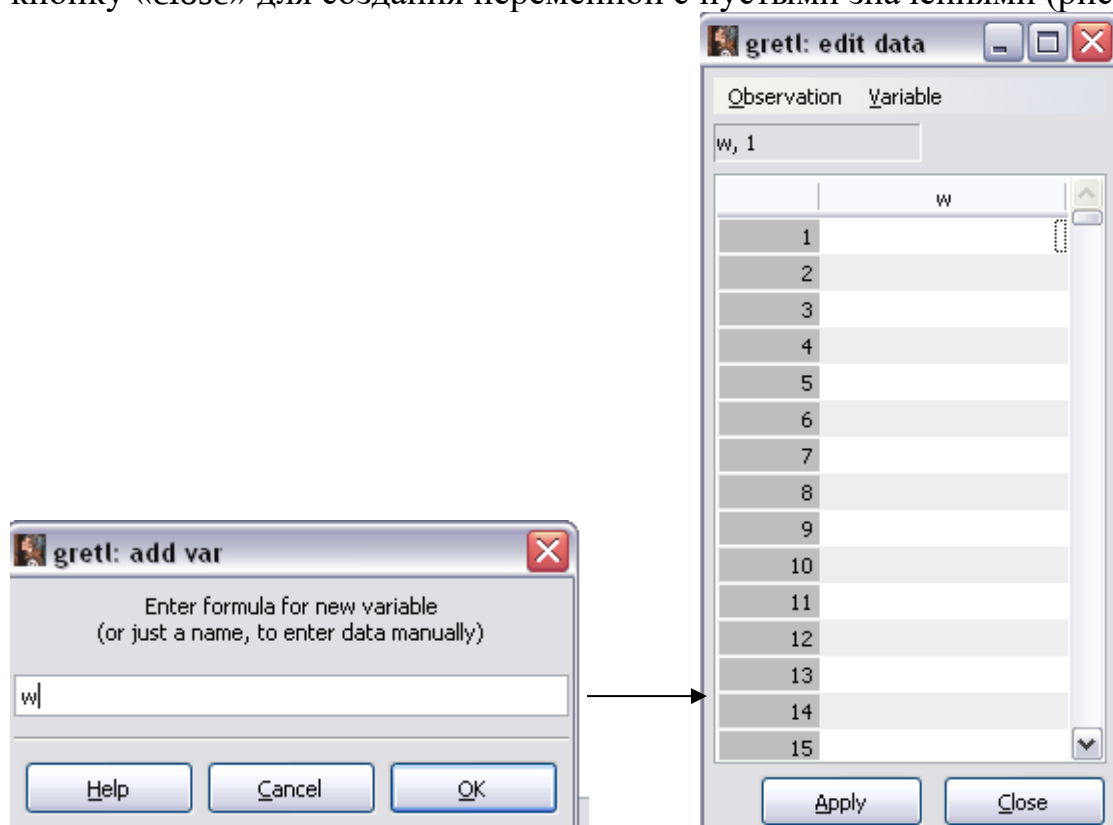


Рисунок 14 – Ввод имени переменной

Данная переменная *w* с пустыми значениями и без описания появится в списке открытого набора данных. Введём описание переменной, обратившись к команде **Variable\Edit Attributes**.

В открывшемся окне (рисунок 15) введём в поле “Description” описание «weights» (веса) и нажмём ОК.

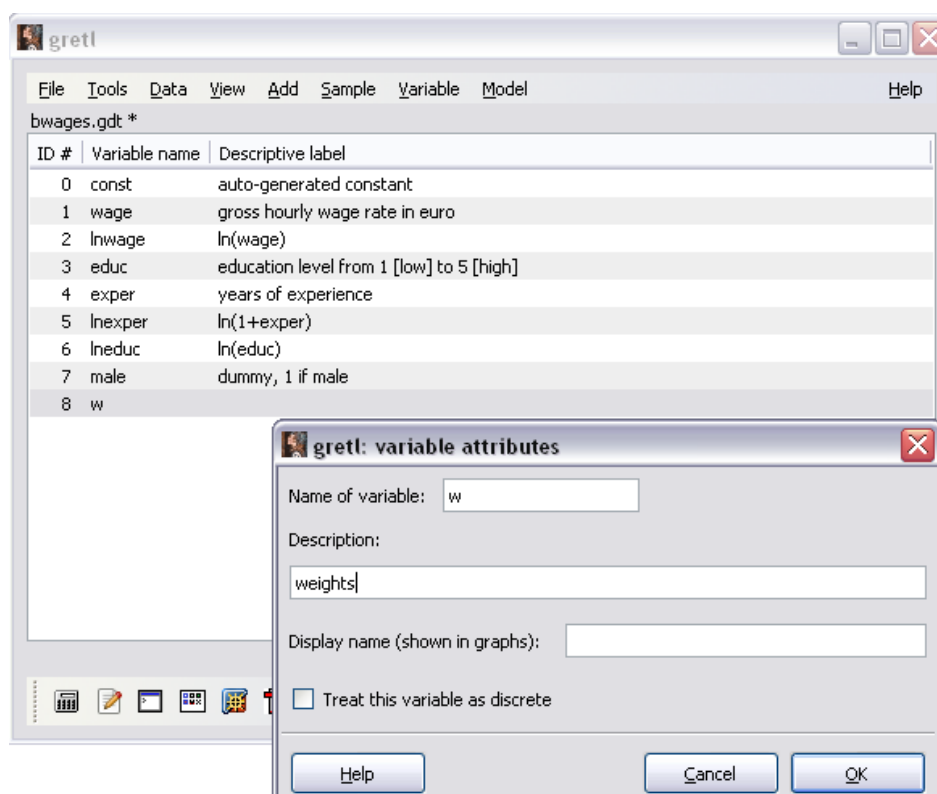
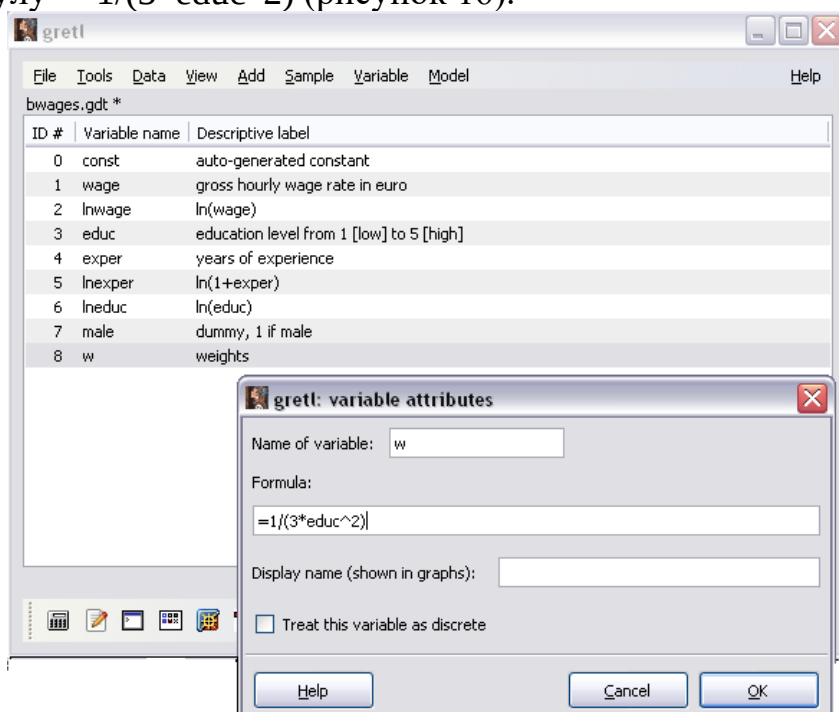
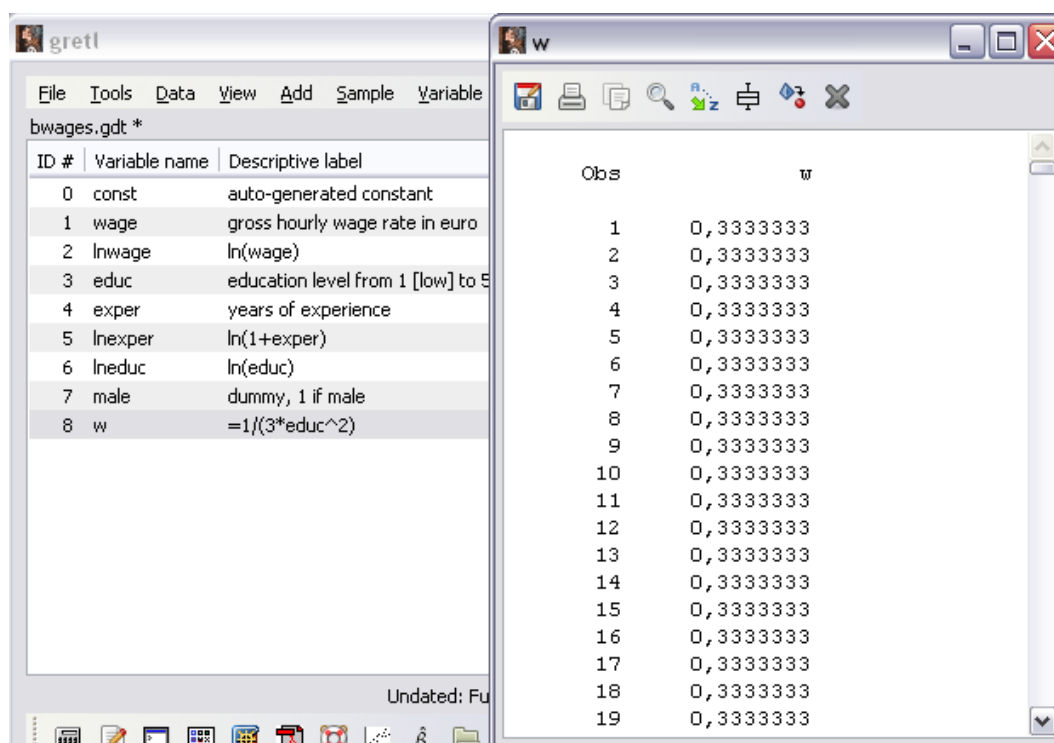


Рисунок 15 – Изменение атрибутов переменной

Введём формулу для переменной w , повторно обратившись к команде **Variable\Edit Attributes**. Теперь в поле “Description” вместо описания «weights» введём формулу $=1/(3*educ^2)$ (рисунок 16).

Рисунок 16 – Ввод формулы для определения переменной w

Этап 2. Просмотрим значения созданной переменной w , выбрав её в списке щелчком мыши и обратившись к команде **Data\Display Values** (или двойным щелчком мыши по её названию в списке) (рисунок 17).



The screenshot shows the Gretl interface with two windows. The main window displays a list of variables for the file 'bwages.gdt'. The variable 'w' is defined as $w = 1/(3 * educ^2)$. A secondary window titled 'w' displays the values of this variable for 19 observations, all of which are 0,3333333.

ID #	Variable name	Descriptive label
0	const	auto-generated constant
1	wage	gross hourly wage rate in euro
2	lnwage	ln(wage)
3	educ	education level from 1 [low] to 5
4	exper	years of experience
5	lnexper	ln(1+exper)
6	lneduc	ln(educ)
7	male	dummy, 1 if male
8	w	$=1/(3 * educ^2)$

Obs	w
1	0,3333333
2	0,3333333
3	0,3333333
4	0,3333333
5	0,3333333
6	0,3333333
7	0,3333333
8	0,3333333
9	0,3333333
10	0,3333333
11	0,3333333
12	0,3333333
13	0,3333333
14	0,3333333
15	0,3333333
16	0,3333333
17	0,3333333
18	0,3333333
19	0,3333333

Рисунок 17 – Просмотр значений переменной w

Этап 3. Для оценки модели $wage = \alpha_0 + \alpha_1 \cdot educ + u$ методом ВМНК (OLS) выполним команду **Model\Other Linear Models\Weighted Least Squares....**(рисунок 18) и в открывшемся окне спецификации модели (рисунок 19) выберем при помощи кнопок «Choose» и «Add» переменные: wage (з.пл) - зависимая переменная; w – веса наблюдений; educ (уровень образования) – независимая переменная.

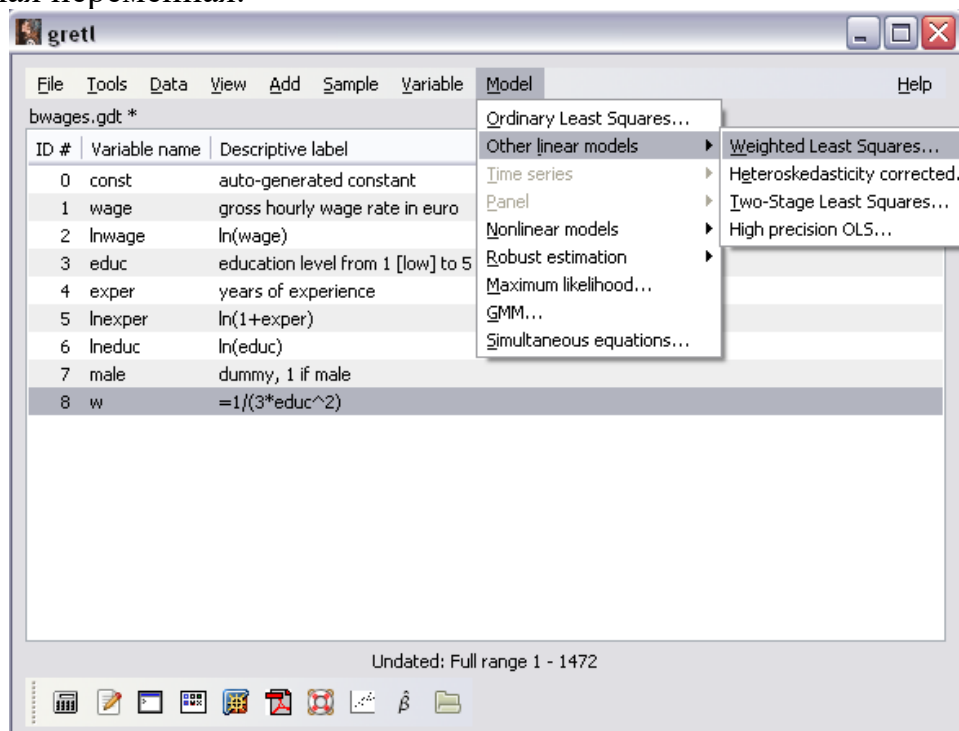


Рисунок 18 – Реализация ВМНК (OLS) в Gretl

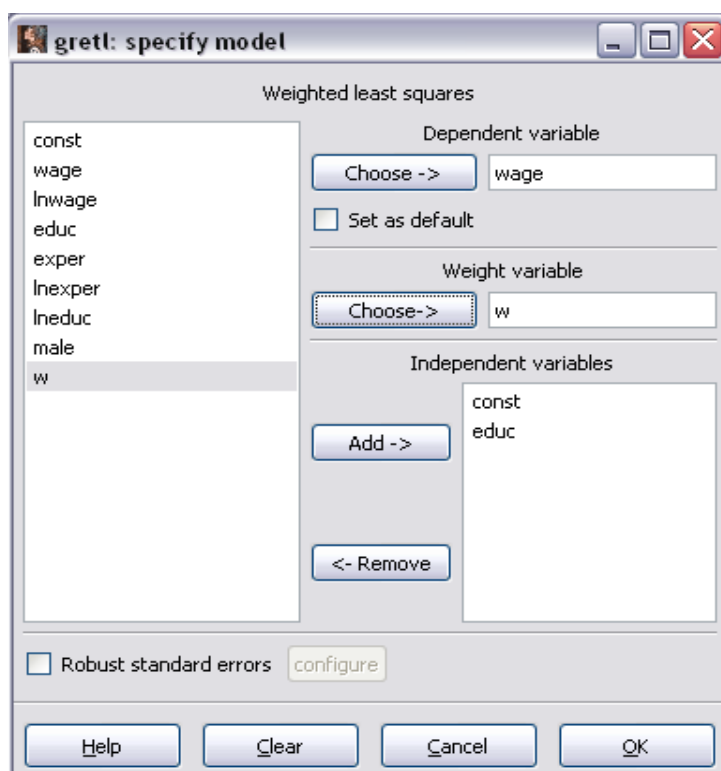


Рисунок 19 – Оценивание модели $wage = \alpha_0 + \alpha_1 \cdot educ + u$ при помощи ВМНК (OLS)

Результаты оценивания представлены в окне результатов моделирования (рисунок 20).

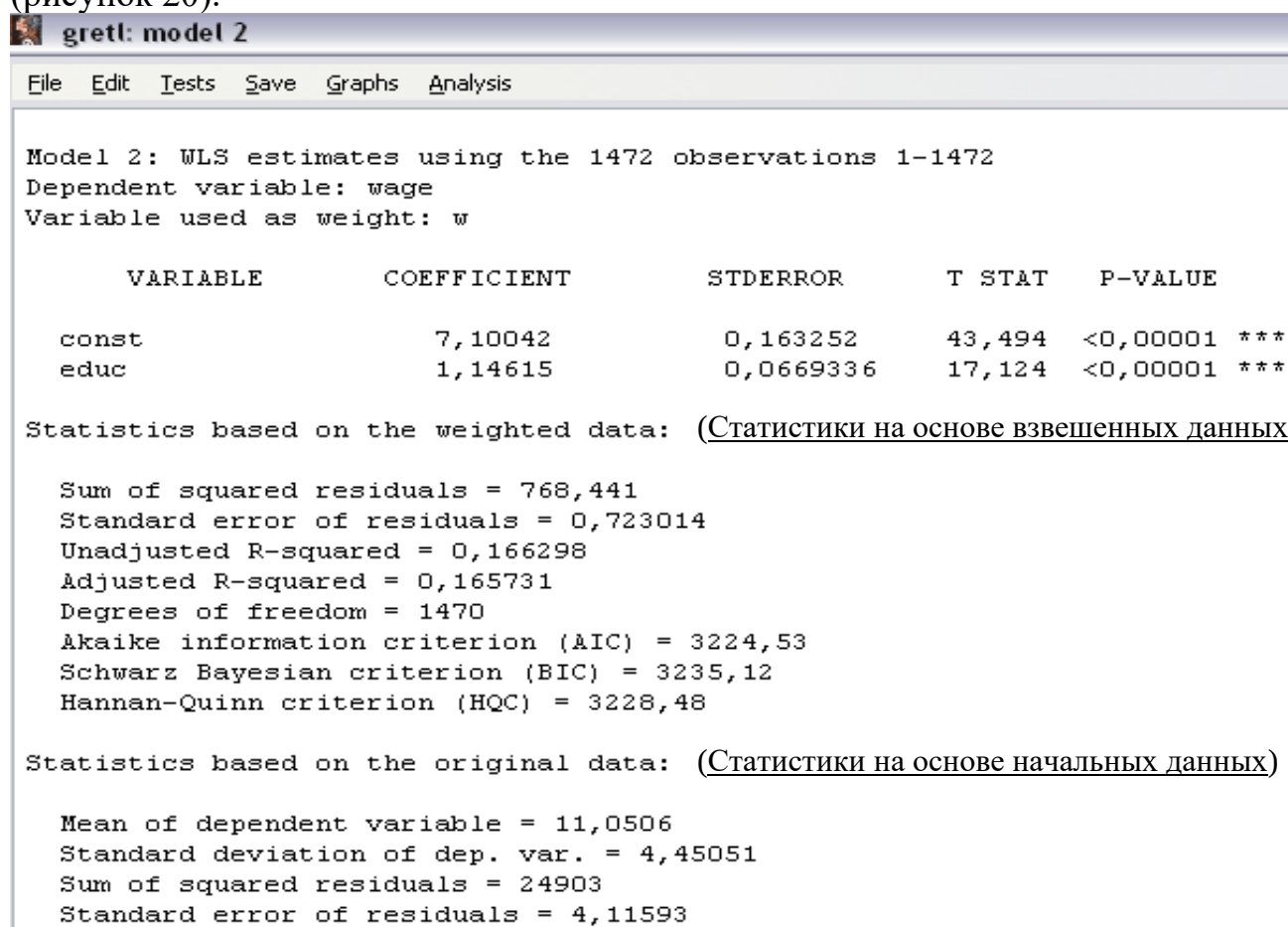


Рисунок 20 – Результаты оценивания модели $wage = \alpha_0 + \alpha_1 \cdot educ + u$ методом ВМНК (WLS)

Сравнивая модели, полученные методом ВМНК $wage = 7,10042 + 1,14615 \cdot educ$ (рисунок 20) и методом 1МНК $wage = 6,18513 + 1,44018 \cdot educ$ (рисунок 7) можно сделать вывод о том, что при использовании 1МНК оценка коэффициента α_1 была завышена, а оценки α_0 и T-STAT – занижены, но на правильности решения о принятии альтернативной гипотезы это не отразилось (значения P-VALUE не изменились).

Этап 4. Воспользуемся вторым способом оценки модели $wage = \alpha_0 + \alpha_1 \cdot educ + u$ методом ВМНК (OLS), выполнив команду **Model\Other Linear Models\Heteroskedasticity Corrected...** (рисунок 21).

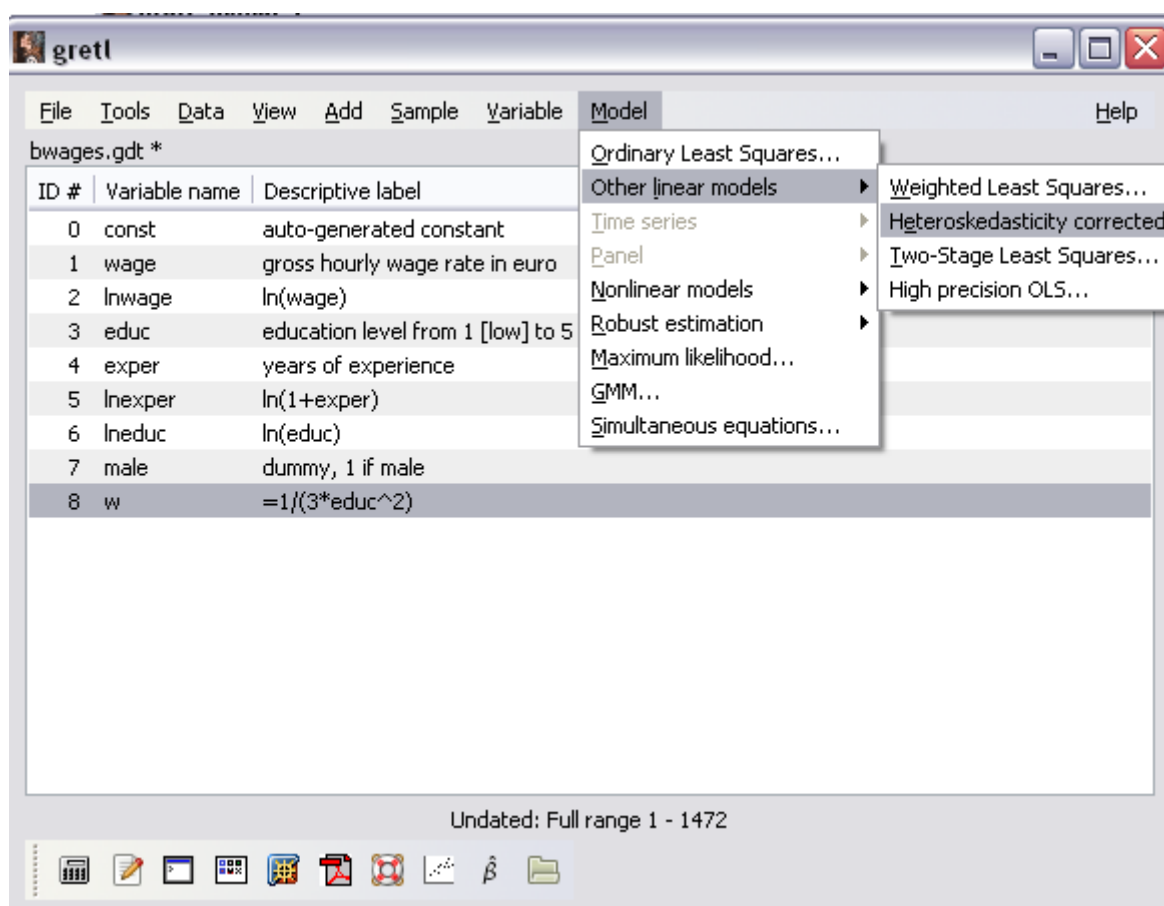


Рисунок 21 – Альтернативный способ оценки модели методом ВМНК (WLS)

Заполним окно спецификации также, как на предыдущем этапе за исключением ввода переменной – веса w , и нажмём кнопку ОК (рисунок 22).

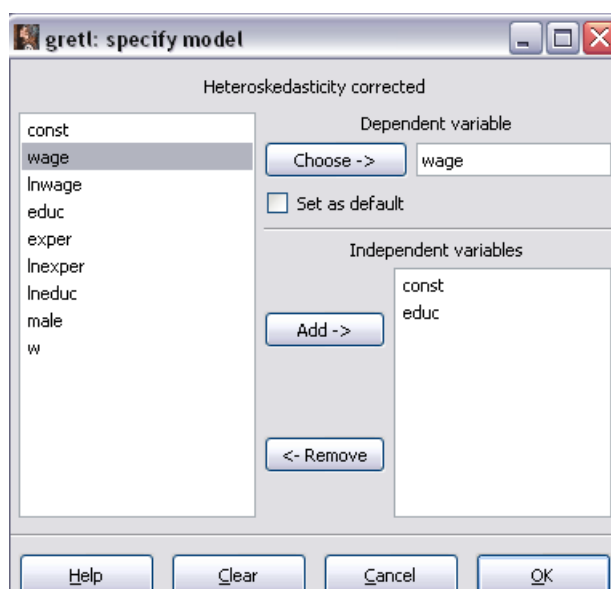


Рисунок 22 – Оценивание модели $wage = \alpha_0 + \alpha_1 \cdot educ + u$ при помощи ВМНК (OLS) с автоматическим определением весов наблюдений

Сравнивая полученную модель $wage = 6,74951 + 1,2509 \cdot educ$ (рисунок 23) с моделью, оцененную с применением 1МНК $wage = 6,18513 + 1,44018 \cdot educ$ (рисунок 7) можно отметить расхождения в значениях показателей регрессионной таблицы, но на правильности решения о принятии альтернативной гипотезы это также не отразилось (значения P-VALUE не изменились).

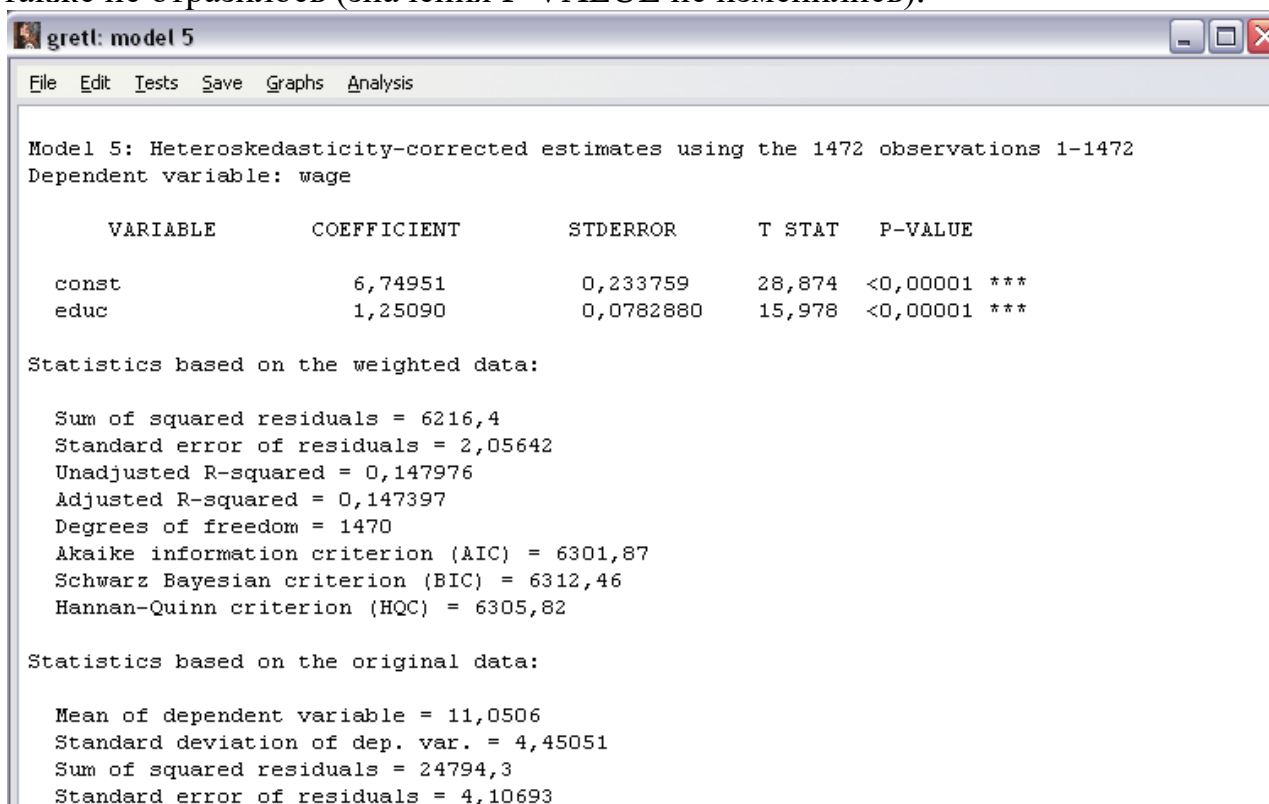
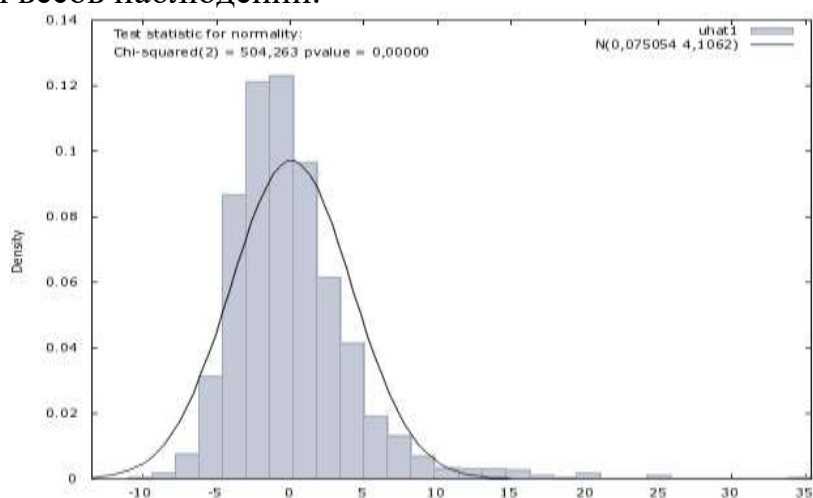


Рисунок 23 – Результаты оценивания модели $wage = \alpha_0 + \alpha_1 \cdot educ + u$ методом ВМНК (WLS) с автоматическим определением весов наблюдений

Проверим условие корректности применения рассматриваемого метода – получение нормального распределения остатков, - обратившись к команде **Tests\normality of residual...** меню окна модели (рисунок 23).

Рисунок 24 показывает, что данное условие не соблюдается $P\text{-value}=0,0000$, что меньше уровней значимости 1% и 5%, поэтому отвергаем нулевую гипотезу о нормальности распределения остатков.

Вывод: В результате исследования 1479 бельгийских семей выявлена прямая зависимость заработной платы (*wage*) от уровня образования (*educ*). Наилучшей моделью, описывающей рассматриваемую зависимость, является модель $wage = 7,10042 + 1,14615 \cdot educ$, полученная методом ВМНК (WLS) с заданием пользователем весов наблюдений.



Frequency distribution for uhat1, obs 1-1472
number of bins = 29, mean = 0,0750541, sd = 4,10624

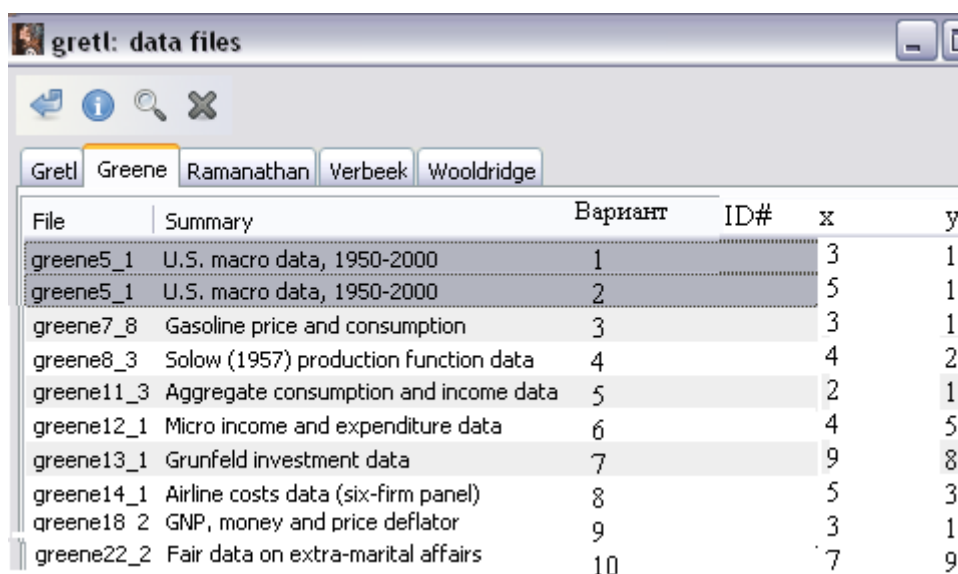
interval	midpt	frequency	rel.	cum.
< -9,3878	-10,187	1	0,07%	0,07%
-9,3878 - -7,7893	-8,5885	4	0,27%	0,34%
-7,7893 - -6,1908	-6,9900	18	1,22%	1,56%
-6,1908 - -4,5922	-5,3915	73	4,96%	6,52%
-4,5922 - -2,9937	-3,7930	204	13,86%	20,38%
-2,9937 - -1,3952	-2,1945	285	19,36%	39,74%
-1,3952 - 0,20332	-0,59594	289	19,63%	59,38%
0,20332 - 1,8018	1,0026	227	15,42%	74,80%
1,8018 - 3,4004	2,6011	144	9,78%	84,58%
3,4004 - 4,9989	4,1996	97	6,59%	91,17%
4,9989 - 6,5974	5,7981	45	3,06%	94,23%
6,5974 - 8,1959	7,3967	31	2,11%	96,33%
8,1959 - 9,7944	8,9952	16	1,09%	97,42%
9,7944 - 11,393	10,594	8	0,54%	97,96%
11,393 - 12,991	12,192	7	0,48%	98,44%
12,991 - 14,590	13,791	7	0,48%	98,91%
14,590 - 16,189	15,389	6	0,41%	99,32%
16,189 - 17,787	16,988	2	0,14%	99,46%
17,787 - 19,386	18,586	1	0,07%	99,52%
19,386 - 20,984	20,185	4	0,27%	99,80%
20,984 - 22,583	21,783	0	0,00%	99,80%
22,583 - 24,181	23,382	0	0,00%	99,80%
24,181 - 25,780	24,980	2	0,14%	99,93%
25,780 - 27,378	26,579	0	0,00%	99,93%
27,378 - 28,977	28,177	0	0,00%	99,93%
28,977 - 30,575	29,776	0	0,00%	99,93%
30,575 - 32,174	31,374	0	0,00%	99,93%
32,174 - 33,772	32,973	0	0,00%	99,93%
>= 33,772	34,571	1	0,07%	100,00%

Test for null hypothesis of normal distribution:
Chi-square(2) = 504,263 with p-value 0,00000

Рисунок 24 – Результаты теста на нормальность распределения остатков модели $wage = 6,74951 + 1,2509 \cdot educ$

4. ПОРЯДОК ВЫПОЛНЕНИЯ ЛАБОРАТОРНОЙ РАБОТЫ

1. Открыть набор данных File\Open Data\ Sample File (закладка Greene), Рисунок 25, в соответствии с номером варианта и общую текстовую информацию о нём.
2. Использовать 1МНК (OLS) для предварительной оценки параметров модели $y = \alpha_0 + \alpha_1 \cdot x + u$, оценить значимость параметров модели.
4. Определить, присутствует ли гетероскедастичность остатков модели.
5. В случае обнаружения гетероскедастичности провести оценку параметров модели методом ВМНК (WLS) двумя способами, оценить значимость параметров моделей.
6. Сравнить результаты, полученные с использованием методов 1МНК и ВМНК.



File	Summary	Вариант	ID#	x	y
greene5_1	U.S. macro data, 1950-2000	1		3	1
greene5_1	U.S. macro data, 1950-2000	2		5	1
greene7_8	Gasoline price and consumption	3		3	1
greene8_3	Solow (1957) production function data	4		4	2
greene11_3	Aggregate consumption and income data	5		2	1
greene12_1	Micro income and expenditure data	6		4	5
greene13_1	Grunfeld investment data	7		9	8
greene14_1	Airline costs data (six-firm panel)	8		5	3
greene18_2	GNP, money and price deflator	9		3	1
greene22_2	Fair data on extra-marital affairs	10		7	9

Рисунок 25 – Варианты заданий: название файла и номер варианта

5. СОДЕРЖАНИЕ ОТЧЕТА О ВЫПОЛНЕНИИ ЛАБОРАТОРНОЙ РАБОТЫ

- 1) Название и цель работы.
- 2) Постановка задачи.
- 3) Этапы выполнения задачи в Gretl.
- 4) Выводы.

РЕАЛИЗАЦИЯ МЕТОДА ГЛАВНЫХ КОМПОНЕНТ В СРЕДЕ GRETL 1.7.1

1. ЦЕЛЬ РАБОТЫ

Целью данной работы является получение практических навыков обработки финансово-экономической информации с использованием алгоритмов пакета Gretl 1.7.1, обеспечивающих выполнение метода главных компонент.

2. ТЕОРЕТИЧЕСКИЕ СВЕДЕНИЯ

Метод главных компонент был предложен Пирсоном в 1901 году и затем вновь открыт и детально разработан Хоттелингом (1933г.) Данный метод применяется, например, для сжатия объемов хранимой информации и упрощения её интерпретации или сравнения многомерных исследуемых объектов, позволяя снизить размерность исходного признакового пространства $x_1, \dots, x_p$ (x_i - исходный признак) и перейти к новым агрегированным признакам $y_1, \dots, y_{p'}$ (y_j -главная компонента), $p' < p$. При этом новые показатели $y_1, \dots, y_{p'}$ представляют собой линейные комбинации исходных $x_1, \dots, x_p$, коррелированных между собой, формула (1).

$$\begin{cases} y_j(x) = w_{1j} \left(\frac{x_1 - \bar{x}_1}{\sigma_1} \right) + \dots + w_{pj} \left(\frac{x_p - \bar{x}_p}{\sigma_p} \right); \\ \sum_{i=1}^p w_{ij}^2 = 1 & (j = 1 \dots p); \\ \sum_{i=1}^p w_{ij} w_{ik} = 0 & (j, k = 1 \dots p, j \neq k); \end{cases} \quad (1)$$

где $\bar{x}_j$ и σ_j — среднее арифметическое и среднеквадратическое отклонение признака x_i ;

w_{ij} - коэффициенты главных компонент, максимизирующие дисперсию y_j , которые находятся из уравнения $(S - \lambda E) \cdot \vec{w} = 0$, имеющее решение, если $|S - \lambda E| = 0$, где S - ковариационная (или корреляционная) матрица;

λ_i - собственные числа матрицы S - равны дисперсиям проекций множества объектов на оси главных компонент (или диагоналям эллипса рассеяния), рисунок 1. На рисунке 1 “р” точек сосредоточены в трехмерном пространстве двух систем координат: переменных x_1, x_2, x_3 и главных компонент y_1, y_2, y_3 . При этом оси OY_1, OY_2, OY_3 проходят через центр тяжести эллипсоида рассеяния.

Традиционный алгоритм расчета главных компонент включает переход от исходной матрицы наблюдений X к ковариационной (или корреляционной) матрице S между исходными признаками $x_1, \dots, x_p$, далее к расчету собственных

чисел λ_i . Основываясь на наибольших собственных числах, наилучшим образом объясняющих исходное пространство признаков, производится переход к главным компонентам путём определения их коэффициентов $w_j = (w_{1j}, \dots, w_{pj})'$, максимизирующих дисперсию проекций множества объектов на оси главных компонент. Таким образом, выбираются только те главные компоненты, изменчивость которых покрывает большую часть изменчивости $x_1, \dots, x_p$.

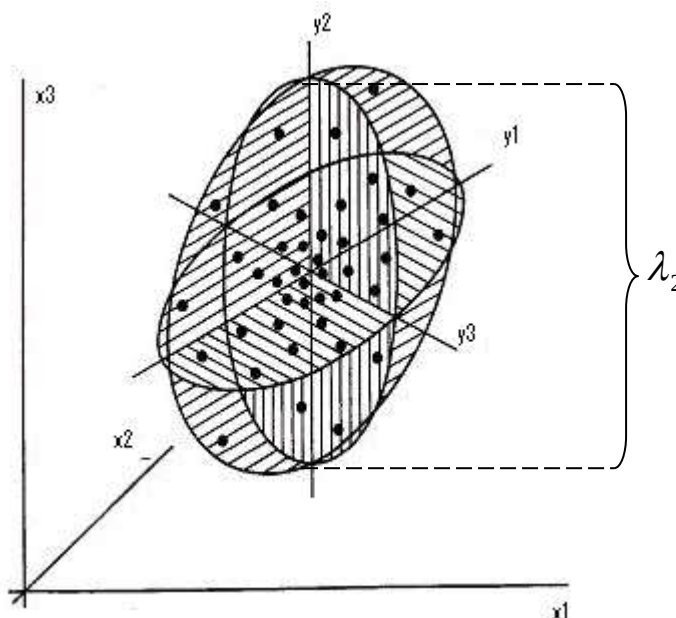


Рисунок 1 – Диаграмма рассеяния x_1, x_2, x_3

3. ПРИМЕР ПРАКТИЧЕСКОЙ РЕАЛИЗАЦИИ МЕТОДА ГЛАВНЫХ КОМПОНЕНТ С ИСПОЛЬЗОВАНИЕМ СИСТЕМЫ GRETL

Рассмотрим задачу, стоящую перед международным рейтинговым агентством, например, Standard&Poor's или Moody's Investors Service по формированию рейтинга устойчивого развития предприятия (обобщённого показателя устойчивости предприятия для их сравнения) на примере компаний Ford Motor Company и General Motors.

Пусть устойчивое предприятие – это компания, существующая более 100 лет, а к составляющим устойчивого развития относятся:

1. Финансовая составляющая, y_1 ;
2. Продукт и потребители, y_2 ;
3. Качество отношений с персоналом, y_3 ;
4. Окружающая среда и безопасность, y_4 .

Тогда предприятие – это многомерный объект, характеризующийся вектором функций $a = (y_1, y_2, y_3, y_4)'$, где каждый элемент вектора y_i – главная компонента (с наибольшим вкладом), определяемая набором исходных признаков $x_1, \dots, x_p$. Тогда собственное число λ_i , соответствующее каждому y_i ($i=1..4$), является рейтингом предприятия в соответствующей области (1,2,3 или 4), а средняя арифметическая значений собственных чисел (соответствующих y_1 ,

у₂, у₃, у₄) составит обобщённый показатель (рейтинг) устойчивого развития предприятия в целом. Сравним данные рейтинги для компаний Ford Motor Company и General Motors с использованием пакета GRETl.

3.1. Исходная информация

Автомобильная корпорация Ford Motor Company (ОАО) www.ford.com была основана в 1903 году Генри Фордом и одна из немногих пережила Великую Депрессию, постоянно находясь под контролем одной семьи в течение более 100 лет, стала одной из крупнейших и наиболее прибыльных автомобилестроительных компаний мира. Компания получила известность как первая в мире применившая классический автосборочный конвейер. Штаб-квартира Ford находится в Мичигане (США), компания является производителем и дистрибьютором автомобильной техники на 200 рынках в шести странах мира, имеет более 280000 работников и более 100 заводов по всему миру. Компания выпускает широкий спектр легковых и коммерческих автомобилей под марками Ford, Jaguar, Land Rover, Lincoln, Mercury, Volvo, Aston Martin, Mazda. Ford — третий по объёму выпуска автопроизводитель в мире после General Motors и Toyota. В настоящее время президентом компании и исполнительным директором является Alan Mulally, а председателем совета директоров - William Clay Ford. Основными конкурентами предприятия являются General Motors, Toyota Motor Corp, DaimlerChrysler, PSA/Peugeot-Citroën, Hyundai Motor, Honda Motor Co., Nissan Motor Co., Renault SA. Выручка компании в 2007 году составила \$172,455 млрд, чистый убыток — \$2,7 млрд.

General Motors (GM) (ОАО). <http://www.gm.com> — один из основных конкурентов Ford Motor Company, основана в 1908 году в Мичигане (США). В настоящее время имеет около 284000 работников, с 1931 года является крупнейшей в мире (по объёму выпуска) автомобилестроительной компанией с заводами в более 33 странах мира. Председатель совета директоров и главный управляющий — Рик Вагонер (Rick Wagoner). Производимые автомобильные марки: Buick, Cadillac, Chevrolet, Hummer, Opel, Pontiac, и т.д. В 2007 году чистый убыток составил \$38,7 млрд, а выручка составила \$181 млрд.

Определим исходные признаки $x_1, \dots, x_p$ для каждой составляющей устойчивого развития предприятия:

1. Финансовая составляющая (у₁):

X₁₁ - Net income, bln\$ (Чистая прибыль, млрд. \$);

X₂₁ - Sales&Revenues, bln\$ (Выручка, млрд.\$);

X₃₁ - Stockholders' Equity, bln\$ (Акционерный капитал, млрд.\$)

X₄₁ - Gross cash, bln\$ (Валовая наличность, млрд.\$)

X₅₁ - Cash dividends, \$ (Наличные дивиденды, \$)

X₆₁ - Total assets, bln\$ (Совокупные активы, млрд.\$)

X₇₁ - Shareholders return, % (Доход акционеров, %)

2. Продукт и потребители (у₂):

X₁₂ - Market share US, % (Доля рынка в США, %);

X₂₂ - Market share Europe (Доля рынка в США, %);

X₃₂ - Customer satisfaction, % (Удовлетворённость потребителя, %)

X₄₂ - Customer Loyalty, US % (Верность потребителя, США %)

X₅₂ - Vehicles sold, units (Объём продаж, шт.)

3. Качество отношений с персоналом (y₃):

X₁₃ - Personnel full-time (Численность персонала на полной ставке, чел.);

X₂₃ - Employee satisfaction, % (Удовлетворённость работников, %);

X₃₃ - Laybor cost per hour, \$ (З.пл. в час, \$)

4. Окружающая среда и безопасность (y₄):

X₁₄ - U.S. Corporate Average Fuel Economy (Экономия топлива, США, миль на галлон);

X₂₄ - Energy efficiency index, % (Индекс эффективности энергопотребления, %);

X₃₄ - Charitable contributions, mln \$ (Благотворительные взносы, млн. \$);

X₄₄ - Global Fatalities (Число смертельных случаев на производстве).

Фактические значения вышеперечисленных показателей для компании Ford Motor Company представлены в таблицах 1-4 соответственно, источником данных являются годовые отчёты по устойчивому развитию www.ford.com/go/sustainability и финансовые отчёты <http://www.ford.com/about-ford/investor-relations/company-reports> предприятия, размещённые на его веб-сайте. Фрагменты данных отчётов представлены в приложении А. Знак «*» обозначает отсутствие данных и должно быть удалено перед обработкой в Gretl

Таблица 1 – Показатели финансовой составляющей компании Ford

Year	Net income, bln\$	Sales& Revenues, bln\$	Stockholders equity, bln\$	Gross cash, bln\$	Cash dividends, \$	Total assets, bln\$	Shareholders return, %
2002	0.9	167	7.633	25.3	0.4	281.851	-39
2003	0.2	166.095	13.459	25.9	0.4	303.361	79
2004	3.038	172.316	17.437	23.6	0.4	299.686	-6
2005	1.44	176.896	13.442	25.1	0.4	275.936	-45
2006	-12.613	160.123	3.465	33.9	0.25	290.217	1

Таблица 2 – Показатели составляющей «Продукт и потребители» компании Ford

Year	Market share US, %	Market share Eur, %	Customer satisfaction, %	Customer loyalty US, %	Vehicles sold, units
2002	21.1	10.9	72	48.5	6973000
2003	20.5	10.7	73	49.9	6720000
2004	19.3	10.9	74	47.5	6842000
2005	18.2	10.8	73	45.2	6767000
2006	17.1	10.6	74	43.3	6597000

Таблица 3 - Показатели составляющей «Качество отношений с персоналом» компании Ford

Year	Personnel full-time	Employee satisfaction, %	Laybor cost per hour, \$
2002	350321	59	52.6
2003	*	58	61.4
2004	*	61	62.9

2005	300000	62	64.9
2006	283000	62	70.5

Таблица 4 - Показатели составляющей «Окружающая среда и безопасность» компании Ford

Year	Fuel economy, US	Energy efficiency index, %	Charitable contributions, mln\$	Global Fatalities
2002	23.2	89.7	84	1
2003	23.6	91.7	78	3
2004	22.8	87.8	78	2
2005	24.1	83.4	80	2
2006	23.8	78.4	58	6

Создадим на рабочем столе файл Excel с исходными данными dataFord.xls, перенеся каждую из вышеприведённых таблиц 1-4 на отдельный лист созданной рабочей книги Excel как показано на рисунке 2.

	A	B	C	D	E	F	G	H
	Year	Net income, bln\$	Sales & Revenues, bln\$	Stockholders equity, bln\$	Gross cash, bln\$	Cash dividends, \$	Total assets, bln\$	Shareholders return, %
1	Year							
2	2002	0.9	167	7.633	25.3	0.4	281.851	-39
3	2003	0.2	166.095	13.459	25.9	0.4	303.361	79
4	2004	3.038	172.316	17.437	23.6	0.4	299.686	-6
5	2005	1.44	176.896	13.442	25.1	0.4	275.936	-45
6	2006	-12.613	160.123	3.465	33.9	0.25	290.217	1
7								
8								
32								

Рисунок 2 - Файл исходных данных dataFord.xls

Импортируем созданный файл в пакет Gretl, выбрав пункт меню **Open data\Import\Excel** и файл dataFord.xls на рабочем столе (рисунок 3). Затем в открывшемся окне нажмём кнопку ОК, подтвердив заданные по умолчанию настройки начала импорта с первой ячейки A1 первого листа Finance файла dataFord.xls (рисунок 4). Получим набор данных из семи переменных $x_{11} \dots x_{71}$ таблицы 1.

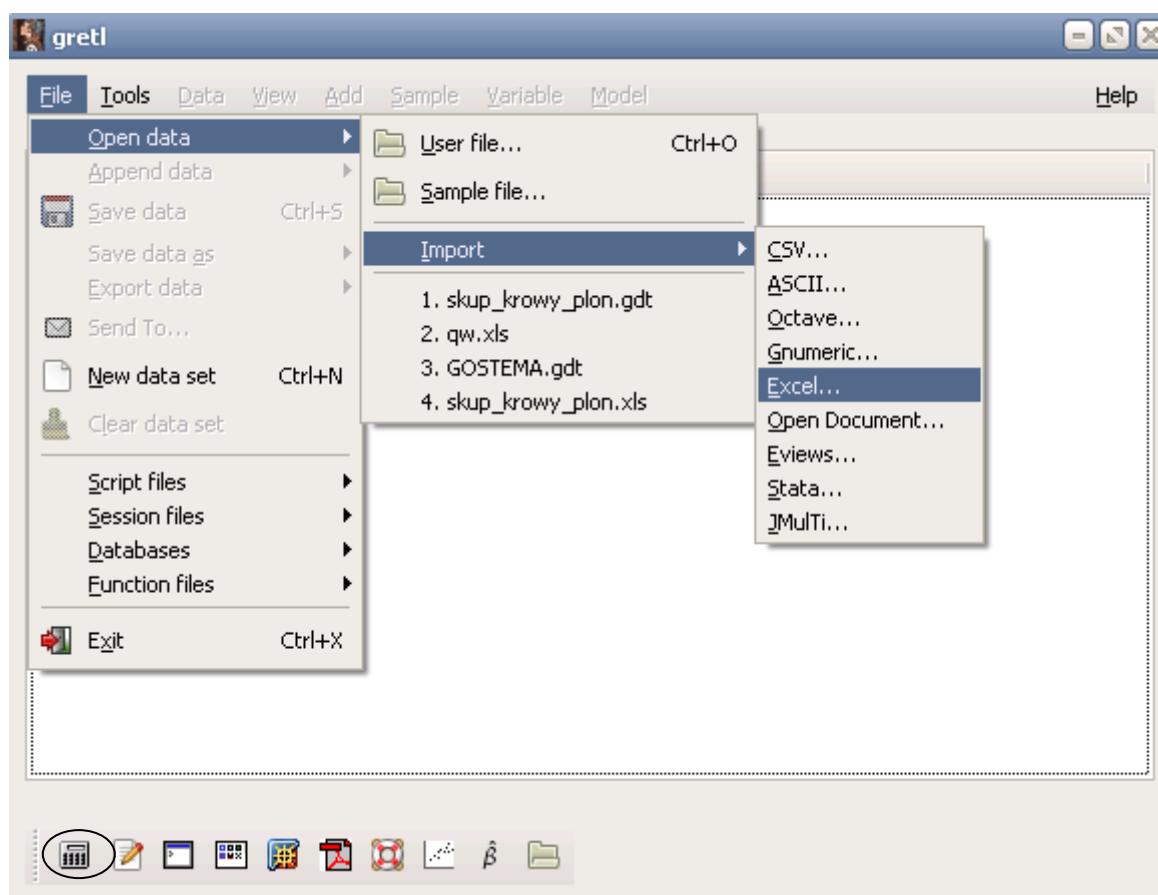


Рисунок 3 – Импорт исходных данных из файла dataFord.xls

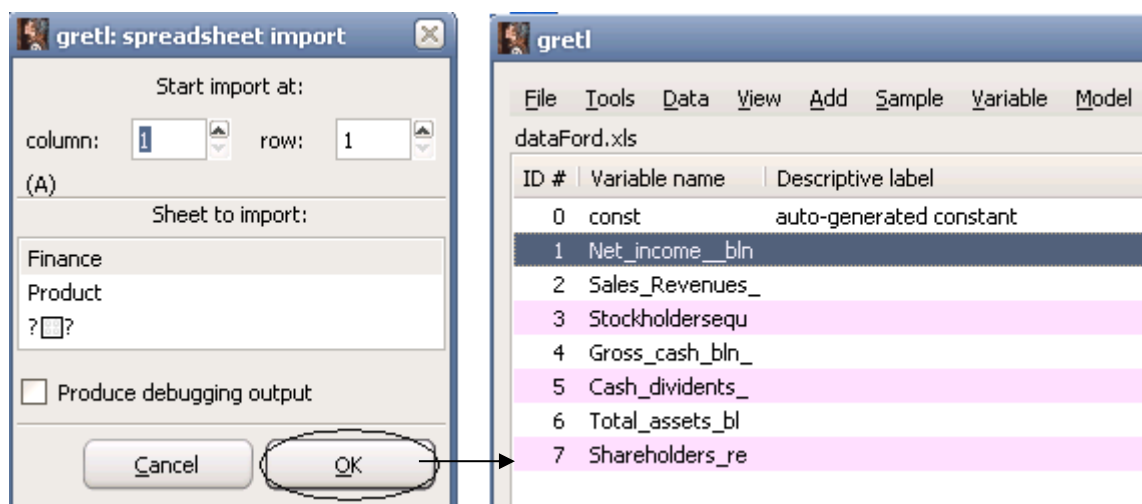


Рисунок 4 – Построение набора данных Ford.gdt

Для каждой переменной введём описание «Finance x_{i1} », $i=1\ldots 7$. Выберем пункт меню **Variable>Edit Attributes** и заполним поле Description (рисунок 5).

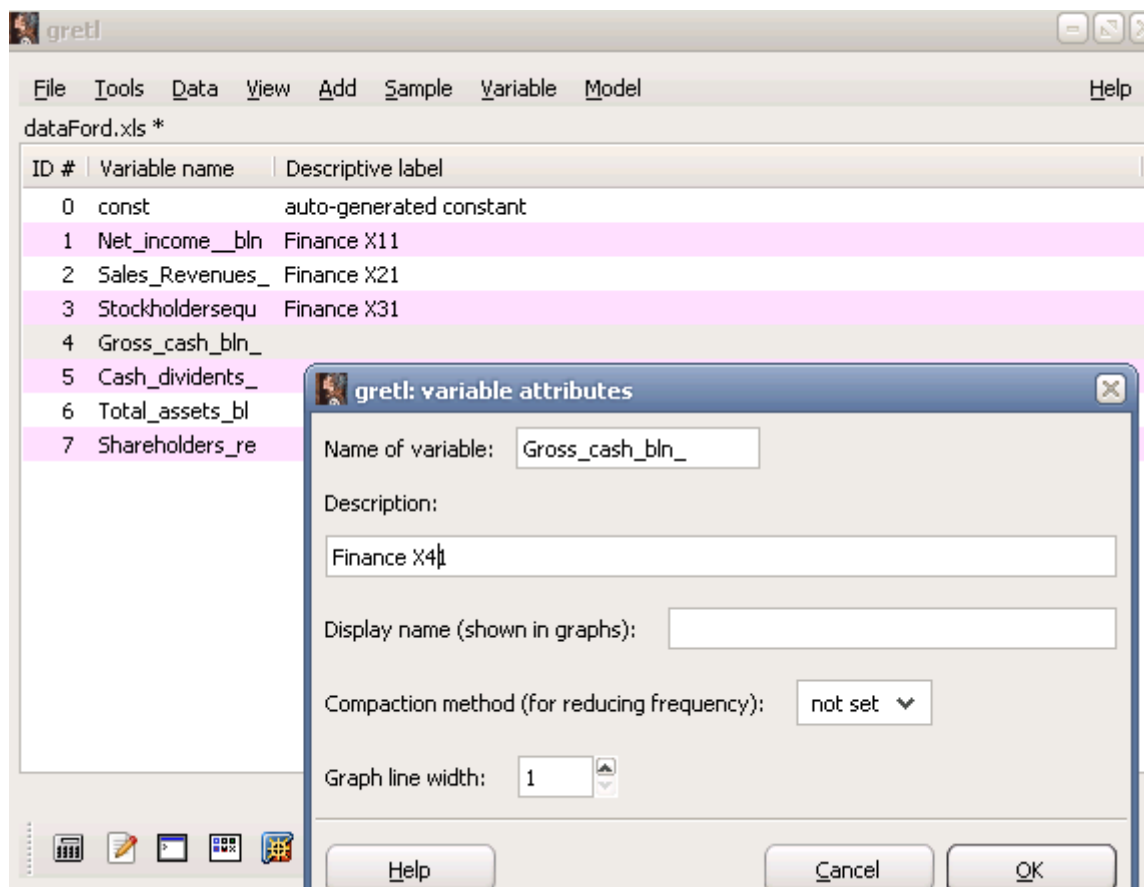


Рисунок 5 – Ввод описания переменных

Добавим в созданный набор данных следующий лист “Product” файла dataFord.xls, выбрав пункт меню **File\Append Data\Excel** и соответствующий файл (рисунок 6). В открывшемся окне выберем лист «Product» на нажмём кнопку ОК. В результате данные таблицы 2 будут добавлены к созданному набору данных Gretl (рисунок 7).

Аналогичным описанному выше способу добавим описания переменных «Product x_{i2} », $i=1...5$. Повторим соответствующие действия добавления для листов “Personnel” и “Ecology”.

Сохраним полученный набор данных, выбрав пункт меню **File\Save data** и введя имя файла Ford.gdt.

Откроем введенные данные в режиме просмотра и редактирования, выбрав пункт меню **View\Icon View** и дважды щёлкнув иконку **Data Set**.

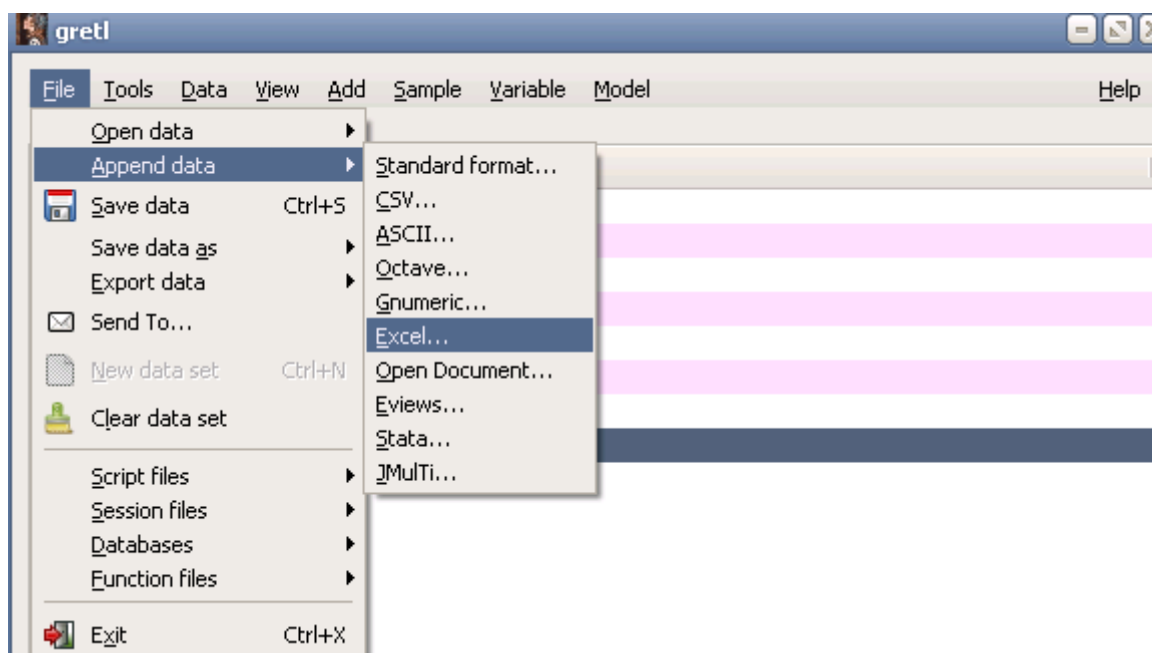


Рисунок 6 – Добавление дополнительных данных в созданный набор

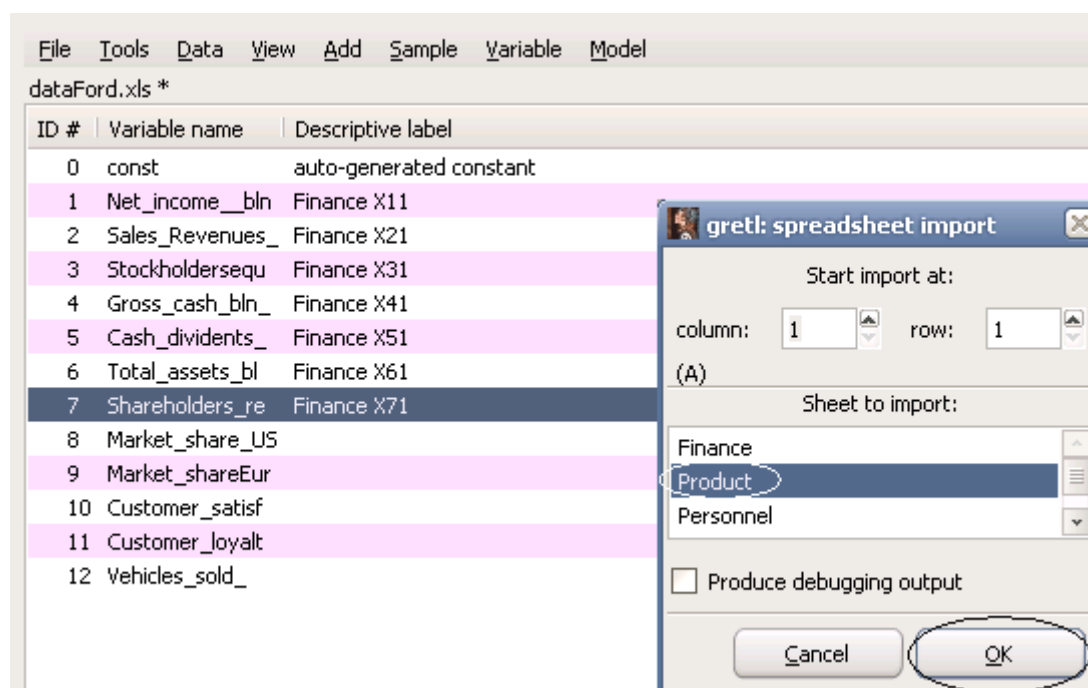


Рисунок 7 – Добавление данных таблицы 2 к созданному набору

3.2. Построение главных компонент

Поскольку метод главных компонент основан на коррелированности исходных признаков, перед построением главных компонент необходимо проверить наличие корреляции между $x_1 \dots x_p$ для каждой составляющей устойчивого развития. Построим матрицу корреляций, используя соответствующие функции пакета Gretl. Коэффициенты линейной корреляции Пирсона рассчитываются при помощи команды главного меню **View\Correlation**

Matrix. Выберем данный пункт меню (шаг 1.), в открывшемся окне при помощи кнопки Select выберем переменные $x_{11} \dots x_{71}$ (шаг 2.) и нажмём кнопку ОК (шаг 3.), как показано на рисунке 8.

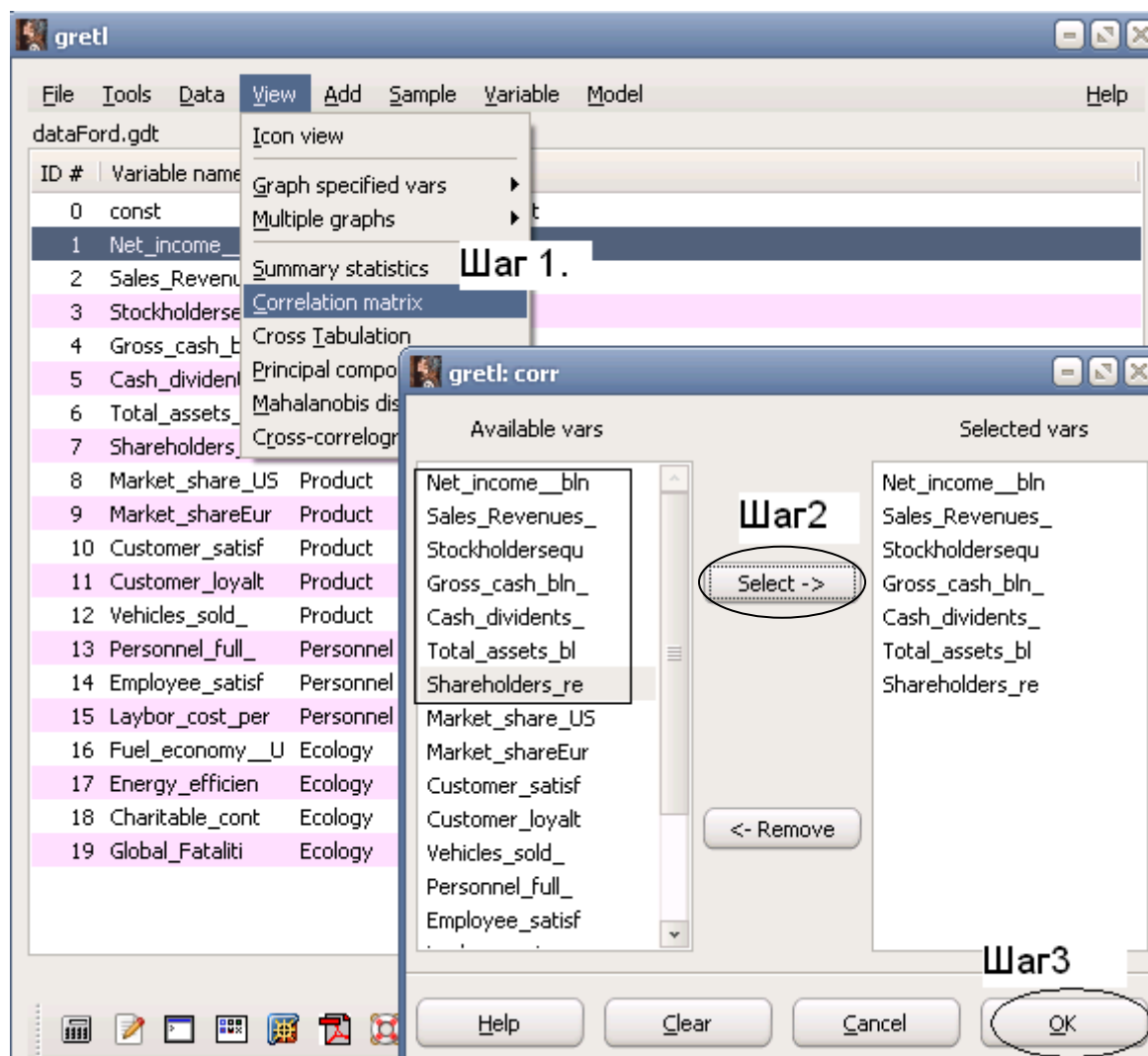


Рисунок 8 – Определение коэффициентов корреляции $x_{11} \dots x_{71}$

Оценив коэффициенты корреляции, отражённые в появившемся окне (рисунок 9), отметим, что большинство признаков коррелированы между собой за исключением:

Total assets (Совокупные активы)- Net Income (Чистая прибыль);
 Total assets (Совокупные активы)-Gross Cash (Валовая наличность);
 Total assets (Совокупные активы)- Cash Dividends (Наличные дивиденды);
 Shareholders return (Доход акционеров) - Net Income (Чистая прибыль);
 Shareholders return (Доход акционеров) - Gross Cash (Валовая наличность);
 Shareholders return (Доход акционеров) - Cash Dividends (Наличные дивиденды);
 для которых коэффициент корреляции примерно равен 0.

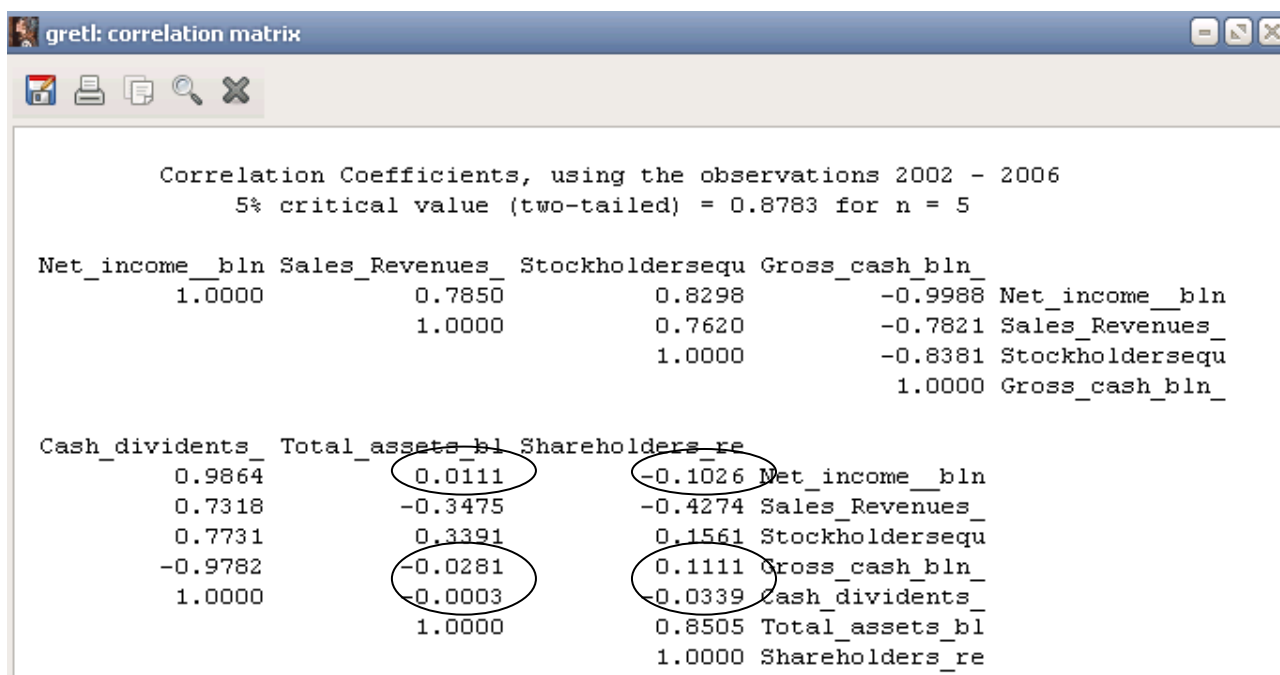


Рисунок 9 – Значения коэффициентов корреляции между переменными $X_{11} \dots X_{71}$

Т.о. необходимо исключить из рассмотрения переменные Total assets и Shareholders Return, тогда повторное обращение к функции **View\Correlation Matrix** даст отличные от нуля коэффициенты корреляции (рисунок 10).

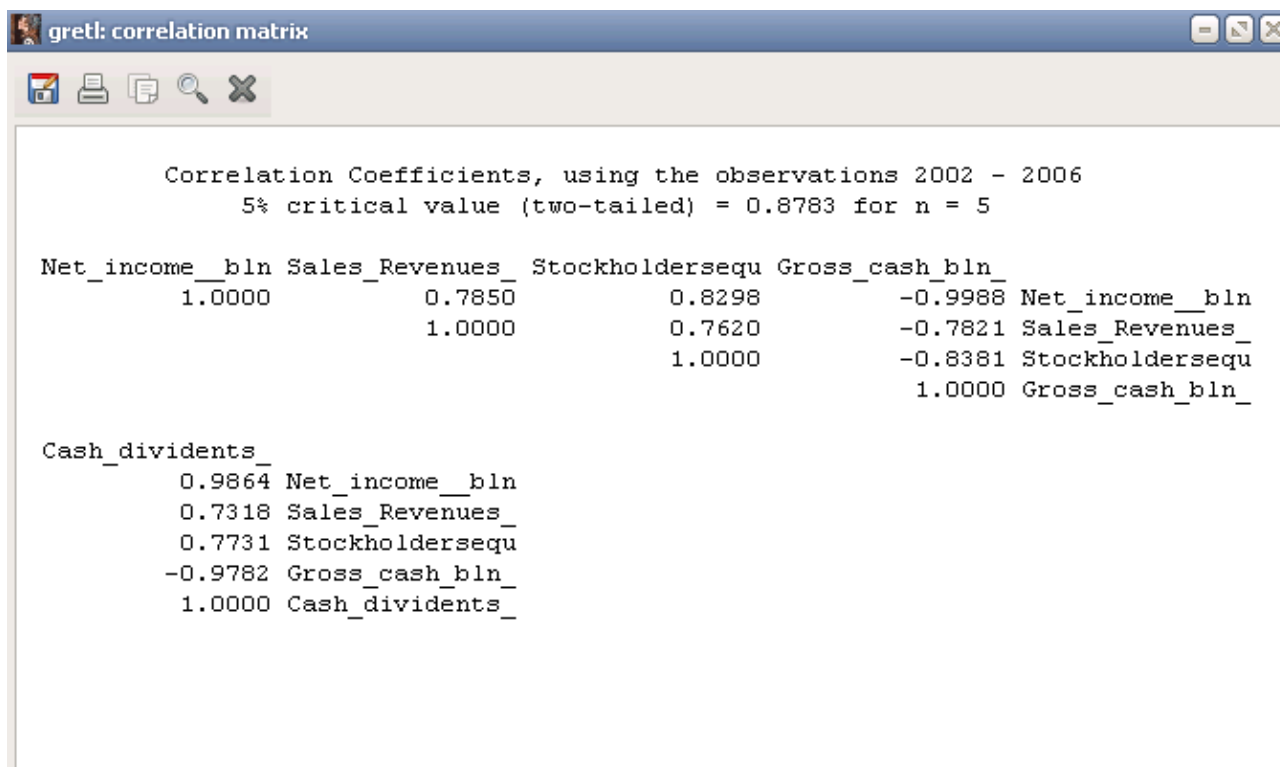


Рисунок 10 - Расчёт коэффициентов корреляции между переменными $X_{11} \dots X_{51}$

Построим главные компоненты по сокращённому пространству исходных признаков $x_{11} \dots x_{15}$, выбрав пункт меню **View\Principal Components** и выбрав при помощи кнопки **Select** соответствующие переменные $x_{11} \dots x_{15}$ (рисунок 11).

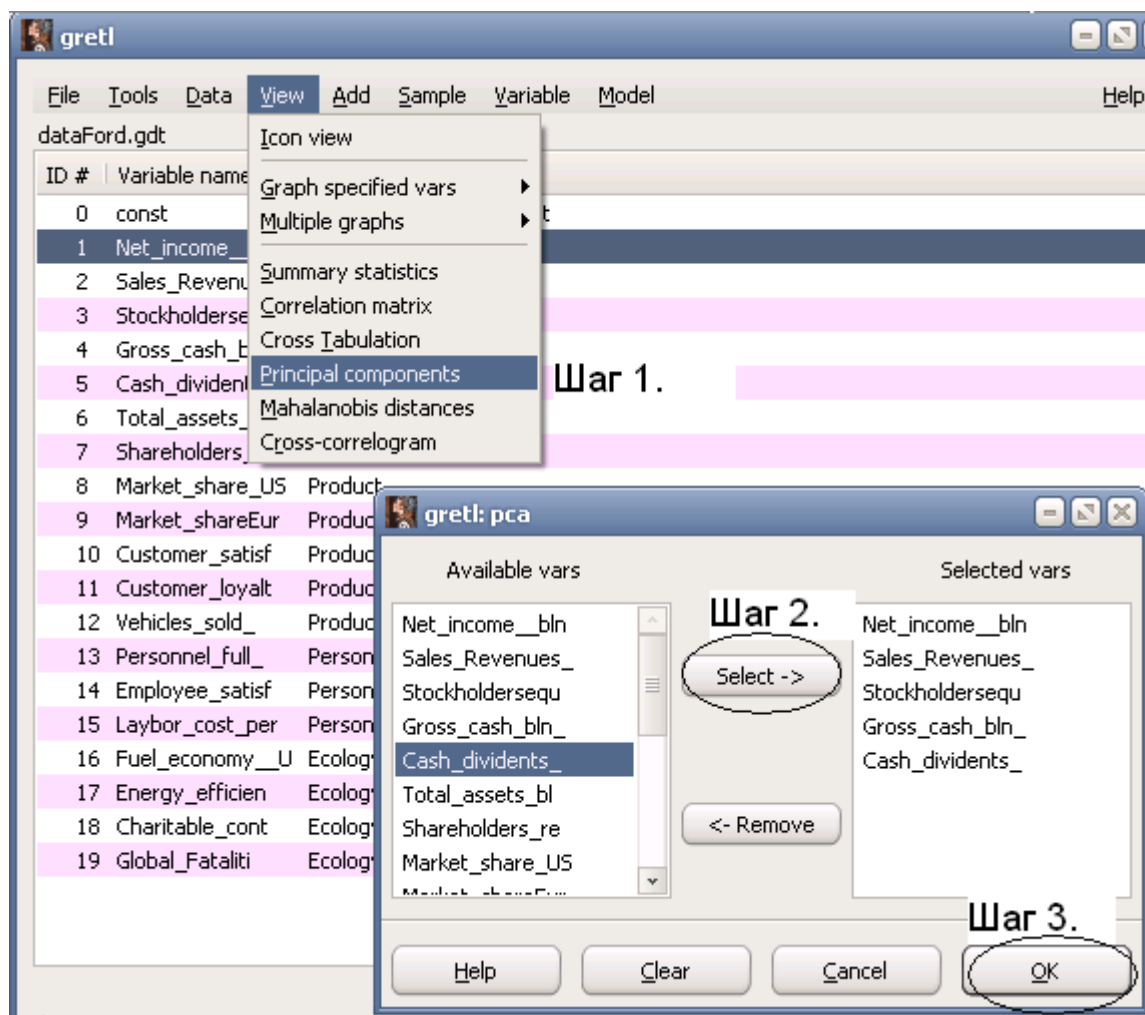


Рисунок 11- Построение главных компонент для финансовой составляющей

Полученные результаты моделирования представлены на рисунке 12. В таблице 5 приведены показатели и атрибуты окна результатов моделирования (рисунок 12).

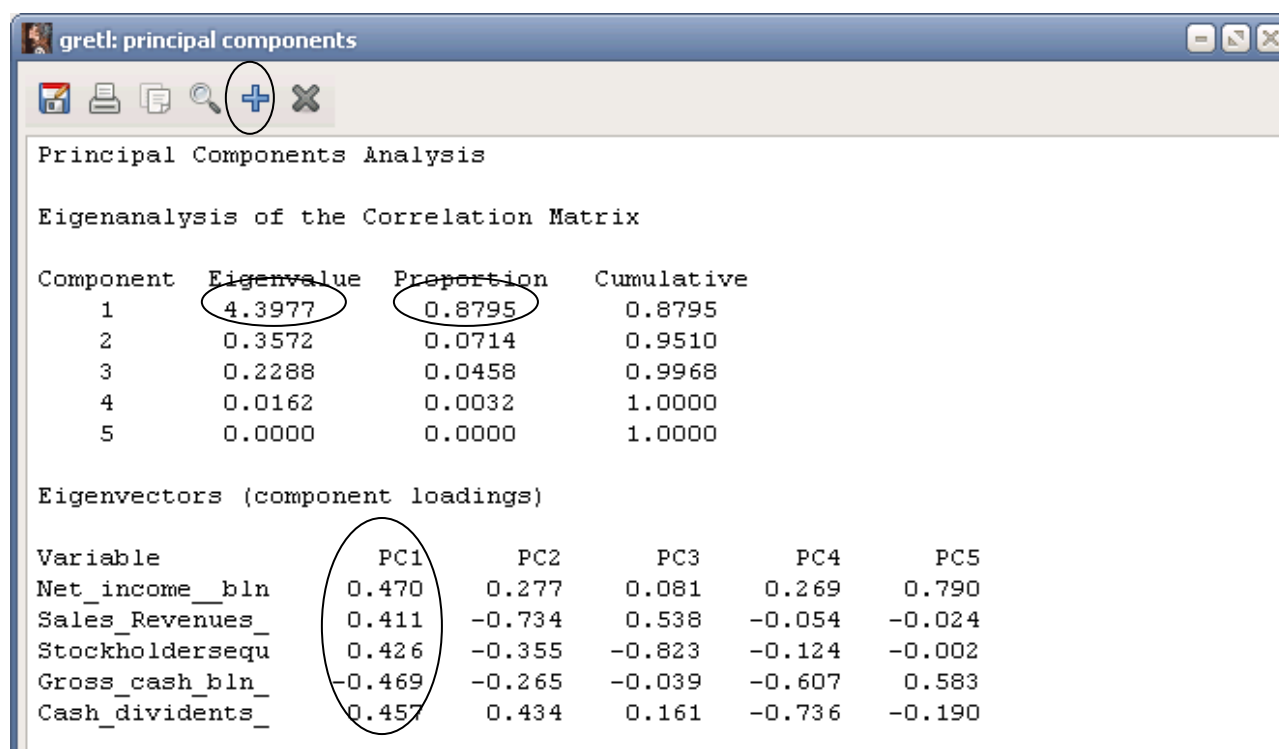


Рисунок 12 – Окно результатов моделирования методом главных компонент

Таблица 5 – Значения атрибутов и показателей окна, представленного на рисунке 12

Атрибут или показатель	Объяснение
Principal Component Analysis	– анализ главных компонент;
Eigenanalysis of the Correlation Matrix	– нахождение собственных чисел корреляционной матрицы;
Component	– номер главной компоненты;
Eigenvalue	– собственное число λ_i корреляционной матрицы, которое характеризует абсолютный вклад соответствующей главной компоненты y_i в суммарную дисперсию исходных признаков $x_1, \dots, x_p$ (весомость y_i);
Proportion	– процентный вклад главной компоненты y_i в дисперсию $x_1, \dots, x_p$, (доля λ_i от общей суммы собственных чисел);
Cumulative	- процентный вклад главной компоненты нарастающим итогом, %;
Eigenvectors	- собственные вектора $w_j = (w_{1j}, \dots, w_{pj})'$ корреляционной матрицы, коэффициенты главных компонент;
Variable	- переменная (исходный признак x_i);
PC _i (Principal Component)	- главная компонента y_i

Выберем главную компоненту PC_1 (или y_1) с максимальным собственным числом Eigenvalue (λ_1).

Столбцы PC_i показывают значения коэффициентов главных компонент $w_j = (w_{1j}, \dots, w_{pj})'$, по данным столбца PC_1 построим первую главную компоненту y_1 , формула (2).

$$y_1 = 0.47NetIncome + 0.411SalesRevenues + 0.426StockholdersEquity - 0.469GrossCash + CashDividends; \quad (2)$$

где X_{11} - Net income, bln\$ (Чистая прибыль, млрд. \$);

X_{21} - Sales&Revenues, bln\$ (Выручка, млрд.\$);

X_{31} - Stockholders' Equity, bln\$ (Акционерный капитал, млрд.\$);

X_{41} - Gross cash, bln\$ (Валовая наличность, млрд.\$);

X_{51} - Cash dividends, \$ (Наличные дивиденды, \$).

Первое (наибольшее) значение столбца Eigenvalue рисунка 12 показывает абсолютный вклад $\lambda_1 = 4.3977$ первой главной компоненты PC_1 (или y_1) в общую дисперсию наблюдаемых признаков $X_{11} \dots X_{51}$.

Первое (наибольшее) значение столбца Proportion рисунка 12 показывает относительный вклад первой главной компоненты PC_1 (или y_1) равный 87,95% ($=4.3977/(4.3977+0.3572+0.2288+0.0162+0.0000)*100\%$) в общую дисперсию наблюдаемых признаков $X_{11} \dots X_{51}$, следовательно остальными главными компонентами $PC_2 \dots PC_5$ (с незначительными вкладами в общую дисперсию) можно пренебречь, а PC_1 рассматривать как главную компоненту с наибольшей значимостью (весомостью).

Тогда выбираем PC_1 (или y_1), формула (2), как первый элемент вектора функций Ford Motor Company, характеризующий финансовую составляющую устойчивого развития в целом, а соответствующее собственное число $\lambda_1 = 4.3977$ будем рассматривать как финансовый рейтинг предприятия.

Сохраним новую переменную - главную компоненту y_1 в созданном наборе данных Ford.gdt, нажав на знак "+" меню окна, представленного на рисунке 12, и кнопку ОК.

Аналогичным описанному выше способу построим главные компоненты для остальных составляющих устойчивого развития «продукт и потребители» (y_2), «качество отношений с персоналом» (y_3), «окружающая среда и безопасность» (y_4) и получим следующие значения, представленные формулами (3-5).

$$y_2 = 0.492MarketShareUS + 0.416MarketShareEur - 0.405CustomerSatisf + 0.432CustomerLoyal + 0.483VehiclesSold, \quad (3)$$

где $\lambda_2 = 3.6503$ рейтинг составляющей «Продукт и потребители»;

X_{12} - Market share US, % (Доля рынка в США, %);

X_{22} - Market share Europe (Доля рынка в США, %);

X_{32} - Customer satisfaction, % (Удовлетворённость потребителя, %)

X_{42} - Customer Loyalty, US % (Верность потребителя, США %)

X_{52} - Vehicles sold, units (Объём продаж, шт.).

$$y_3 = -0.582 \text{Personnel}_{Full} + 0.572 \text{EmployeeSatisf} + 0.578 \text{LaborCost}, \quad (4)$$

где $\lambda_3 = 2.9468$ рейтинг составляющей «Качество отношений с персоналом»;

X_{13} - Personnel full-time (Численность персонала на полной ставке, чел.);

X_{23} - Employee satisfaction, % (Удовлетворённость работников, %);

X_{33} - Laybor cost per hour, \$ (З.пл. в час, \$).

$$y_4 = -0.512 \text{FuelEconomy} - 0.623 \text{EnergyEfficiency} + 0.592 \text{GlobalFatalities}, \quad (5)$$

где $\lambda_4 = 2.1022$ «Окружающая среда и безопасность»;

X_{14} - Ford U.S. Corporate Average Fuel Economy (Экономия топлива, США, миль на галлон);

X_{24} - Energy efficiency index, % (Индекс эффективности энергопотребления, %);

X_{44} - Global Fatalities (Число смертельных случаев на производстве).

Т.о. получен вектор функций $\mathbf{a}=(y_1, y_2, y_3, y_4)'$, для предприятия Ford Motor Company.

Нажав на пиктограмму системного калькулятора в нижнем левом углу стартового экрана Gretl (рисунок 3), рассчитаем общий рейтинг устойчивого развития Ford $\lambda_{CP-ford} = (\lambda_1 + \lambda_2 + \lambda_3 + \lambda_4) / 4 = 3.27425$ как среднее арифметическое вышеперечисленных рейтингов отдельных составляющих. Данный показатель будем использовать для сопоставления значений рейтингов нескольких предприятий при сравнительном анализе их уровня устойчивого развития.

4. ПОРЯДОК ВЫПОЛНЕНИЯ ЛАБОРАТОРНОЙ РАБОТЫ

1. Выполнить пример, рассмотренный в п.3.

2. Сравнить рейтинги устойчивого развития двух предприятий согласно варианту (таблица 6):

- Перенести исходные данные по каждому предприятию в зависимости от варианта в файл data.xls и импортировать data.xls в систему Gretl как описано в вышеприведённом примере. Перед импортом в систему Gretl файла data.xls удалить в нём все знаки «*», символизирующие пропущенные наблюдения, а также представить показатели, рассчитанные в йенах, в долларах (1yen=0.00863 US\$).

- Оценить степень корреляции между исходными признаками и в случае возможности применения метода главных компонент для каждого предприятия построить главные компоненты и выбрать наиболее весомую из них для финансовой составляющей (y_1) устойчивого развития и составляющей «Продукт и потребители» (y_2), определить соответствующие частные рейтинги и рассчитать общий рейтинг устойчивого развития предприятия.

- Ознакомиться с данными финансовых отчётов и отчётов по устойчивому развитию сравниваемых компаний, представленными в приложении А. Провести сравнительный анализ предприятий по всем рассчитанным рейтингам.

Таблица 6 – Варианты заданий

Вариант	Компании для сравнения и исходные данные
---------	--

1	Ford Motor Company – General Motors (таблицы 1,2,7,8)
2	Ford Motor Company - Nissan Motor Co (таблицы 1,2,9,10)
3	Ford Motor Company - Toyota Motor Co (таблицы 1,2,11,12)
4	General Motors - Nissan Motor Co (таблицы 7,8,9,10)
5	General Motors - Toyota Motor Co. (таблицы 7,8,11,12)
6	Nissan Motor Co - Toyota Motor Co.(таблицы 9,10,11,12)

Таблица 7 - Показатели финансовой составляющей компании General Motors

Year	Net income, bln\$	Net Sales& Revenues, bln\$	Stockholders equity, bln\$	Gross cash,bln\$	Cash dividends,\$	Total assets,bln\$	Earnings per share,\$
2002	1.736	186.763	6.637	20.32	2	369.053	3.36
2003	3.822	185.837	24.867	32.554	2	448.507	6.71
2004	2.701	195.351	27.886	23.3	2	479.603	4.78
2005	-10.417	194.655	14.653	20.4	2	474.156	-18.42
2006	-1.978	207.349	-5.441	26.4	1	186.192	-3.5

Таблица 8 - Показатели составляющей «Продукт и потребители» General Motors

Year	Market Share, Global, %	Vehicles sold worldwide, units
2002	*	8411000
2003	*	8098000
2004	14.3	9098000
2005	14.1	9051000
2006	13,5	9181000

Таблица 9 - Показатели финансовой составляющей компании Nissan Motor Co.

Year	Net income, bln yen	Net sales, bln yen	Shareholders' equity, bln yen	Gross cash, bln yen	Cash dividends, yen	Total assets, bln yen	Return on equity, %
2002	495	6829	1808.3	269.817	14	7349.2	4
2003	504	7429	2024	194.164	19	7859.9	4.6
2004	512	8577	2465.75	289.784	24	9848.523	22.8
2005	518	9428	3087.983	404.212	29	11481.426	18.7
2006	461	10469	3277.956	469.388	34	11481.426	*

Таблица 10 - Показатели составляющей «Продукт и потребители» Nissan Motor Co.

Year	Domestic Market Share, %	Vehicles sold worldwide, units	Vehicles produced worldwide, units
2002	13.9	2771000	2586602
2003	14.2	3057000	2883409
2004	14.6	3389000	*
2005	14.4	3569000	3340827
2006	13.2	3483000	3428981

Таблица 11 - Показатели финансовой составляющей компании Toyota Motor Co.

Year	Net income, bln yen	Net revenues, bln yen	Return on capital, %	Dividends per share, yen	Return on equity, %	Return on assets, %
2002	615.8	15000	5.4	28	7.8	3.1
2003	750	15501	6.2	36	10.4	3.8
2004	1162	17294	8.4	45	15.2	5.5
2005	1171	18551	7.6	65	13.6	5.1
2006	1372	21036	7.9	90	14	5.2

Таблица 12 - Показатели составляющей «Продукт и потребители» Toyota Motor Co.

Year	Domestic Market Share, %	Vehicles sold, units	Vehicles produced, units
2002	38.1	5909000	5983000
2003	*	6113000	5850000
2004	40.8	6719000	6513000
2005	*	7408000	7232000
2006	37.5	7974000	7711000

5. СОДЕРЖАНИЕ ОТЧЕТА О ВЫПОЛНЕНИИ ЛАБОРАТОРНОЙ РАБОТЫ

- 1) Название и цель работы.
- 2) Постановка задачи.
- 3) Этапы выполнения задачи в Gretl.
- 4) Выводы.

АНАЛИЗ ВРЕМЕННЫХ РЯДОВ В СРЕДЕ GRETЛ 1.7.1

1. ЦЕЛЬ РАБОТЫ

Целью данной работы является получение практических навыков анализа временных рядов с использованием инструментария системы Gretl 1.7.1. для объяснения механизма и составления прогноза экономических явлений.

2. ТЕОРЕТИЧЕСКИЙ РАЗДЕЛ

Поскольку экономические условия ведения бизнеса изменяются во времени, специалистам по управлению необходимо отслеживать динамику этих изменений для успешной реализации деловых операций. Одним из приемов, которым менеджеры могут воспользоваться при оценке эффективности будущих управленческих решений, являются методы прогнозирования, основанные на анализе временных рядов, цель которых – предсказать с той или иной степенью надежности будущие события и учесть этот прогноз при планировании тех или иных управленческих решений. Прогноз по временным рядам предусматривает определение прогнозного значения переменной исключительно на основе прошлых и текущих значений этой же переменной.

Специалисты предприятия должны уметь прогнозировать спрос на свою продукцию, предпочтения потребителей, будущий объем продаж и эффективность рекламных кампаний; на уровне правительственных структур требуется оценка будущего уровня инфляции, безработицы, объема производства.

Временным рядом называют последовательность значений экономического показателя, упорядоченных во времени. Например, временными рядами будут серия ежедневных наблюдений в течение некоторого периода за ценами товара при закрытии торгов на бирже, дневные объемы выпуска товара, месячные показатели инфляции или индекса потребительских цен, ежеквартальные оценки валового национального продукта или средних зарплат, ежегодные данные об объеме, выручке и прибыли компании.

Временные ряды делятся на нестационарные и стационарные. Стационарные ряды имеют постоянные по времени среднее, дисперсию и автокорреляции. Временные ряды, не являющиеся стационарными, называются нестационарными.

Временные ряды также делятся на одномерные и комплексные, в первом случае подгонка фактических данных осуществляется к одной модели, во втором - к нескольким.

Цель эконометрического анализа временных рядов состоит в построении по возможности простых моделей, адекватно описывающих имеющиеся ряды наблюдений и составляющих базу для решения, в первую очередь, следующих задач: объяснение механизма формирования уровней ряда, построение прогноза будущих значений временного ряда.

Среди наиболее распространённых методов анализа временных рядов можно выделить:

- Анализ тренда (позволяет отфильтровать шум и периодические колебания, преобразуя данные в относительно гладкую кривую, для выявления тенденции ряда и прогноза его будущих значений);
- Декомпозиция временного ряда (позволяют выделить в ряде тренд, сезонную, циклическую и случайную составляющую для проведения его структурного анализа);
- Автокорреляционный анализ (позволяет определить периодические компоненты ряда);
- ADL (анализ распределенных лагов, используется если требуется предсказать значения одного ряда на основе значений другого);
- ARIMA (описывает большинство экономических процессов, используются для составления краткосрочных прогнозов по историческим данным), включает метод авторегрессии AR;
- ARIMA с интервенцией (необходимость в такого рода анализе возникает, когда с некоторого момента резко изменяется поведение ряда в силу внешних причин);
- Спектральный (Фурье) анализ (служит для определения скрытых периодичностей в данных стационарных временных рядов).

2.1. Анализ тренда

Экономические процессы чаще всего имеют многокомпонентную структуру (1).

$$y_t = I_t + s_t + v_t + u_t, \quad (1)$$

где I_t - тренд, представляющий собой ход кривой, со спокойным гладким характером, которая описывает долговременные изменения и определяет главное направление развития;

s_t - сезонная компонента, кратковременные регулярные колебания;

v_t - циклическая компонента, долгосрочные регулярные колебания.

u_t - случайная составляющая временного ряда, отражающая воздействие многочисленных факторов случайного характера.

Отметим, что стационарные временные ряды не содержат тренда, сезонной и циклической компонент, а каждый следующий их уровень образуется как сумма среднего уровня ряда и некоторой (положительной или отрицательной) случайной компоненты.

Функции тренда для одномерных временных рядов могут представляться полиномами различных степеней и другими функциями относительно переменной времени $t=1,2,\dots,n$ (моделями кривых роста), рисунок 1.

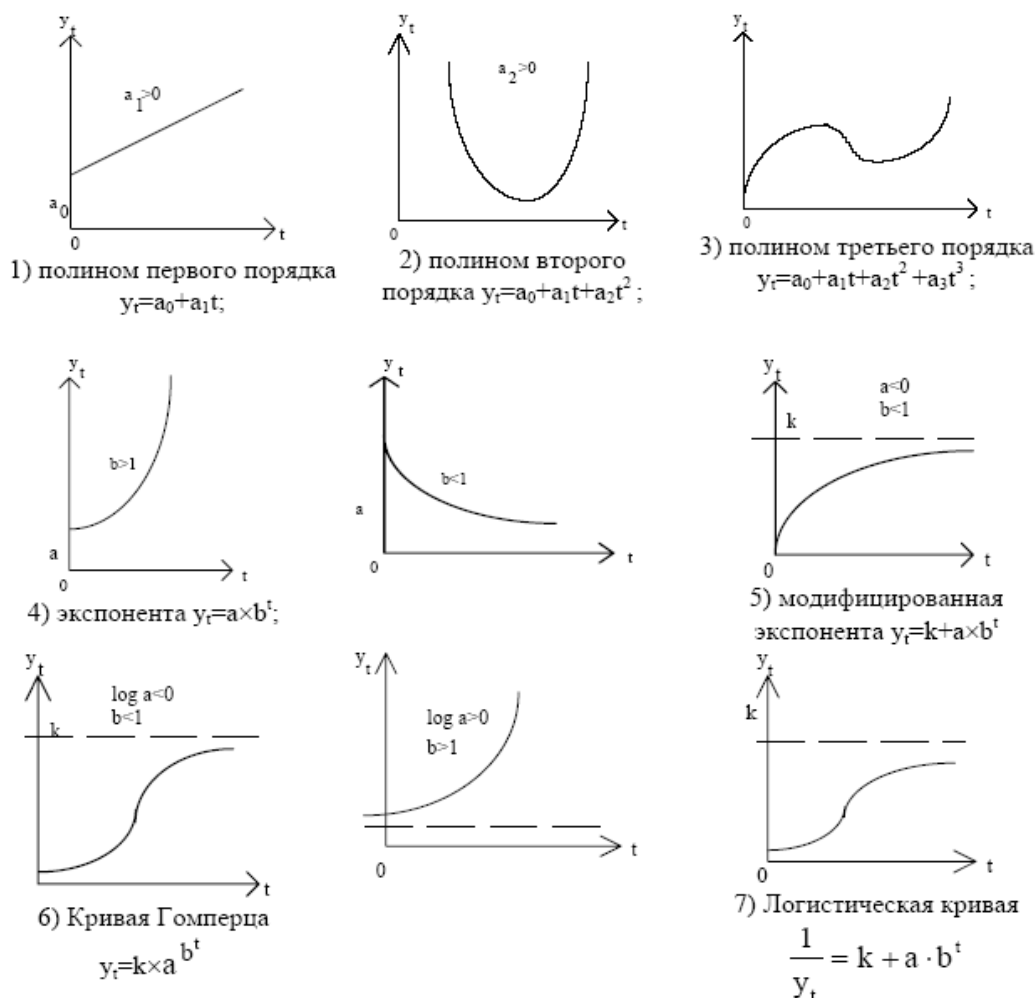


Рисунок 1 - Основные функций тренда и их графики

Анализ (выделение) тренда состоит из двух этапов:

- сглаживание временных рядов методом простой или взвешенной скользящей средней;
- применение моделей кривых роста для описания функции тренда и прогнозирования ряда.

Сглаживание используется как первый шаг при выделении (анализе) тренда временных рядов, содержащих значительную ошибку u_t , а также для удаления сезонной компоненты s_t . Суть различных приемов сглаживания сводится к замене фактических уровней временного ряда расчетными уровнями, которые подвержены колебаниям в меньшей степени. Сглаживание всегда включает некоторый способ локального усреднения данных, при котором несистематические компоненты и периодические колебания взаимно погашают друг друга.

Общий метод сглаживания – метод скользящего среднего, в котором каждый член ряда заменяется простым или взвешенным средним g соседних членов, где g - длина интервала сглаживания. При этом отфильтровывают шум и преобразуются данные в относительно гладкую кривую. Простое скользящее среднее определяется по формуле (2).

$$\hat{y}_t = \frac{\sum_{i=t-p}^{t+p} y_i}{2p+1} = \frac{y_{t-p} + y_{t-p+1} + \dots + y_{t+p-1} + y_{t+p}}{2p+1}, \quad (2)$$

где y_i - фактическое значение i -го уровня временного ряда;

$2p+1$ – длина интервала сглаживания g ;

$\hat{y}_t$ - значение скользящего среднего в момент t .

Взвешенное скользящее среднее определяется по формуле взвешенного среднего арифметического (значениям y_i присваиваются веса).

На втором этапе анализа тренда применяют одну из моделей кривых роста, которая позволяет описать функцию тренда, т.е. получить выровненные (теоретические) значения уровней ряда для его прогнозирования. Среди кривых роста одной из наиболее распространённых является класс полиномов (3).

$$y_t = a_0 + a_1 t + a_2 t^2 + \dots + a_p t^p, \quad (3)$$

где a_i ($i=0 \dots p$) – параметры многочлена;

t - независимая переменная (время).

Оценки параметров полиномиальных моделей (3) определяются методом наименьших квадратов (МНК), суть которого состоит в нахождении таких параметров, при которых сумма квадратов отклонений расчетных значений уровней $\tilde{y}_i$ от фактических y_i значений была бы минимальной, формула (4). Параметры экспоненциальных и S-образных кривых находятся более сложными методами.

$$\sum_i (y_i - \tilde{y}_i)^2 \rightarrow \min. \quad (4)$$

2.2. Декомпозиция временного ряда

Методы декомпозиции позволяют выделить тренд, сезонную, циклическую и случайную составляющие временного ряда, формула (1). Многие экономические процессы характеризуются явлением сезонности и требуют исключения её влияния для их дальнейшего изучения, например, при оценивании тенденции развития (тренда).

Центральными статистическими управлениями широко используются две процедуры: X-12-ARIMA, применяемая американским Bureau of the Census (Бюро переписей США) и TRAMO/SEATS, рекомендованная EUROSTAT.

Пакет программ Gretl позволяет использовать обе процедуры для структурного анализа ряда и исключения влияния сезонности. Для этого необходимо установить два дополнительных приложения: x12a_install.exe и ts_install.exe, доступные на сайте www.kufel.torun.pl/ru

2.3. Анализ сезонности. Коррелограмма

В общем виде периодическая зависимость может быть формально определена как корреляционная зависимость между каждым i -м и $(i-p)$ -м элементом ряда, при этом p называют лагом (сдвигом, запаздыванием).

Данную зависимость оценивает автокорреляционная функция АСФ (Autocorrelation function), вычисляемая как последовательность корреляций между рядом и им же, сдвинутым на $1, 2, \dots, p, \dots$ временных точек, лагов. АСФ представляет собой коэффициенты автокорреляции, формула (5), для последовательности лагов из определенного диапазона, т.е. зависимость между разнесёнными по времени наблюдениями.

$$r(t, s) = \frac{M[(y_t - \bar{y}_t)(y_s - \bar{y}_s)]}{\sigma_t \sigma_s}, \quad (5)$$

где M - математическое ожидание;

σ_t - среднеквадратическое отклонение y_t ;

$\bar{y}_t$ - среднее значение y_t ;

t и s – различные моменты времени, где $t-s = p$, (p – лаг);

y_t – уровень временного ряда в момент времени t ;

$r(t, s)$ - автокорреляционная функция временного ряда y_t (t и s – переменные).

Т.о. коэффициент автокорреляции при лаге 1 есть коэффициент корреляции между y_t и y_{t-1} . Если ряд стационарен, то данный коэффициент равняется коэффициенту корреляции между y_t и y_{t-2} ; y_t и y_{t-3} и т.д.

Сезонные составляющие временного ряда могут быть найдены с помощью коррелограммы, которая представляет собой график зависимости значений автокорреляционной (АСФ) или частной автокорреляционной (РАСФ) функции от величины лага p (порядка коэффициента автокорреляции).

Однако, если коррелированы y_t и y_{t-1} , а также y_{t-1} и y_{t-2} , то и y_t и y_{t-2} также коррелированы, т.е. корреляция при лаге 1 "вызывает" корреляцию при лаге 2, а, соответственно, и при больших лагах. Поэтому для исследования периодичности и получения более «чистой» картины периодических зависимостей используется РАСФ (Partial autocorrelation function), которая представляет собой «чистую» зависимость между наблюдениями. Частная автокорреляция на данном лаге аналогична обычной автокорреляции, за исключением того, что при вычислении из нее удаляется влияние автокорреляций с меньшими лагами, например РАСФ при лаге 3 равняется корреляции рядов y_t и y_{t-3} , причем считается исключенной их корреляция с рядами y_{t-1} и y_{t-2} .

Т.о. РАСФ позволяет оценить порядок запаздывания процесса p для модели авторегрессии $AR(p)$, которая будет рассмотрена ниже.

2.4. Метод авторегрессии

Один из методов, используемый для прогнозирования по стационарным временным рядам, основан на авторегрессионных моделях. Обычно обнаруживается, что значения отклика в некоторой точке временного ряда сильно коррелированы с несколькими предшествующими и/или последующими значениями. Действительно, для многих явлений их современное состояние функционально определяется предшествующими состояниями системы. Как было рассмотрено выше, подобные связи принято называть автокорреляцией. Значимый коэффициент частной автокорреляции при наибольшем лаге (n) указывает на существование в анализируемом процессе авторегрессии AR (n) порядка, равного величине данного лага, формула (6).

$$y_t = a_0 + a_1 y_{t-1} + a_2 y_{t-2} + \dots + a_n y_{t-n} + u_t, \quad (6)$$

Оценив параметры модели (6) получаем механизм прогнозирования, основанный на предположении, что возникшая связь между значениями сохранится некоторое время в будущем.

2.5. Спектральный (Фурье) анализ

Цель спектрального анализа - разложить комплексные стационарные временные ряды с циклическими компонентами на несколько основных синусоидальных функций с определенной длиной волн, появление которых особенно существенно и значимо, формула (7). В результате анализа можно обнаружить всего несколько основных периодических компонент (функций синусов или косинусов) в изучаемом временном ряду, который, на первый взгляд, выглядит как случайный шум, что позволит изучить интересующее явление.

$$y_t = a_0 + \sum_{i=1}^q [a_i \cos(2\pi f_i t) + b_i \sin(2\pi f_i t)], \quad (7)$$

где a_0 – константа;

a_i, b_i – амплитуды соответствующих функций, коэффициенты регрессии, которые показывают степень, с которой соответствующие функции синусов и косинусов коррелируют с фактическими данными;

$2\pi f_i$ – круговая частота соответствующей функции, радиан.

f_i – частота Фурье или угловая частота, обратная периоду.

i – номер соответствующих гармоник (косинуса и синуса) с определённым значением частоты;

N – число наблюдений;

q – для нечётного числа наблюдений $q=(N-1)/2$ – число различных синусов и косинусов, для чётного $q= N/2$ и существует q значений косинусов и $q-1$ значений синусов.

Для нахождения частот основных периодических составляющих (синусов и косинусов) временного ряда вычисляется периодограмма, формула (8), путём суммирования квадратов коэффициентов a_i и b_i для каждой частоты и умножения на $N/2$. Значения периодограммы на графике изображаются в зависимости от частот или периодов.

$$I(f_i) = \frac{N}{2}(a_i^2 + b_i^2), \quad (8)$$

$$i = 1, 2, \dots, q$$

где $I(f_i)$ значение периодограммы на частоте f_i ;
 N - общая длина ряда.

Сглаженная периодограмма, состоящая из усреднённых методом взвешенного или простого скользящего среднего значений периодограммы, представляет собой функцию спектральной плотности и используется для тех же целей. Одним из методов взвешенного скользящего среднего является метод Бартлетта (Bartlett). В интервале сглаживания, для каждой частоты, веса для взвешенного скользящего среднего значений периодограммы вычисляются как: $w_i = 1 - (i/p)$ (для $i = 0$ до p); $w_{-i} = w_i$ (для $i \neq 0$); $p = (g-1)/2$, где g - длина интервала сглаживания.

3. ОПИСАНИЕ СРЕДСТВ АНАЛИЗА ВРЕМЕННЫХ РЯДОВ СИСТЕМЫ GRETL

3.1. Пример построения полиномиальной модели тренда

Для построения модели полиномиального тренда в Gretl 1.7.1 из меню **Model** выбирается пункт **Ordinary Least Squares** (рисунок 2).

Продemonстрируем выбор степени полинома относительно переменной времени t при построении модели тренда на примере ежемесячных данных об уровне безработицы в Польше с 1993 по 2005 год (файл macro_1993_2005 доступен на сайте www.kufel.torun.pl/ru).

Откроем набор исходных данных macro_1993_2005 на закладке KUFEL, выбрав пункт меню **File\Open Data\Sample file** или создадим его, перенеся данные Приложения А в файл macro.xls и импортировав его в пакет Gretl. Для этого в меню Gretl выберем команду **File\Open Data\Import\Excel**, в появившемся окне укажем номер строки и столбца начала таблицы Excel и нажмём кнопку ОК, в следующем окне нажмём кнопку YES, затем выберем тип данных временной ряд (time series), нажмём кнопку FORWARD, отметим периодичность данных MONTHLY, нажмём кнопку FORWARD, введём год и месяц начала сбора данных 1993:01, нажмём кнопку FORWARD и ОК.

Сохраним созданный набор данных (рисунок 2) **File\Save Data** как файл macro_1993_2005 на рабочем столе.

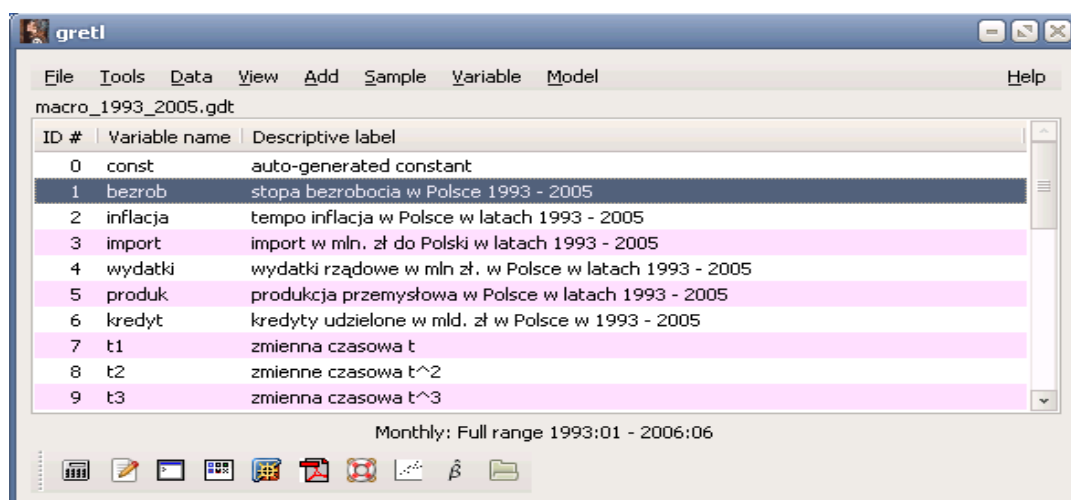


Рисунок 2 – Набор данных macro_1993_2005

Обратимся к команде **Model\Ordinary Least Squares** (метод наименьших квадратов). Для построения модели линейного тренда в открывшемся окне спецификации модели (рисунок 3) при помощи кнопки Choose выберем зависимую переменную Y (Dependent variable) – уровень безработицы (bezrob) и добавим объясняющую переменную X при помощи кнопки ADD (переменная времени t1 указывает номер месяца), нажмём кнопку ОК.

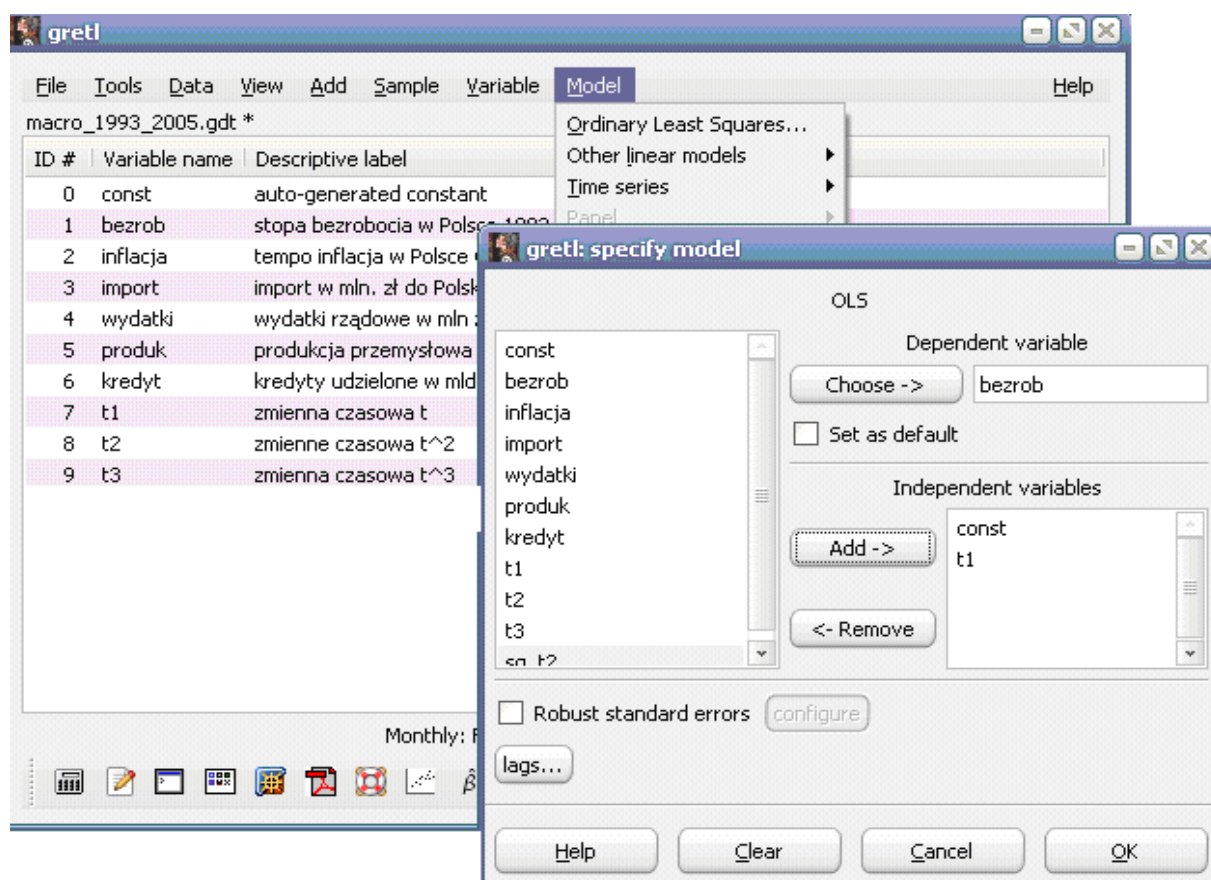


Рисунок 3 – Построение модели линейного тренда

Результаты оценивания линейного тренда показаны в окне результатов моделирования (рисунок 4) и на графике (рисунок 5). Полученная модель $y_t = 12,7823 + 0,0314t_1$ адекватна, параметры значимы (по t-критерию Стьюдента) на уровне 1%.

Из визуального анализа графика (рисунок 5) можно сделать вывод, что исходные данные представляют собой нестационарный временной ряд с выраженной цикличностью данных: ряд имеет тенденцию (непостоянное среднее), непостоянную дисперсию, существенные периодические колебания.

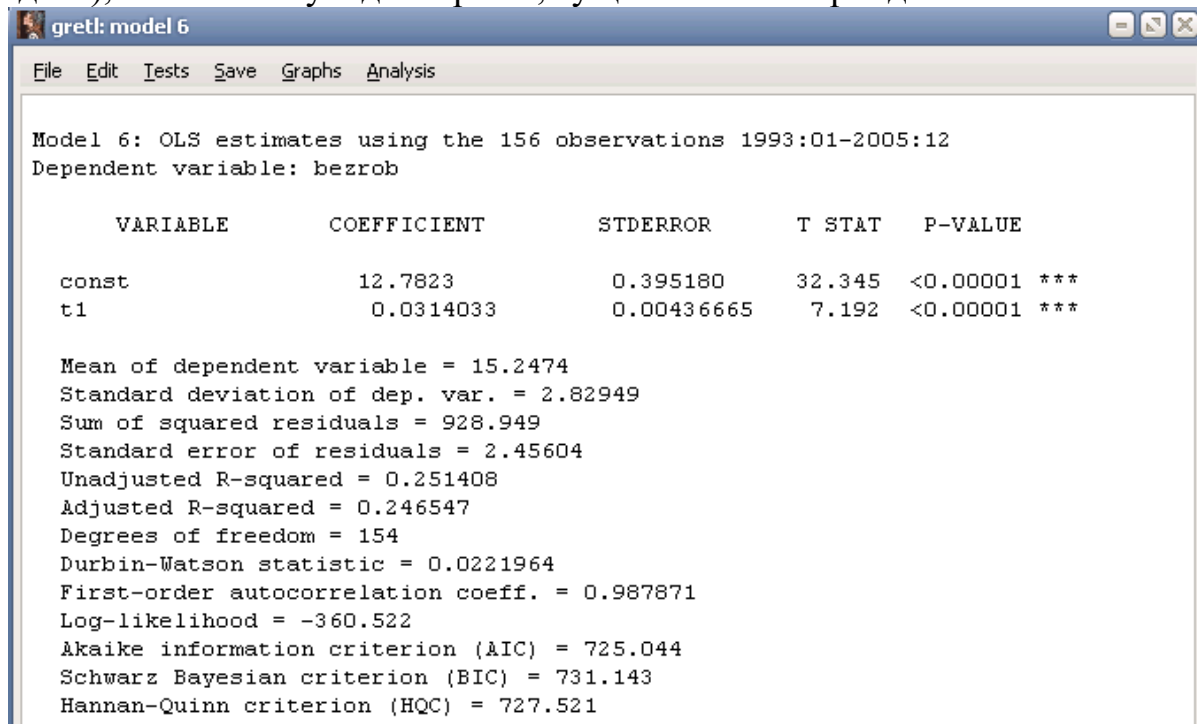


Рисунок 4 – Окно результатов моделирования линейного тренда ряда bezrob

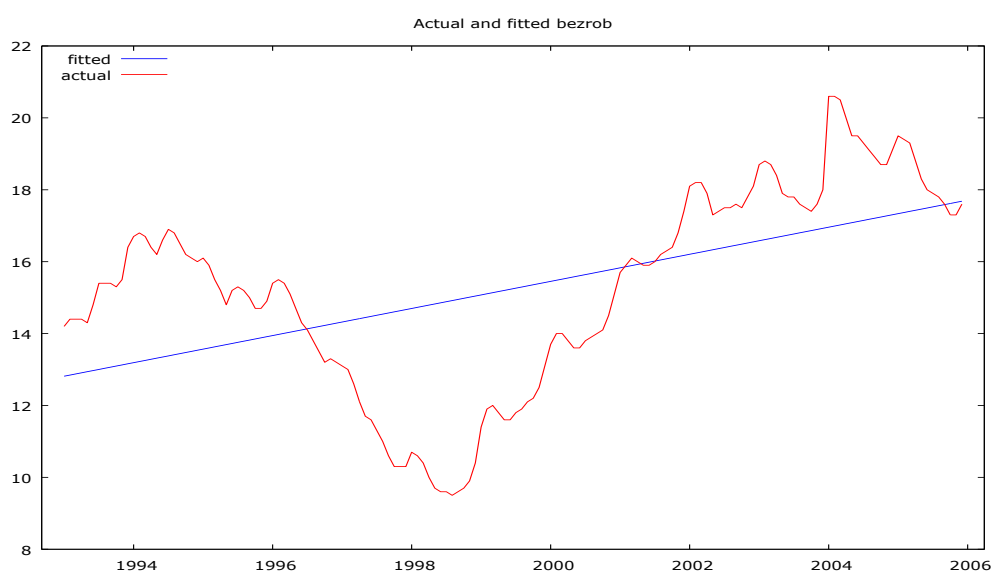


Рисунок 5 – Фактические данные и график линейного тренда

Для прогнозирования будущих значений временного ряда bezrob необходимо добавить соответствующее число «пустых» наблюдений к набору данных, например, 6 для составления прогноза на январь-июнь 2006 года. Для этого

обратимся к команде **Data\Add observations...**, введём число добавляемых наблюдений в открывшемся окне (6) и нажмём кнопку **ОК** (рисунок 6).

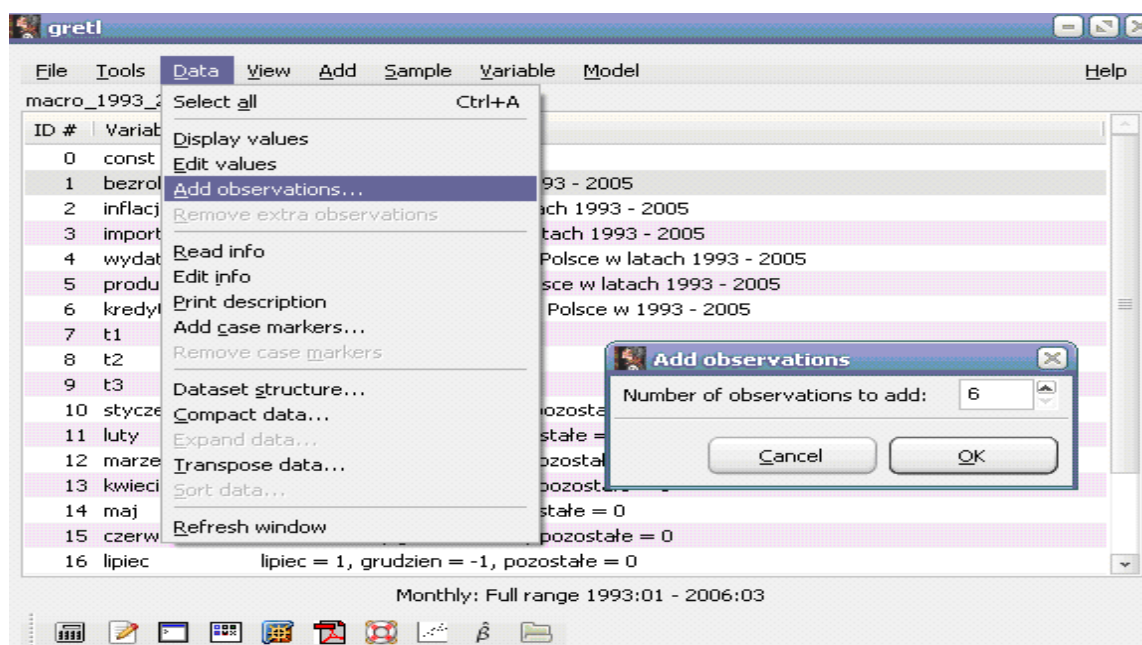


Рисунок 6 – Добавление «пустых» наблюдений к переменным набора данных **macro_1993_2005**

Затем повторим описанную выше последовательность действий по построению модели линейного тренда. В окне результатов моделирования (рисунок 4) выберем пункт **Forecasts** меню **Analysis**. В открывшемся окне введём общее число наблюдений до составления прогноза 156 и период для составления прогноза с января по июнь 2006 года (рисунок 7), нажмём кнопку **ОК**. Результаты прогнозирования представлены на рисунке 7 и 8.

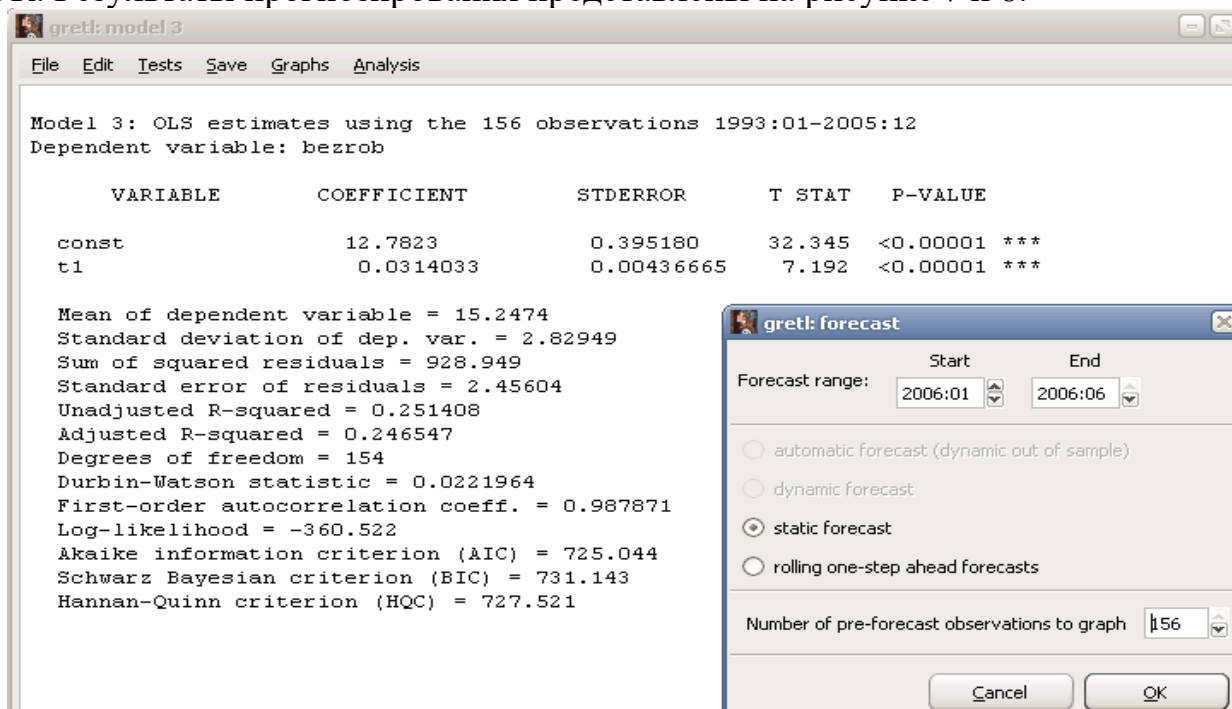
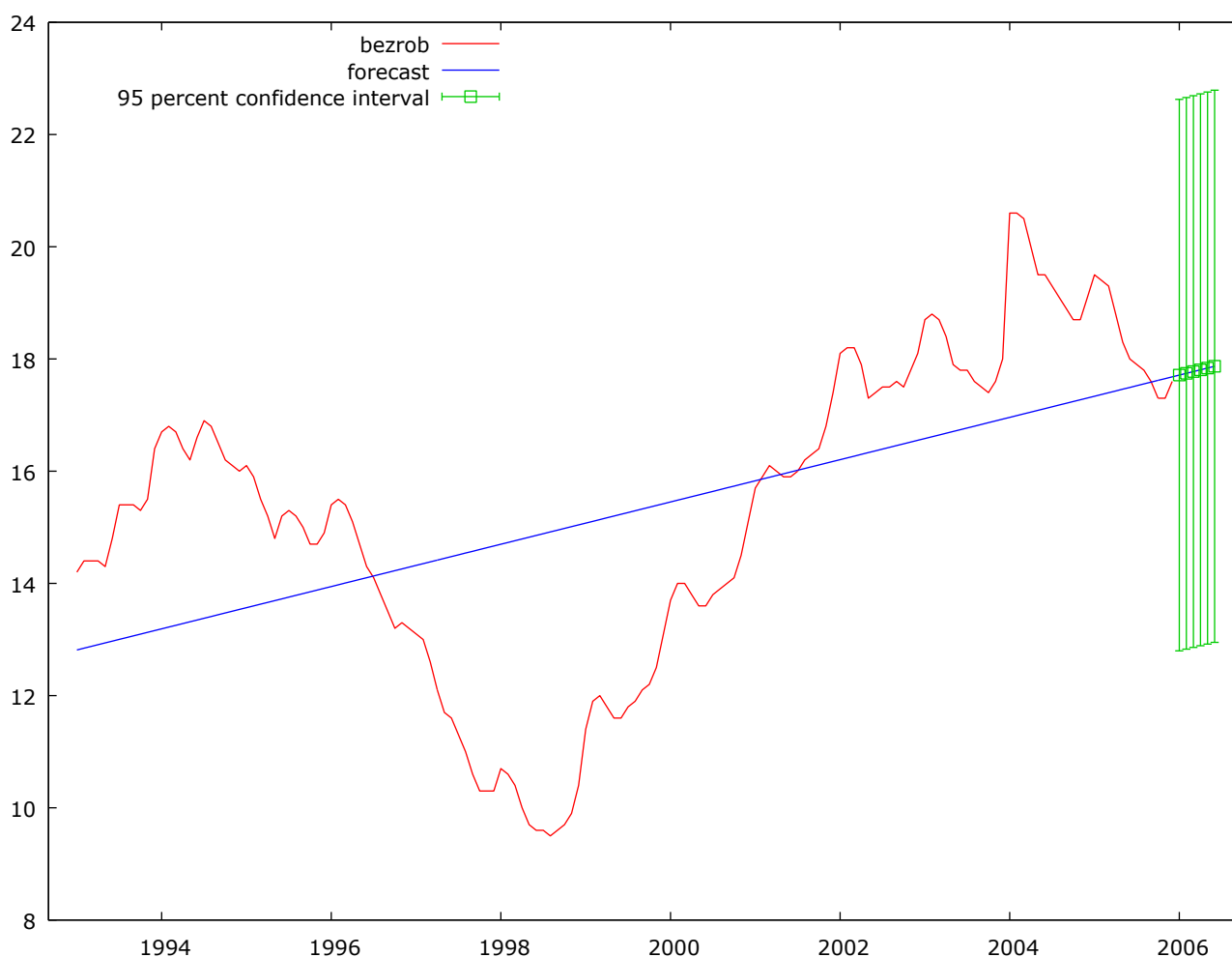


Рисунок 7 – Прогнозирование временного ряда **bezrob**



For 95% confidence intervals, $t(154, .025) = 1.975$				
Obs	bezrob	Prediction	std. error	95% confidence interval
Наблюдение		Прогноз		Доверительный интервал
2006:01	undefined	17.7126	2.48763	(12.7983, 22.6269)
2006:02	undefined	17.7440	2.48824	(12.8285, 22.6595)
2006:03	undefined	17.7754	2.48885	(12.8587, 22.6921)
2006:04	undefined	17.8068	2.48947	(12.8889, 22.7247)
2006:05	undefined	17.8382	2.49010	(12.9191, 22.7574)
2006:06	undefined	17.8696	2.49073	(12.9492, 22.7900)

Рисунок 8 – Результаты прогнозирования временного ряда bezrob

Построим модель полиномиального тренда четвертого порядка и сравним её с построенной моделью первого порядка.

Просмотрим переменные t_1, t_2, t_3 набора данных `masro_1993_2005`, дважды щёлкнув по их названию левой кнопкой мыши. Отметим, что $t_2 = t_1^2$, а $t_3 = t_1^3$. Для построения полинома четвертого порядка, необходимо добавить переменную $t_4 = t_1^4$

Щелчком мыши выберем переменную t_2 в открытом наборе данных и обратимся ко команде **Add\Squares of selected variables** (рисунок 9), что добавит переменную $sq_t2 = t_4 = t_1^4$ к набору данных.

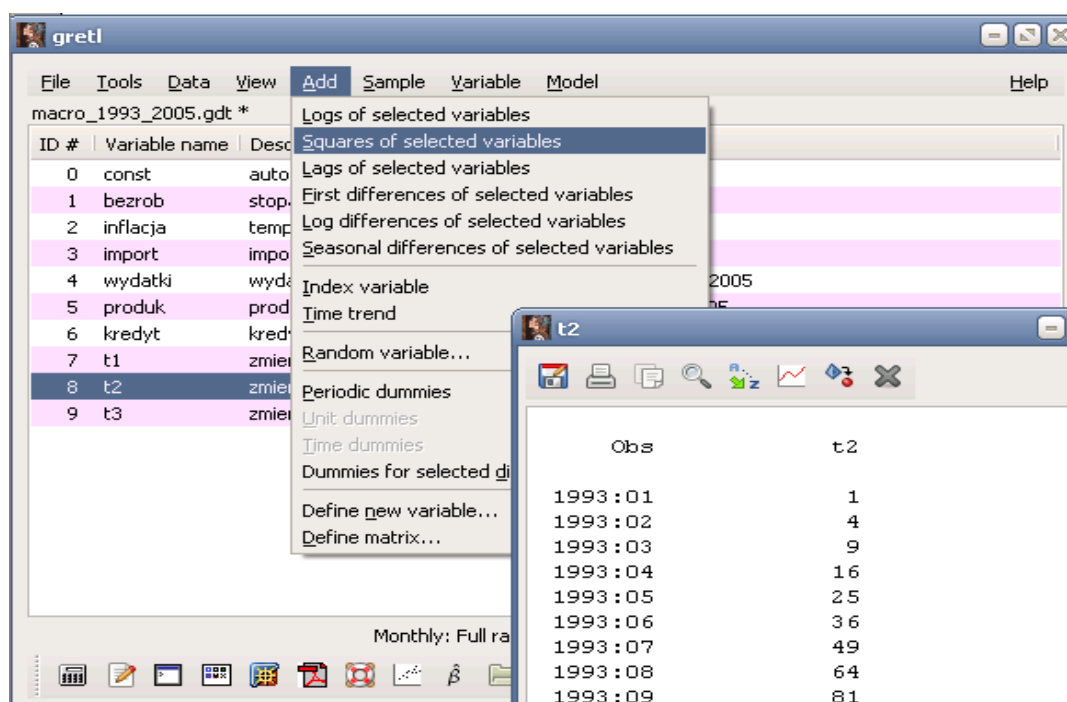


Рисунок 9 – Добавление переменной sq_t2

Повторим рассмотренную выше последовательность действий по построению линейного тренда, добавив в число объясняющих переменных t1,t2,t3, sq_t2 (рисунок 10). Окно результатов моделирования представлено на рисунке 11.

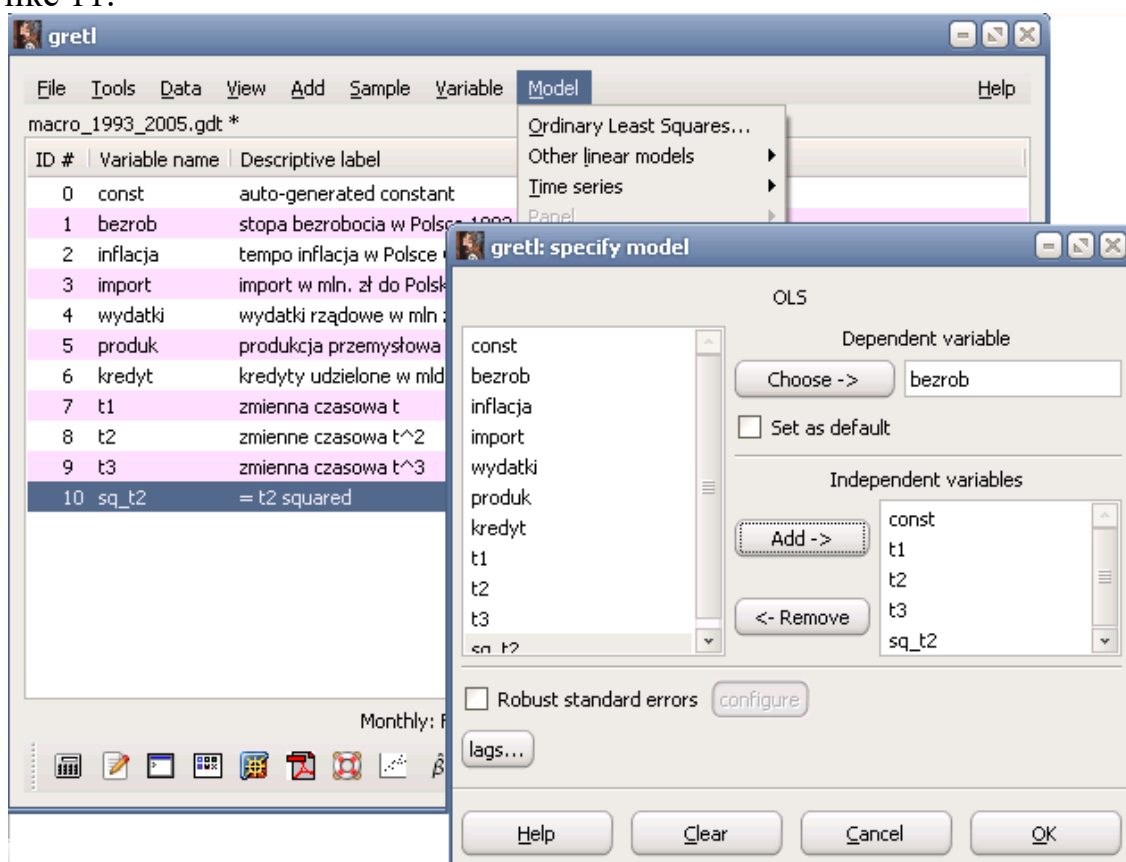


Рисунок 10 – Построение модели полиномиального тренда четвёртого порядка

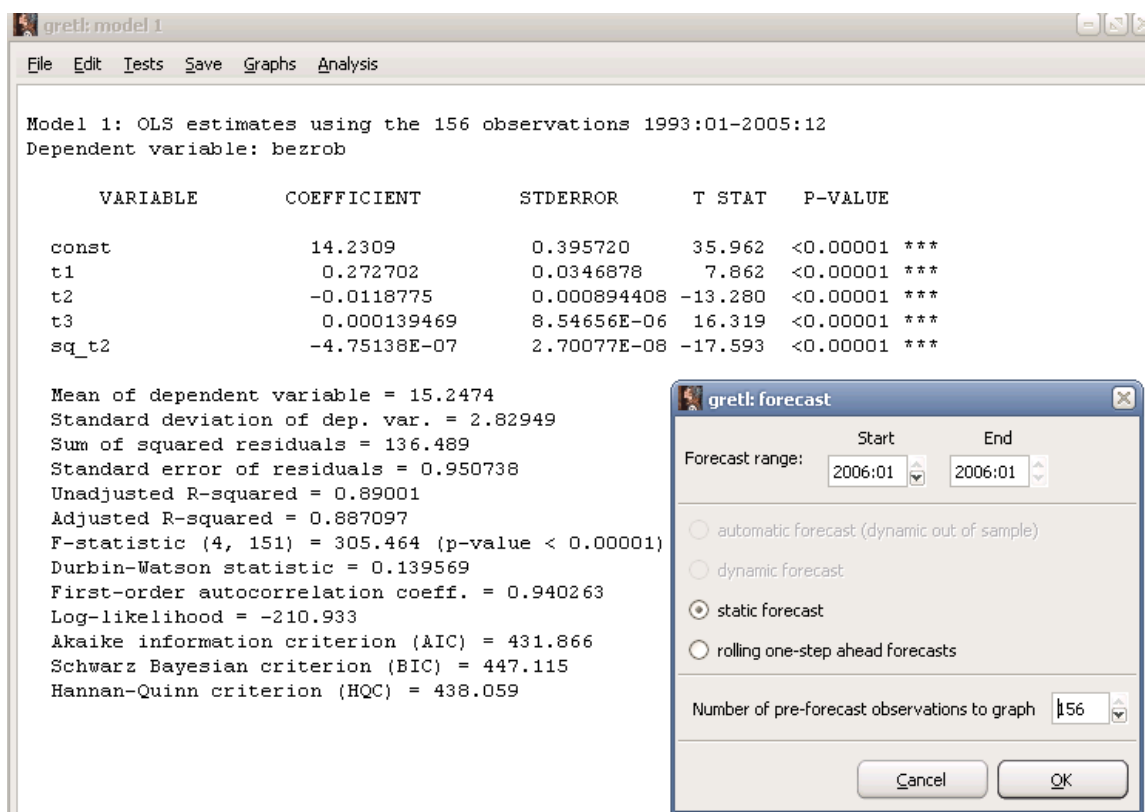
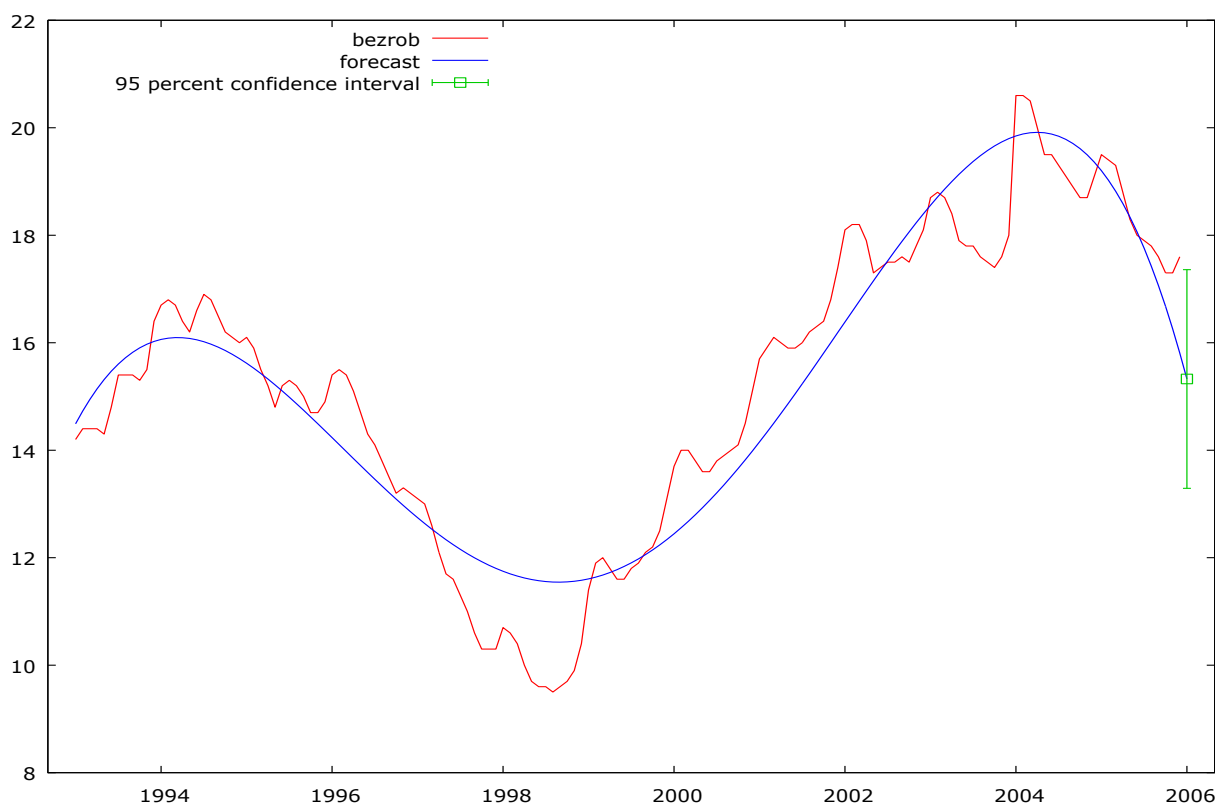


Рисунок 11 – Окно результатов моделирования полиномиального тренда четвёртого порядка ряда bezrob

Повторим рассмотренную выше последовательность действий по прогнозированию рассматриваемого временного ряда (рисунок 11), составив прогноз на 1 месяц (рисунок 12). Для составления прогноза необходимо, чтобы переменные t1, t2, t3, sq_t2 не имели пустых значений на период составления прогноза. Просмотреть и заполнить недостающие значения данных переменных можно обратившись к команде **View/Icon view** и дважды щёлкнув левой кнопкой мыши по иконке **Data set** в режиме просмотра и редактирования набора данных.

Полученная модель полиномиального тренда четвёртого порядка наилучшим образом описывает исходные данные, поскольку сумма квадратов ошибок RSS (Sum of squared residuals) для данной модели (136,489) существенно меньше данного значения для линейного тренда (928,949), в то время как коэффициент детерминации R^2 (Unadjusted R-squared), показывающий процент дисперсии bezrob, описанной моделью, напротив, значительно выше 88,71% >25,14%. При сравнении качества моделей также могут использоваться Т- и F-критерии (большее значение критерия наиболее предпочтительное).

Т.о. можно сделать вывод что построенная полиномиальная модель четвёртого порядка наиболее предпочтительна для прогнозирования рассматриваемого ряда, поэтому выбираем прогноз уровня безработицы в Польше на январь 2006 года равный 15,3257% согласно данной модели (рисунок 12). Реальный уровень безработицы в Польше в данный период составил 15%.



For 95% confidence intervals, $t(151, .025) = 1.976$

Obs	bezrob	prediction	std. error	95% confidence interval
2006:01	undefined	15.3257	1.02980	(13.2911, 17.3604)

Рисунок 12 – Результаты прогнозирования временного ряда bezrob с использованием модели полиномиального тренда четвёртого порядка

3.2. Пример декомпозиции динамики макроэкономических показателей

Проведём декомпозицию временного ряда значений уровня безработицы в Польше с 1993 по 2005 год (bezrob) рассматриваемого набора данных macro_1993_2005 методами TRAMO/SEATS и X-12-ARIMA для

-исключения влияния сезонной компоненты s_t , т.е. получения модифицированного сезонно скорректированного ряда;

- определения сезонной компоненты s_t и её прогноза;

- получения тренд-циклической компоненты ряда $(I_t + v_t)$;

- выделения случайной составляющей временного ряда u_t ;

- структурного анализа ряда в целом.

Сократим число наблюдений набора данных до исходного значения с 1993:01 по 2005:12, обратившись к команде **Sample\Set Range**.

Для использования процедуры X-12-ARIMA необходимо выбрать пункт **X-12-ARIMA analysis** меню **Variable** предварительно выбрав переменную bezrob в открытом наборе данных. Затем в открывшемся окне отметить флажками опции **Seasonally adjusted series**, **Trend/cycle**, **Irregular**, **Generate Graph** и нажать кнопку ОК (рисунок 13).

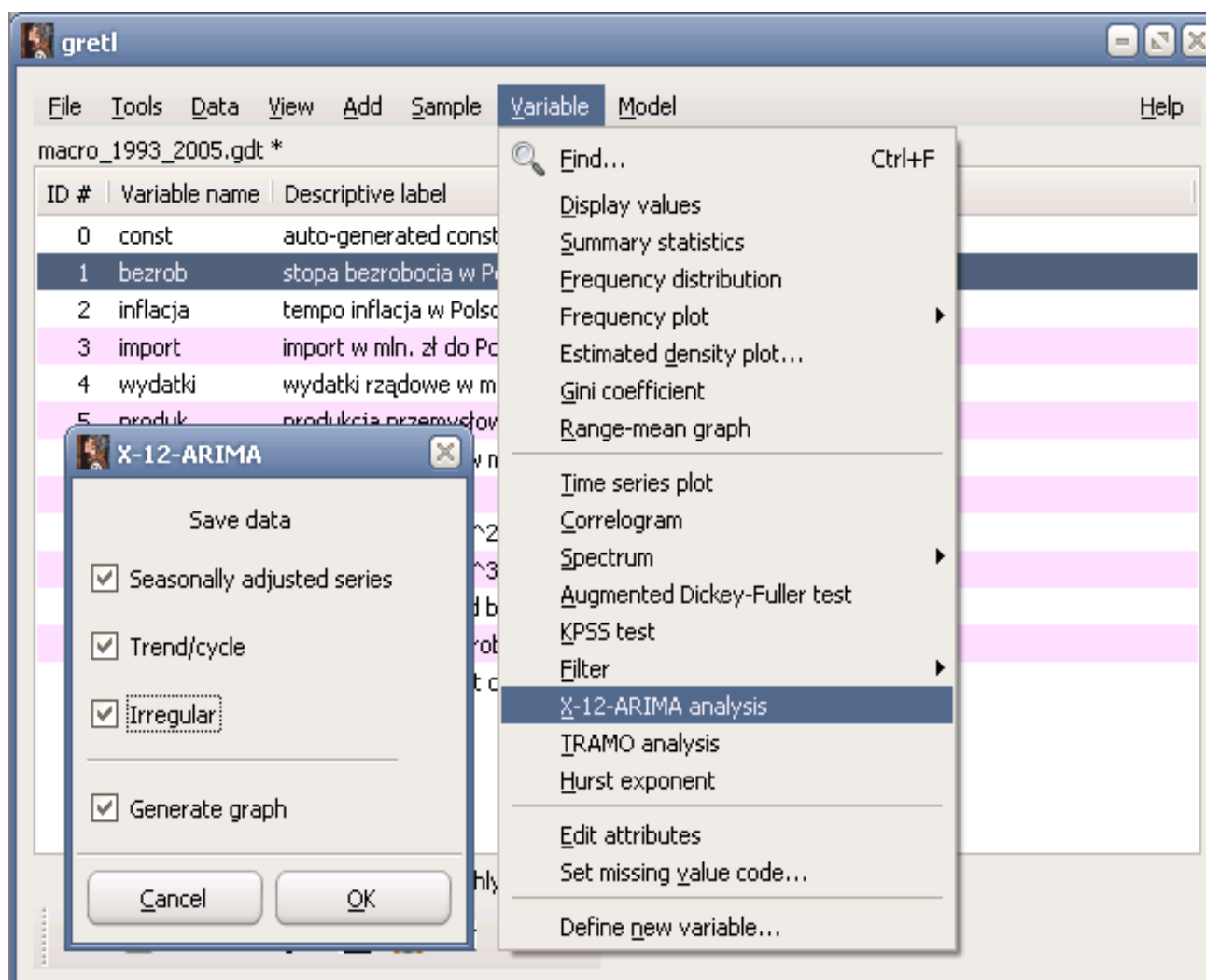


Рисунок 13 – Использование процедуры X-12-ARIMA

Выбор опции **Seasonally adjusted series** (**Ряд без учёта сезонной компоненты**) предполагает, что к исходному ряду будет применена сезонная корректировка, т.е. будет получен ряд без учёта влияния сезонной компоненты s_t , формула (1).

Выбор опции **Trend/cycle** (**Тренд/Цикл**) позволяет выделить трендовую I_t и циклическую v_t составляющие временного ряда, формула (1).

Опция **Irregular** (**Случайная составляющая**) позволяет выделить случайную составляющую временного ряда u_t , формула (1).

Опция **Generate Graph** (**Генерировать график**) выводит графическое отображение вышеперечисленных компонент ряда.

Ниже представлены результаты применения процедуры X-12-ARIMA и их графическое отображение (рисунок 14).

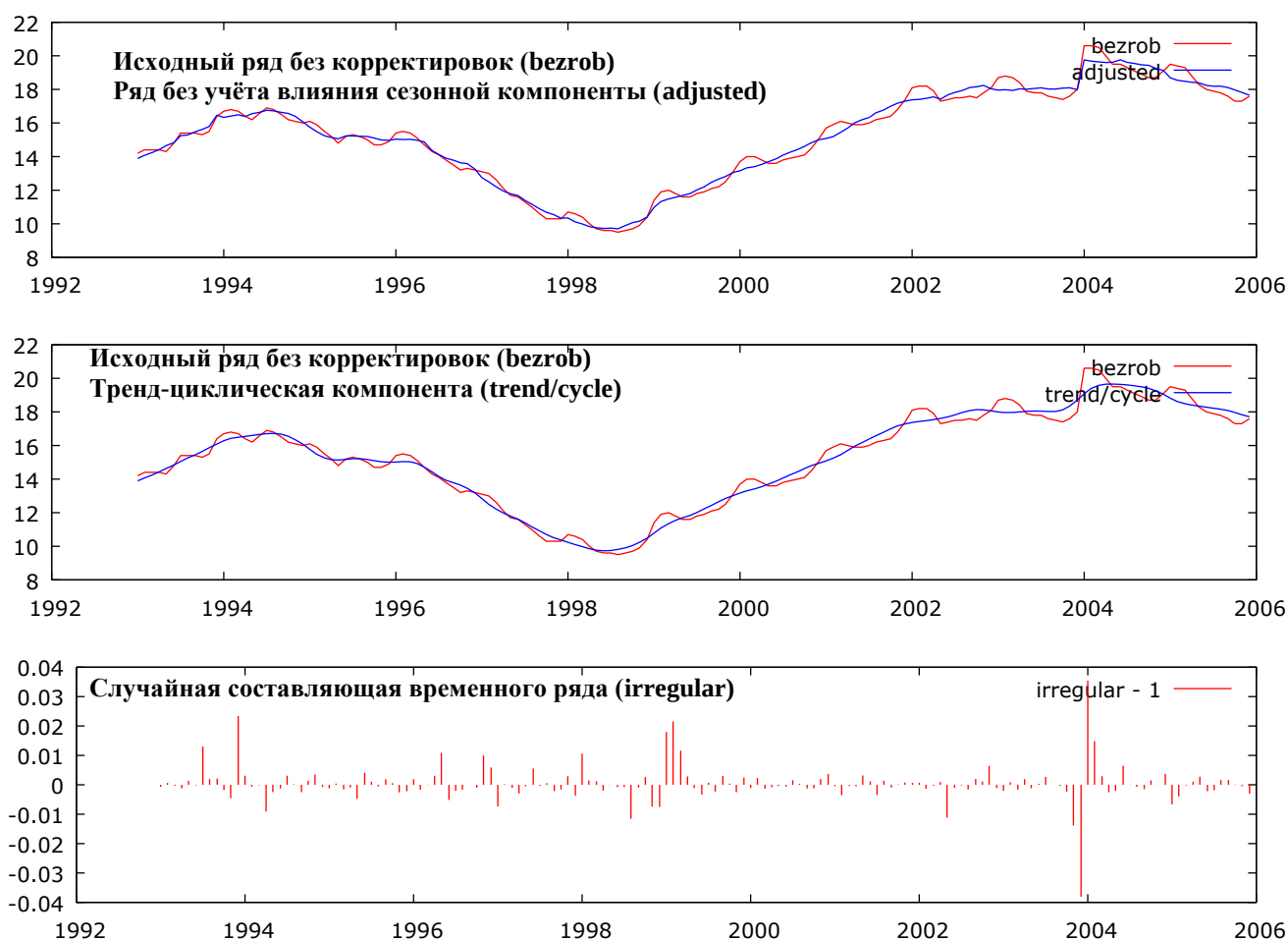


Рисунок 14 – Графическое отображение результатов применения процедуры X-12-ARIMA к временному ряду bezrob

D 10 Final seasonal factors (окончательное оценивание сезонных факторов)

From 1993.Jan to 2005.Dec (период рассмотрения)

Observations 156 (число наблюдений)

Seasonal filter 3 x 3 moving average

месяцы	Jan	Feb	Mar	Apr	May	Jun	
годы	Jul	Aug	Sep	Oct	Nov	Dec	AVGE
1993	102.3	102.3	101.1	99.9	97.6	99.8	
	101.0	100.7	99.5	98.0	98.2	99.6	100.0
1994	102.2	102.4	101.3	100.1	97.8	99.7	
	100.8	100.5	99.2	97.8	98.1	99.6	100.0
1995	102.2	102.7	101.7	100.4	98.3	99.8	
	100.5	99.9	98.7	97.3	98.0	99.5	99.9
1996	102.4	103.2	102.5	100.9	98.8	99.6	
	99.8	99.2	98.0	96.8	97.9	99.6	99.9
1997	102.8	104.0	103.4	101.3	99.2	99.3	
	99.1	98.5	97.3	96.4	97.6	99.7	99.9
1998	103.4	104.7	104.2	101.7	99.3	98.7	
	98.5	97.9	97.0	96.3	97.5	100.0	99.9
1999	103.9	105.1	104.5	101.9	99.3	98.3	
	98.2	97.6	97.0	96.4	97.5	100.4	100.0
2000	104.1	105.0	104.5	102.0	99.3	98.1	
	98.0	97.5	97.1	96.5	97.8	100.6	100.0
2001	104.1	104.7	104.2	102.0	99.3	98.2	
	98.0	97.5	97.2	96.5	97.8	100.6	100.0

2002	<u>104.1</u> 98.2	<u>104.6</u> 97.6	<u>104.2</u> 97.2	<u>101.9</u> 96.4	99.3 97.6	98.4 <u>100.3</u>	100.0
2003	<u>104.1</u> 98.4	<u>104.6</u> 97.6	<u>104.2</u> 97.1	<u>102.0</u> 96.2	99.4 97.3	98.6 <u>100.1</u>	100.0
2004	<u>104.2</u> 98.4	<u>104.7</u> 97.7	<u>104.4</u> 97.1	<u>102.0</u> 96.2	99.4 97.1	98.7 <u>99.9</u>	100.0
2005	<u>104.3</u> 98.4	<u>104.6</u> 97.8	<u>104.4</u> 97.2	<u>102.0</u> 96.3	99.5 97.0	98.7 <u>99.7</u>	100.0
AVGE	<u>103.4</u> 99.0	<u>104.0</u> 98.5	<u>103.4</u> 97.6	<u>101.4</u> 96.7	99.0 97.7	98.9 <u>100.0</u>	

Table Total- 15594.69 Mean- 99.97 Std. Dev.- 2.54
Min - 96.24 Max - 105.08

D 10.A Final seasonal component forecasts (прогноз сезонной компоненты на 2006г.)

From 2006.Jan to 2006.Dec (период рассмотрения)

Observations 12 (число наблюдений)

	Jan Jul	Feb Aug	Mar Sep	Apr Oct	May Nov	Jun Dec	AVGE
2006	<u>104.4</u> 98.4	<u>104.6</u> 97.9	<u>104.4</u> 97.2	<u>102.0</u> 96.3	99.5 97.1	98.7 <u>99.6</u>	100.0

D 11 Final seasonally adjusted data (данные без учёта сезонности)

From 1993.Jan to 2005.Dec (период рассмотрения)

Observations 156 (число наблюдений)

	Jan Jul	Feb Aug	Mar Sep	Apr Oct	May Nov	Jun Dec	TOTAL
1993	14. 15.	14. 15.	14. 15.	14. 16.	15. 16.	15. 16.	180.
1994	16. 17.	16. 17.	16. 17.	16. 17.	17. 16.	17. 16.	198.
1995	16. 15.	15. 15.	15. 15.	15. 15.	15. 15.	15. 15.	183.
1996	15. 14.	15. 14.	15. 14.	15. 14.	15. 14.	14. 13.	172.
1997	13. 11.	12. 11.	12. 11.	12. 11.	12. 11.	12. 10.	138.
1998	10. 10.	10. 10.	10. 10.	10. 10.	10. 10.	10. 10.	120.
1999	11. 12.	11. 12.	11. 12.	12. 13.	12. 13.	12. 13.	144.
2000	13. 14.	13. 14.	13. 14.	14. 15.	14. 15.	14. 15.	168.
2001	15. 16.	15. 17.	15. 17.	16. 17.	16. 17.	16. 17.	195.
2002	17. 18.	17. 18.	17. 18.	18. 18.	17. 18.	18. 18.	213.
2003	18. 18.	18. 18.	18. 18.	18. 18.	18. 18.	18. 18.	216.
2004	20. 20.	20. 20.	20. 19.	20. 19.	20. 19.	20. 19.	234.
2005	19. 18.	19. 18.	18. 18.	18. 18.	18. 18.	18. 18.	219.
AVGE	15. 15.	15. 15.	15. 15.	15. 15.	15. 15.	15. 15.	

Table Total- 2379.60 Mean- 15.25 Std. Dev.- 2.80
Min - 9.70 Max - 19.77

D 12 Final trend cycle (составляющая тренд-цикл)

From 1993.Jan to 2005.Dec (период рассмотрения)

Observations 156 (число наблюдений)

Trend filter 9-term Henderson moving average

I/C ratio 0.30

	Jan Jul	Feb Aug	Mar Sep	Apr Oct	May Nov	Jun Dec	TOTAL
1993	14. 15.	14. 15.	14. 15.	14. 16.	15. 16.	15. 16.	179.
1994	16. 17.	16. 17.	16. 17.	17. 17.	17. 16.	17. 16.	198.
1995	16. 15.	15. 15.	15. 15.	15. 15.	15. 15.	15. 15.	183.
1996	15. 14.	15. 14.	15. 14.	15. 14.	15. 13.	14. 13.	171.
1997	13. 11.	12. 11.	12. 11.	12. 11.	12. 11.	12. 10.	138.
1998	10. 10.	10. 10.	10. 10.	10. 10.	10. 10.	10. 10.	120.
1999	11. 12.	11. 12.	11. 12.	12. 13.	12. 13.	12. 13.	143.
2000	13. 14.	13. 14.	13. 14.	14. 15.	14. 15.	14. 15.	168.
2001	15. 16.	15. 17.	15. 17.	16. 17.	16. 17.	16. 17.	195.
2002	17. 18.	17. 18.	17. 18.	18. 18.	18. 18.	18. 18.	213.
2003	18. 18.	18. 18.	18. 18.	18. 18.	18. 18.	18. 19.	217.
2004	19. 20.	19. 20.	20. 19.	20. 19.	20. 19.	20. 19.	233.
2005	19. 18.	19. 18.	18. 18.	18. 18.	18. 18.	18. 18.	219.
AVGE	15. 15.	15. 15.	15. 15.	15. 15.	15. 15.	15. 15.	
Table Total-	2378.84		Mean-	15.25	Std. Dev.-	2.80	
			Min -	9.73	Max -	19.66	

D 13 Final irregular component (случайная составляющая)

From 1993.Jan to 2005.Dec (период рассмотрения)

Observations 156 (число наблюдений)

	Jan Jul	Feb Aug	Mar Sep	Apr Oct	May Nov	Jun Dec	S.D.
1993	99.9 101.3	100.1 100.2	100.0 100.2	99.9 99.8	100.1 99.5	100.0 102.3	0.8
1994	100.3 100.3	99.9 100.0	100.0 99.8	99.1 100.1	99.8 100.3	99.9 99.9	0.3
1995	99.9 100.1	100.0 99.9	99.8 100.2	99.9 100.1	99.5 99.7	100.4 99.8	0.2
1996	100.2 99.8	99.8 99.8	100.0 100.0	100.3 99.9	101.1 101.0	99.5 100.6	0.5
1997	99.3 100.0	100.0 100.1	99.9 99.8	99.7 99.8	99.9 100.3	100.6 99.6	0.3
1998	101.1 99.9	100.1 98.8	100.1 99.9	99.8 100.3	100.0 99.3	99.9 99.2	0.6

1999	101.8	102.2	101.2	100.3	99.9	99.7	
	100.1	99.8	100.3	100.0	99.7	100.2	0.9
2000	99.9	100.2	99.9	99.9	100.0	99.9	
	100.2	100.0	99.9	99.9	100.2	100.4	0.2
2001	100.0	99.6	100.0	99.9	100.3	100.1	
	99.7	100.1	99.9	100.0	100.1	100.1	0.2
2002	100.1	99.9	100.0	100.1	98.9	99.9	
	100.0	99.8	100.2	100.1	100.6	99.9	0.4
2003	99.8	100.1	99.8	100.2	99.9	100.0	
	100.3	100.0	100.0	99.8	98.6	96.2	1.2
2004	103.5	101.5	100.3	99.7	99.8	100.6	
	100.0	99.9	99.9	100.2	100.0	100.4	1.1
2005	99.3	99.6	100.0	100.1	100.3	99.8	
	99.8	100.2	100.2	100.0	100.0	99.7	0.3
S.D.	1.2	0.7	0.3	0.3	0.5	0.3	
	0.4	0.3	0.2	0.1	0.6	1.3	
Table Total-	15605.70	Mean-	100.04	Std. Dev.-	0.63		
		Min -	96.19	Max -	103.54		

Для использования процедуры TRAMO/SEATS необходимо выбрать пункт **TRAMO analysis** меню **Variable** предварительно выбрав переменную **bezrob** в открытом наборе данных. Затем в открывшемся окне на закладке **OUTPUT** отметить флажками опции **Seasonally adjusted series**, **Trend/cycle**, **Irregular**, **Generate Graph**, как и в процедуре X-12-ARIMA и нажать кнопку ОК (рисунок 15). Результаты применения процедуры TRAMO/SEATS не имеют кардинальных отличий (сходны) с результатам X-12-ARIMA, их графическое отображение представлено на рисунке 16.

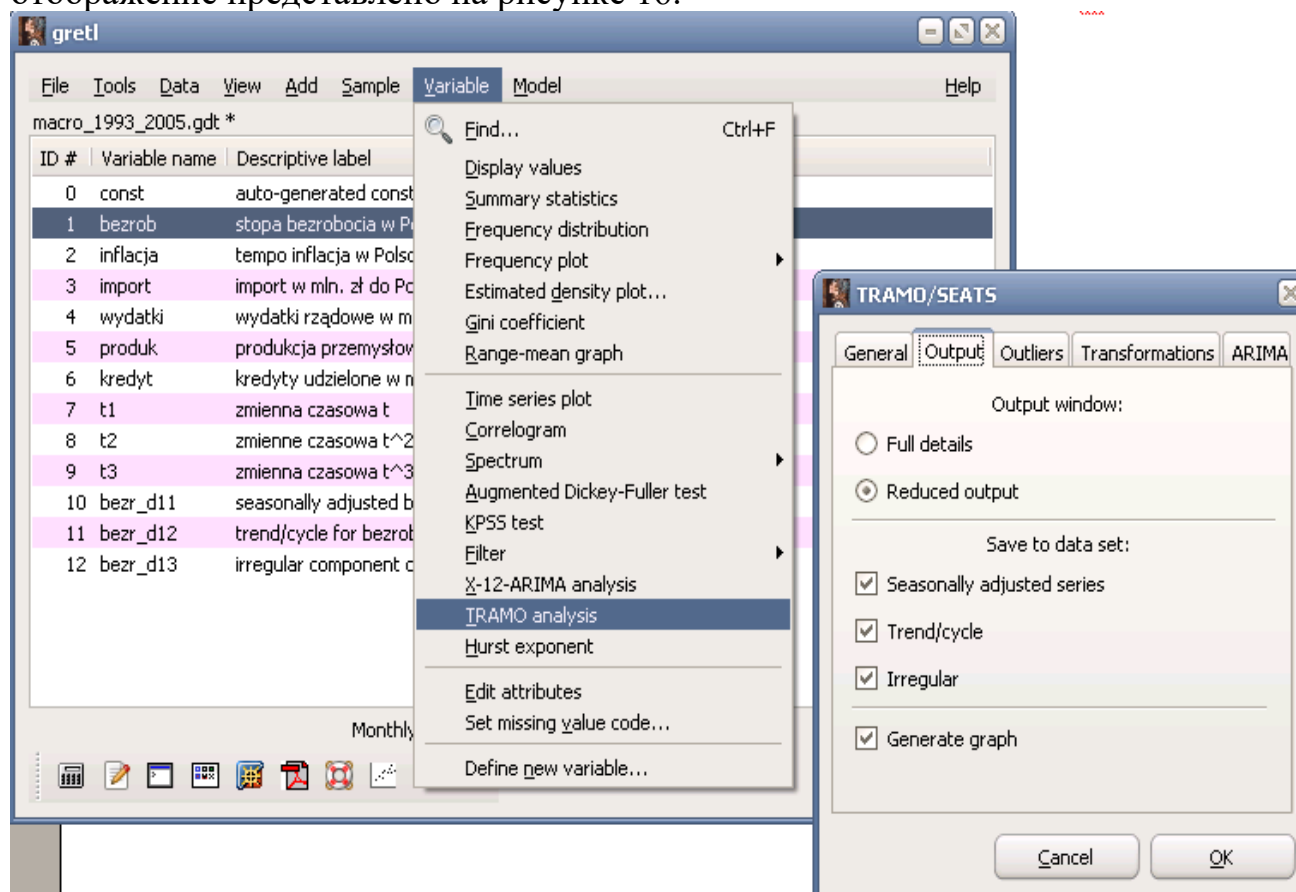


Рисунок 15 – Использование процедуры TRAMO/SEATS

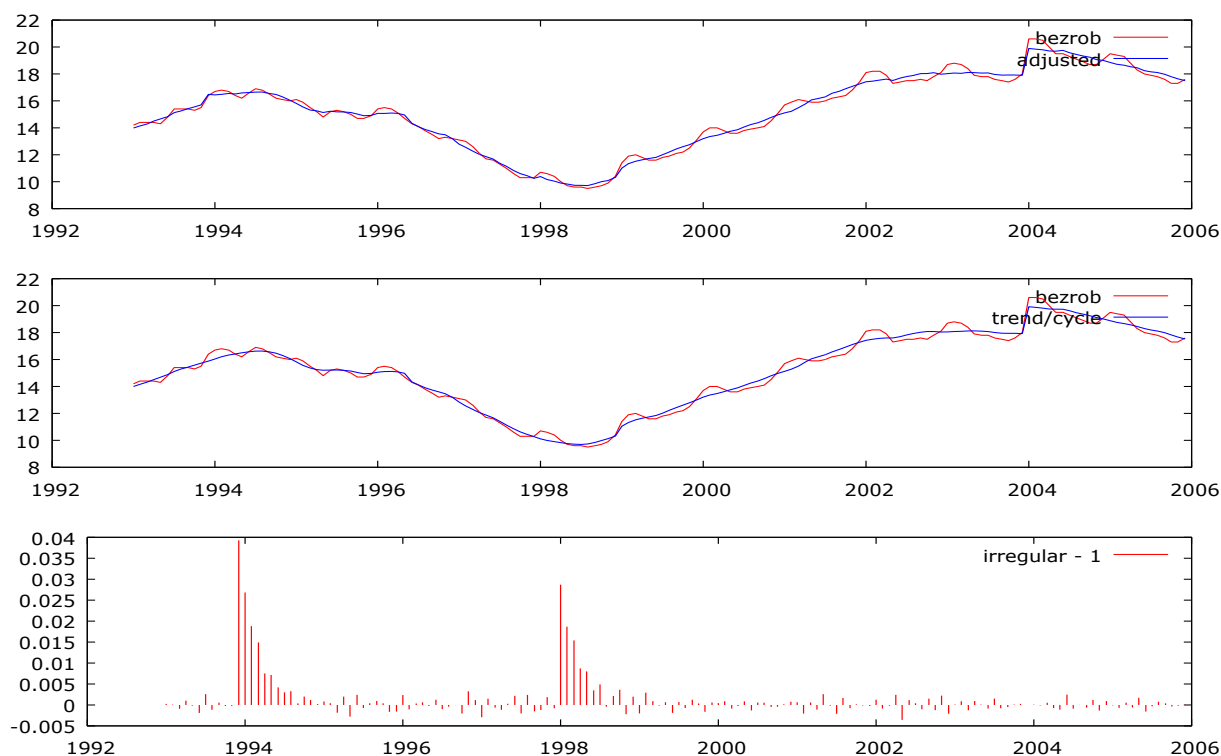


Рисунок 16 – Графическое отображение результатов применения процедуры TRAMO/SEATS к временному ряду bezrob

По результатам применения обеих процедур можно сделать следующие выводы:

1. Тренд-циклическая компонента оказывает существенное влияние на формирование уровней ряда, который демонстрирует циклический поведение. Можно выделить цикл с 10-летним периодом, где пики безработицы приходятся на 1994 и 2004 годы, а спад на 1998 год.

Рассмотрим основные факторы, порождающие циклические эффекты.

Рост уровня безработицы в период 1999-2003г. вызван преимущественно замедлением экономического роста в данный период, в конце которого был достигнут максимальный уровень 20%. Однако с вступлением Польши в Евросоюз в 2004 году удалось возобновить экономический рост, который сопровождался ростом прямых иностранных инвестиций, а также экспорта польских товаров и, следовательно, числом новых рабочих мест. Однако, по сей день в Польше сохраняется самый высокий уровень безработицы из всех стран Евросоюза. Отчасти это объясняется массовым трудоустройством польского населения в других странах Евросоюза с момента вступления в него, а отчасти тем, что рынок труда не может удовлетворить новое поколение – результат бума рождаемости, который наблюдался в период объявленного в 1981 году военного положения.

Трендовая компонента демонстрирует незначительную изменчивость, имеет плавную положительную динамику.

2. Временной ряд также существенно подвержен эффекту сезонности, который проявляется с периодом в один год. Сезонная компонента имеет наибольшие значения и оказывает наибольшее влияние на формирование уровней ряда в декабре, январе, феврале, марте и апреле, т.е. на увеличение уровня безработицы в данный период. Наименьшее влияние данный фактор оказывает в мае, июне, июле, августе, сентябре, октябре и ноябре. Данные сезонные эффекты могут отчасти объясняться привлечением дополнительных кадров для выполнения сезонных работ.

3. Влияние случайной составляющей (различных несистематических факторов) незначительно - максимальное значение составило 0,04 в 1994 году.

3.3. Пример анализа сезонности с применением коррелограммы

Проведём дальнейший анализ выявленной в предыдущем примере сезонности ряда значений уровня безработицы в Польше в 1993-2005 гг., используя метод коррелограммы. Предварительно необходимо щелчком мыши выбрать переменную *bezrob* в списке открытого набора данных *macro_1993_2005*.

В системе Gretl функции автокорреляции и частной автокорреляции можно оценить с применением команды меню **Variable/Correlogram** и указав в появившемся окне значение максимального периода запаздывания (лага $p_{\max}$), который не должен превышать 15-20% длины ряда (примем максимальный лаг 28 в рассматриваемом примере) (рисунок 17).

После нажатия кнопки ОК получим два окна результатов расчёта: графическое и текстовое (рисунок 18 и 19).

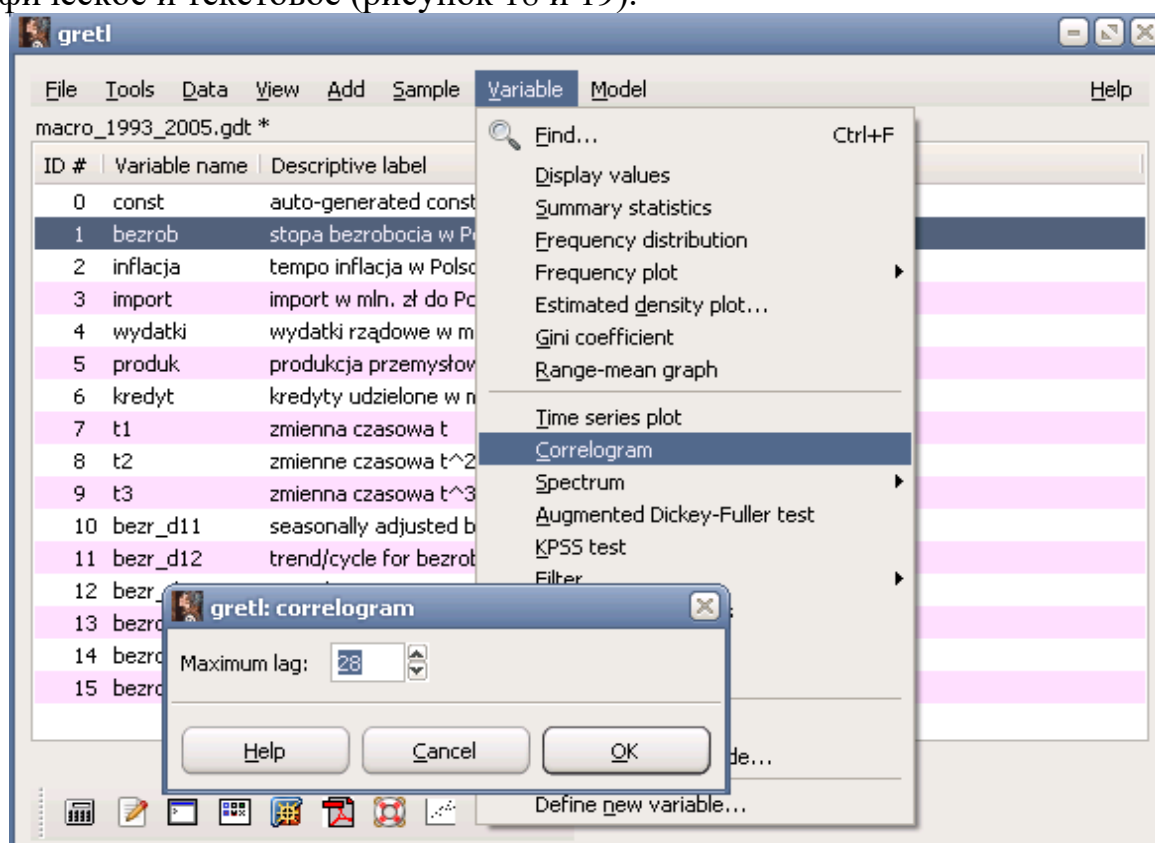


Рисунок 17 – Построение функций ACF и PACF

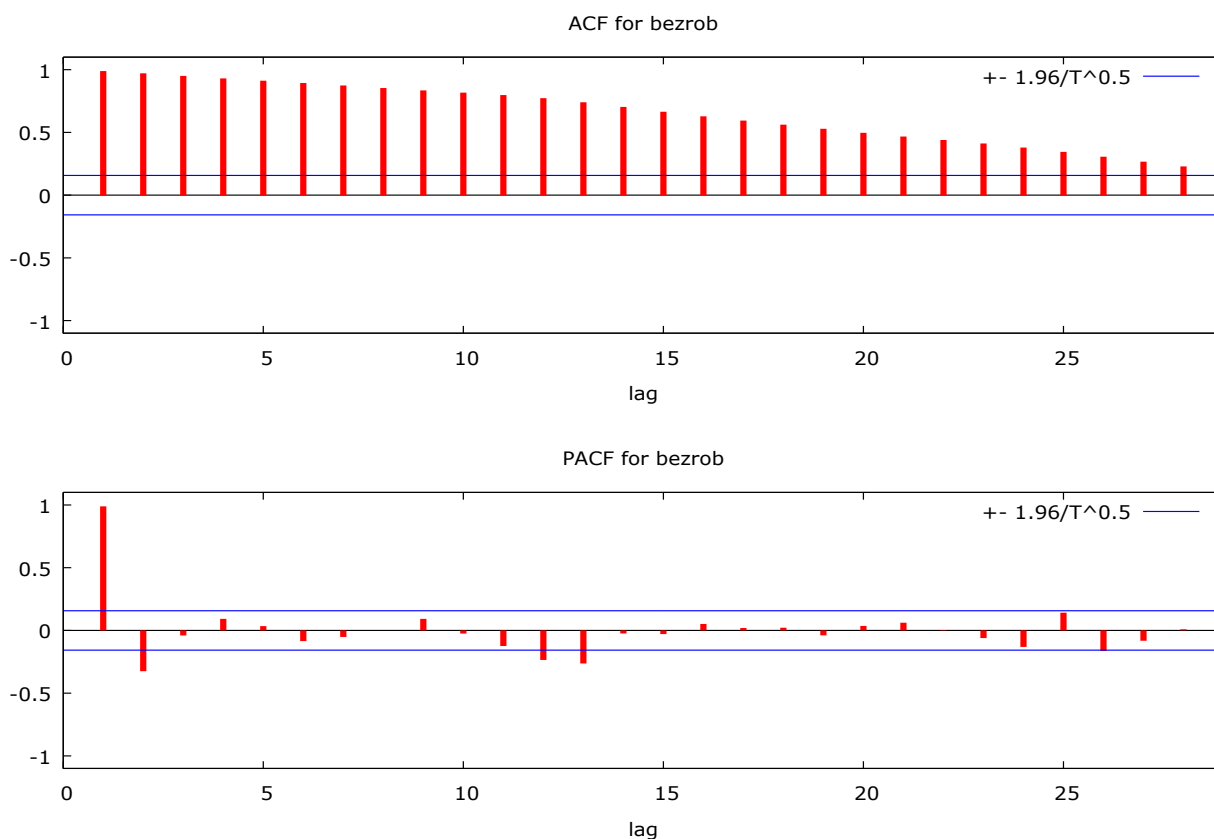


Рисунок 18 – Графики функций ACF и PACF ряда bezrob

gretl: correlogram

Autocorrelation function for bezrob

LAG	ACF		PACF		Q-stat.	[p-value]
1	0.9890	***	0.9890	***	155.5387	[0.000]
2	0.9710	***	-0.3242	***	306.4449	[0.000]
3	0.9502	***	-0.0401		451.9073	[0.000]
4	0.9304	***	0.0923		592.2869	[0.000]
5	0.9123	***	0.0346		728.1516	[0.000]
6	0.8939	***	-0.0846		859.4514	[0.000]
7	0.8742	***	-0.0520		985.8859	[0.000]
8	0.8537	***	-0.0002		1107.2714	[0.000]
9	0.8348	***	0.0922		1224.1138	[0.000]
10	0.8168	***	-0.0248		1336.7586	[0.000]
11	0.7976	***	-0.1242		1444.9081	[0.000]
12	0.7730	***	-0.2354	***	1547.1877	[0.000]
13	0.7402	***	-0.2633	***	1641.6371	[0.000]
14	0.7032	***	-0.0243		1727.4640	[0.000]
15	0.6646	***	-0.0289		1804.6752	[0.000]
16	0.6282	***	0.0517		1874.1548	[0.000]
17	0.5942	***	0.0192		1936.7630	[0.000]
18	0.5616	***	0.0219		1993.0987	[0.000]
19	0.5287	***	-0.0389		2043.3886	[0.000]
20	0.4966	***	0.0359		2088.0808	[0.000]
21	0.4671	***	0.0614		2127.9142	[0.000]
22	0.4396	***	0.0032		2163.4589	[0.000]
23	0.4116	***	-0.0608		2194.8605	[0.000]
24	0.3791	***	-0.1310		2221.7001	[0.000]
25	0.3447	***	0.1418	*	2244.0607	[0.000]
26	0.3064	***	-0.1636	**	2261.8556	[0.000]
27	0.2667	***	-0.0822		2275.4510	[0.000]
28	0.2289	***	0.0090		2285.5419	[0.000]

Рисунок 19 – Окно результатов вычисления ACF и PACF ряда bezrob

Оценочные результаты для ряда значений уровня безработицы в Польше, представленные на рисунках 18 и 19, показывают чистую зависимость между наблюдениями ряда (функция PACF), разнесёнными на 1, 2, 12 и 13 периодов, что позволяет оценить порядок запаздывания процесса p для модели авторегрессии $AR(n) = AR(13)$.

Тест значимости коэффициента автокорреляции, называемый тестом Квенилле (Quenouille), свидетельствует, что если оцениваемый расчётный коэффициент больше критического значения $\pm 1,96/T^{0,5}$ (T -число наблюдений ряда), то присутствует существенная зависимость между процессами для соответствующего лага (запаздывания между процессами). На коррелограмме горизонтальной линией отмечается граница доверительного интервала стандартной погрешности этого коэффициента. За пределы данной границы выходят 1, 2, 12 и 13 значения PACF.

Следовательно, можно сделать вывод о зависимости каждого значения ряда от предыдущего (лаг=1) и от предшествующего предыдущему (лаг=2), а также от значения одноимённого месяца прошлого года (лаг=12) и предшествующего ему (лаг=13), т.е. существует сезонная составляющая в рассматриваемом ряду данных с длиной цикла 12 месяцев.

Например, значение ряда в марте зависит от значения ряда в феврале (1) и январе (2) этого года, а также от значений ряда в марте (12) и феврале (13) прошлого года. Аналогичная зависимость имеет место для каждого значения ряда.

3.4. Пример применения метода авторегрессии

Хотя уравнения авторегрессии $AR(n)$ вычисляются и более эффективны для описания и прогнозирования стационарных процессов, данный метод также применим для нестационарных процессов, особенно, если не стационарность носит однородный характер.

На основе полученных в предыдущем примере данных при анализе PACF построим авторегрессионную модель $AR(13)$. Для упрощения анализа сделаем допущение, что нарушена только одна предпосылка обычного метода наименьших квадратов – имеет место автокорреляция остатков высших порядков.

В данном случае воспользуемся встроенной в Gretl обобщённой процедурой Кохрейна-Оркотта (Generalized Cochrane-Orcutt Iterative procedure) для оценки параметров модели $AR(13)$. Для оценивания параметров с применением этого метода необходимо выбрать команду **Model\Time series\Autoregressive estimation** (рисунок 20).

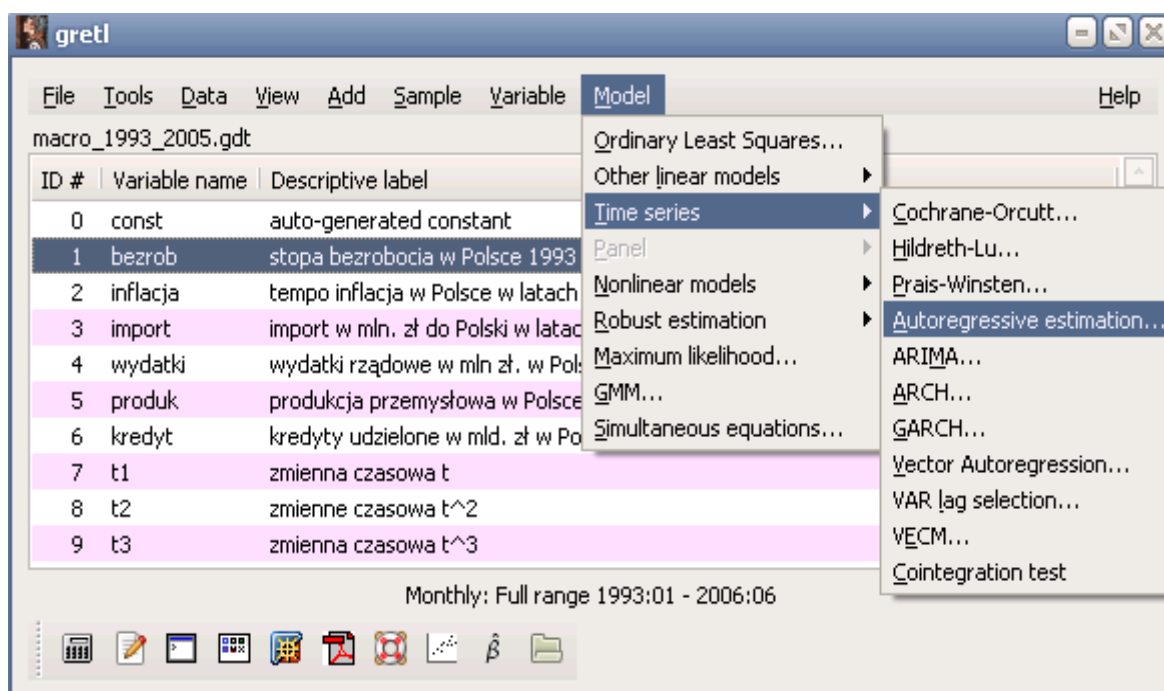


Рисунок 20 – Построение модели авторегрессии

В открывшемся окне (рисунок 21) введём значения зависимой переменной (Dependent variable) – bezrob – при помощи кнопки Choose, перечислим лаги модели List of AR lags 1,2,12,13, которым соответствуют значимые коэффициента частной автокорреляции (рассчитанные в предыдущем примере). В качестве объясняющих переменных введём лаговые значения зависимой переменной: bezrob-1, bezrob-2, bezrob-12, bezrob-13, нажав кнопку LAGS (рисунок 21). В появившемся окне флажками отметим опции Lags of dependent variable и Specific lags, введя лаги 1,2,12,13, нажмём кнопку ОК в обоих окнах.

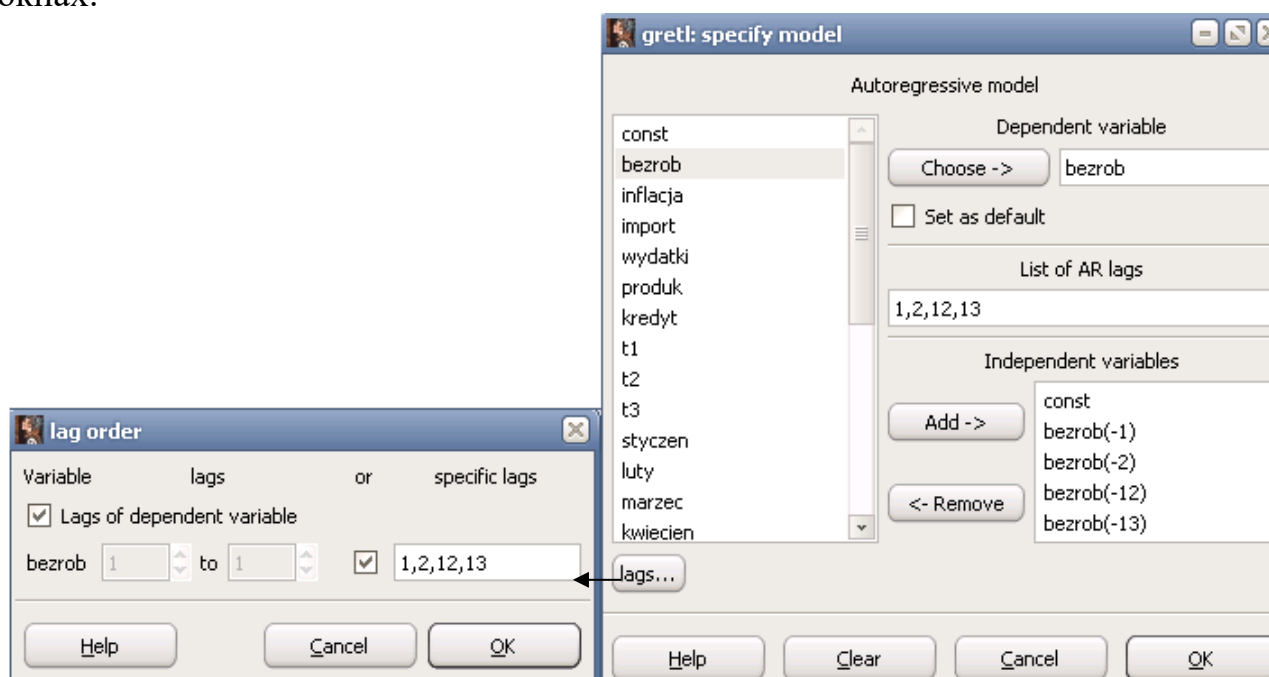
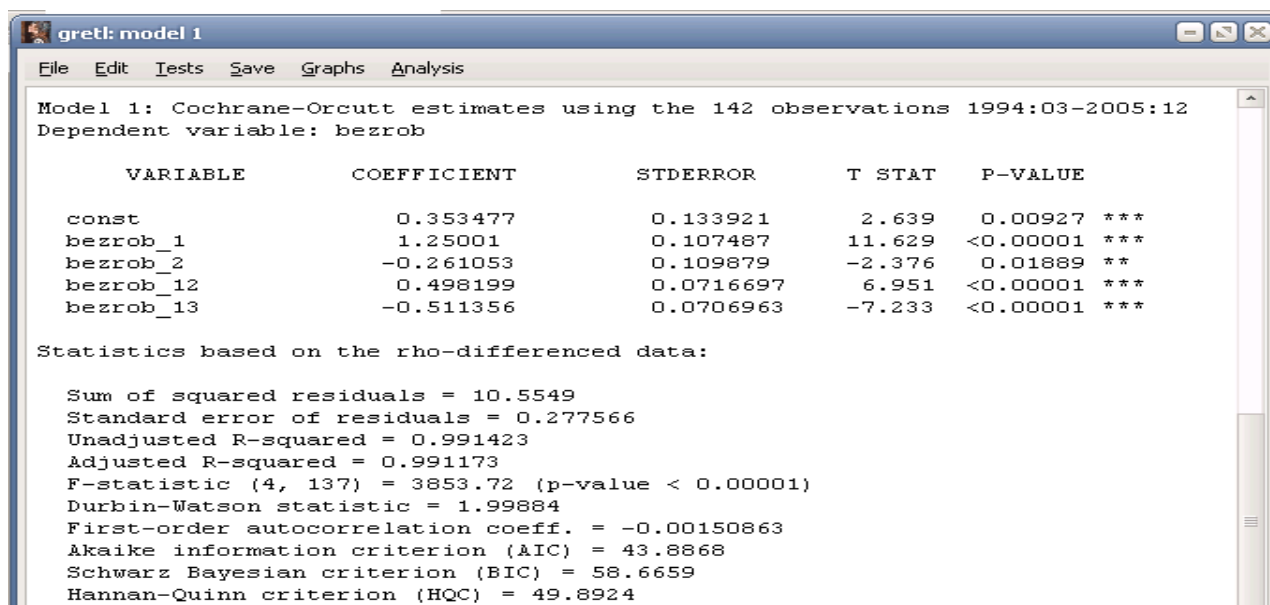


Рисунок 21 – Спецификация авторегрессионной модели

По данным окна результатов моделирования (рисунок 22) отметим, что полученная модель $y_t = 0,35 + 1,25y_{t-1} - 0,26y_{t-2} + 0,5y_{t-12} - 0,51y_{t-13} + u_t$ является адекватной по F-критерию ($p\text{-value} < 0,05$), влияние каждой лаговой переменной существенно по t-критерию ($p\text{-value} < 0,05$) для уровня значимости 5%. Наличие больших значений лагов подтверждает существование выявленной ранее сезонной компоненты (длина цикла - год).



Model 1: Cochrane-Orcutt estimates using the 142 observations 1994:03-2005:12
Dependent variable: bezrob

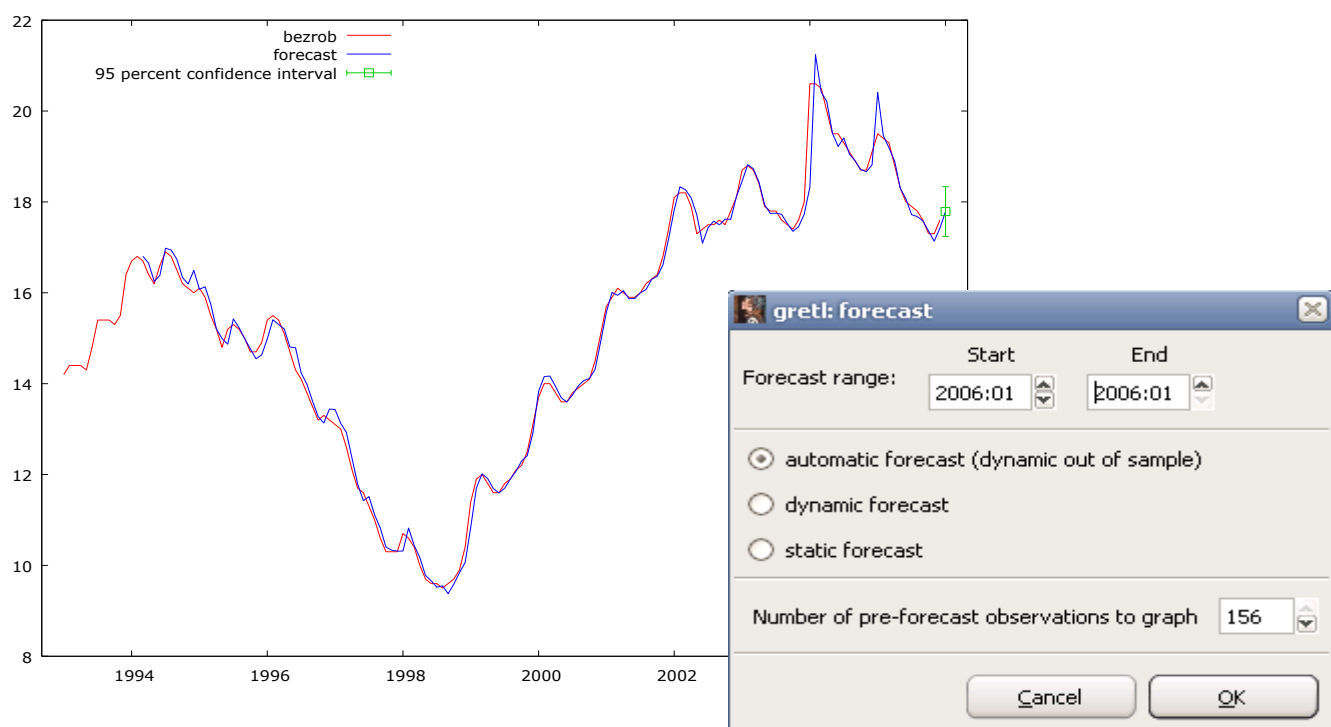
VARIABLE	COEFFICIENT	STDERROR	T STAT	P-VALUE
const	0.353477	0.133921	2.639	0.00927 ***
bezrob_1	1.25001	0.107487	11.629	<0.00001 ***
bezrob_2	-0.261053	0.109879	-2.376	0.01889 **
bezrob_12	0.498199	0.0716697	6.951	<0.00001 ***
bezrob_13	-0.511356	0.0706963	-7.233	<0.00001 ***

Statistics based on the rho-differenced data:

Sum of squared residuals = 10.5549
Standard error of residuals = 0.277566
Unadjusted R-squared = 0.991423
Adjusted R-squared = 0.991173
F-statistic (4, 137) = 3853.72 (p-value < 0.00001)
Durbin-Watson statistic = 1.99884
First-order autocorrelation coeff. = -0.00150863
Akaike information criterion (AIC) = 43.8868
Schwarz Bayesian criterion (BIC) = 58.6659
Hannan-Quinn criterion (HQC) = 49.8924

Рисунок 22 – Окно результатов моделирования с применением метода авторегрессии

Используем авторегрессионную модель для получения прогноза уровня безработицы в январе 2006 года. Для этого обратимся к команде **Analysis\Forecasts** окна результатов моделирования (рисунок 22). Выберем период 2006:1 и число наблюдений ряда 156 и нажмём кнопку ОК в открывшемся окне и получим прогнозное значение уровня безработицы 17.7878 (рисунок 23). Данный прогноз менее точен, чем полученный при прогнозировании с использованием модели полиномиального тренда четвертого порядка, поскольку ряд не стационарен. Можно сделать вывод, что для случая нестационарных рядов данный метод необходимо сочетать с другими методами анализа, например с анализом тренда и т.д.



For 95% confidence intervals, $t(137, .025) = 1.977$

Obs	Bezrob	Prediction	std. error	95% confidence interval
наблюдение	уровень	прогноз	Стандартная ошибка	Доверительный интервал
	безработицы			
2006:01	undefined	17.7878	0.277566	(17.2389, 18.3367)

Рисунок 23 – Прогнозирование временного ряда bezrob с использованием авторегрессионной модели

3.5. Пример применения метода спектрального (Фурье) анализа

Для оценки рассмотренного выше ряда значений уровня безработицы в Польше (1993-2005г.) с позиций спектрального анализа определим функцию спектральной плотности (сглаженную периодограмму) для выявления скрытых периодичностей.

В данном случае использование этого метода для оценки заведомо нестационарного процесса может привести лишь к определению общей формы спектральной плотности, при этом детали (спектральные особенности) окажутся размытыми.

Перед началом процедуры сократим временной ряд до 10 лет **Sample/Set Range**, введём период 1996:01 – 200:12.

Для оценивания функции спектральной плотности в меню выбирается опция **Variable/Spectrum/Sample Periodogram** или **Variable/Spectrum/Bartlett lag window** (рисунок 24 и 25), если используется метод сглаживания Бартлетта. Результаты оценки для первой процедуры представлены в текстовом и графическом виде (рисунок 26).

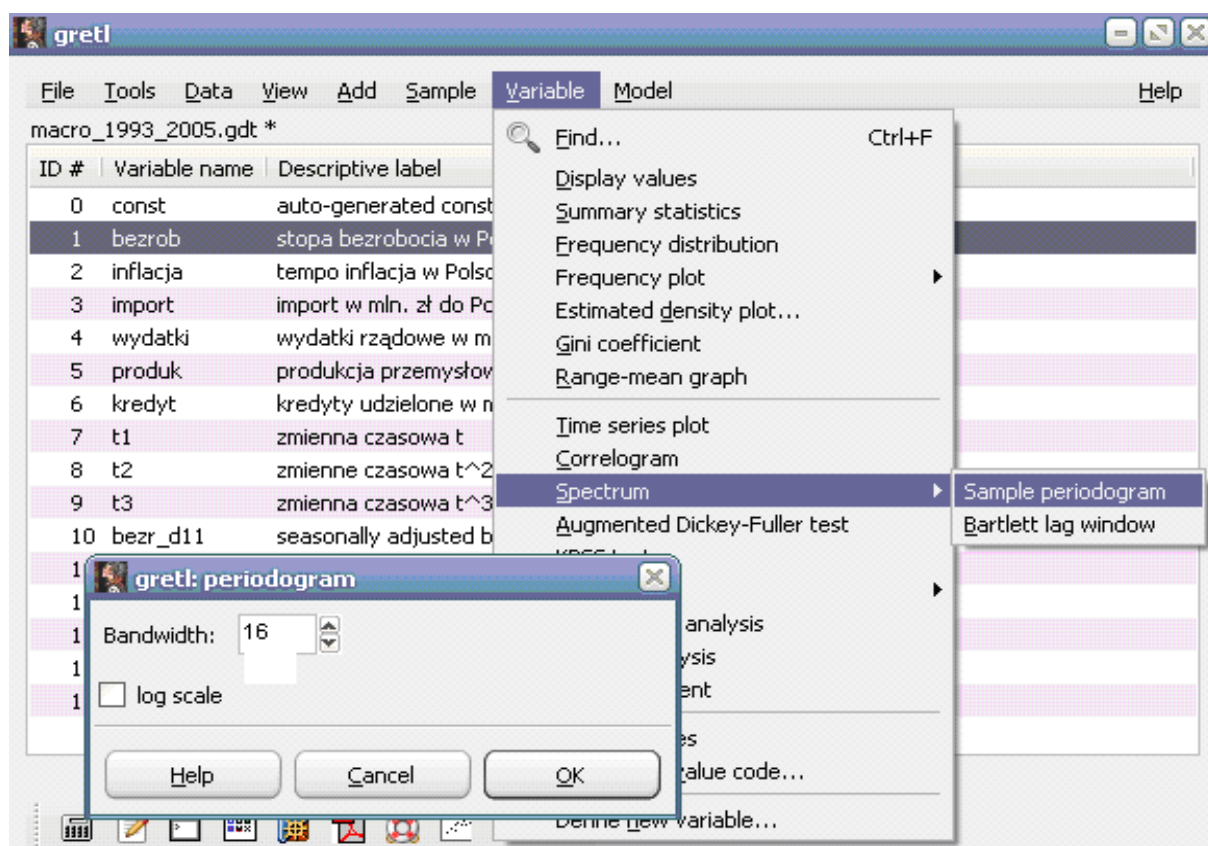


Рисунок 24 – Оценивание функции спектральной плотности

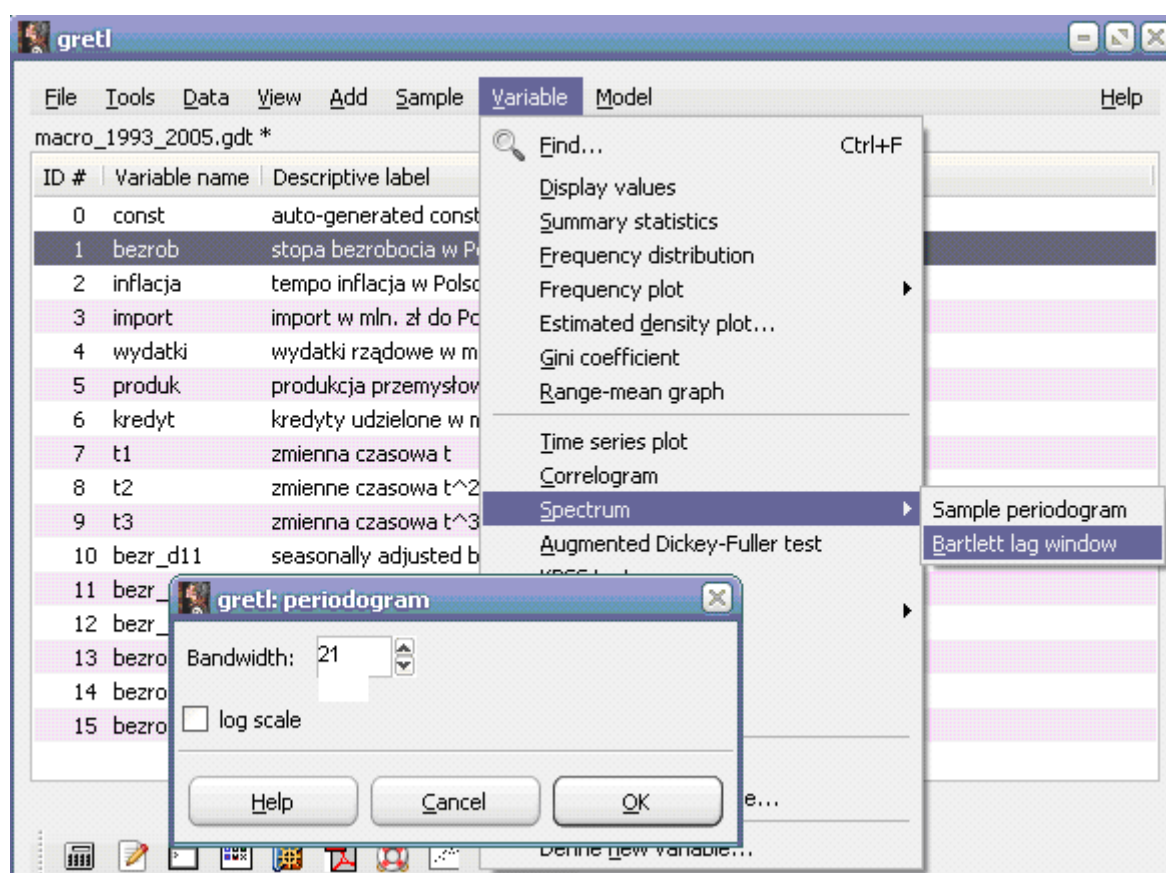


Рисунок 25 – Оценивание функции спектральной плотности с использованием метода Бартлетта (Bartlett).

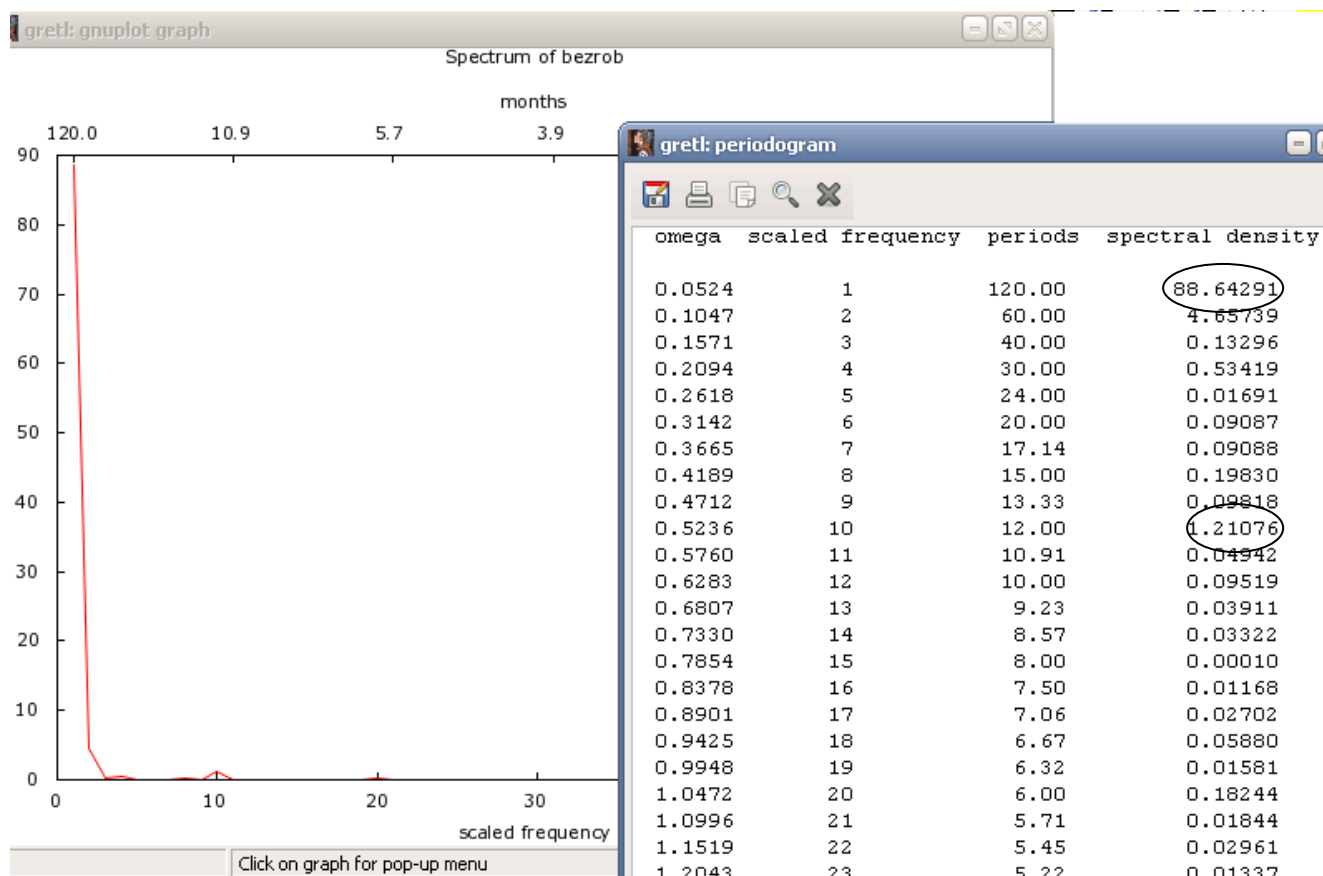


Рисунок 26 – Результаты оценивания функции спектральной плотности

Основные показатели окна результатов оценки (рисунок 26):

- omega- круговые частоты гармоник (функций синус и косинус);
- scaled frequency – номера частот гармоник (i);
- periods – периоды гармоник;
- spectral density- спектральная плотность.

Для данного ряда сглаженная периодограмма содержит резкий подъем в области низких частот, связанный с наличием детерминированной периодичности с очень большим (10-ти летним) периодом. Наличие эффекта сезонности проявляет себя в виде острого пика на 10-й частоте (периодичность с 12-ти месячным периодом), что подтверждает выводы предшествующего анализа данного ряда. Скрытых периодичностей не выявлено.

4. ПОРЯДОК ВЫПОЛНЕНИЯ ЛАБОРАТОРНОЙ РАБОТЫ

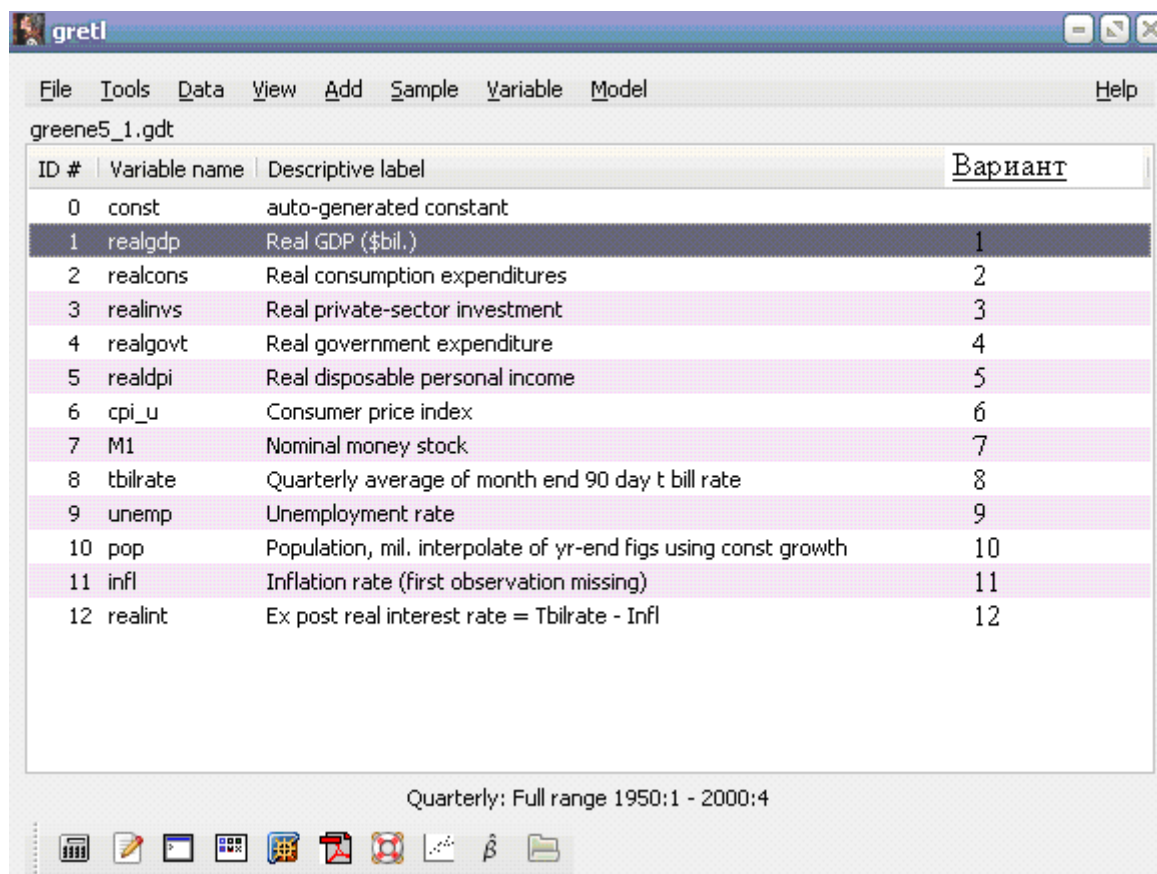
1. Выполнить рассмотренные в п.3.1.-3.5. примеры.
2. Открыть набор данных greene5_1.gdt на закладке Greene (File\Open Data\ Sample File), рисунок 27. Источник данных <http://www.economagic.com>

Просмотреть набор данных (View\Icon View...Data Set) и его описание (Data\Print Description). Просмотреть графическое отображение временного ряда значений одного из макроэкономического показателей - поквартальные данные 1950-2000г. - (Variable\Time series plot) согласно варианту. Для ввода переменной времени t обратиться к команде Add\Index variable.

3. Выбрать методы анализа временного ряда из ниже перечисленных:

- анализ тренда;
- декомпозиция временного ряда;
- метод коррелограммы;
- метод авторегрессии;
- спектральный анализ.

4. Применить выбранные методы анализа и объяснить полученные результаты.



ID #	Variable name	Descriptive label	Вариант
0	const	auto-generated constant	
1	realgdp	Real GDP (\$bil.)	1
2	realcons	Real consumption expenditures	2
3	realinvs	Real private-sector investment	3
4	realgovt	Real government expenditure	4
5	realdpi	Real disposable personal income	5
6	cpi_u	Consumer price index	6
7	M1	Nominal money stock	7
8	tbilrate	Quarterly average of month end 90 day t bill rate	8
9	unemp	Unemployment rate	9
10	pop	Population, mil. interpolate of yr-end figs using const growth	10
11	infl	Inflation rate (first observation missing)	11
12	realint	Ex post real interest rate = Tbilrate - Infl	12

Quarterly: Full range 1950:1 - 2000:4

Рисунок 27 – Исходные данные и варианты заданий

5. СОДЕРЖАНИЕ ОТЧЕТА О ВЫПОЛНЕНИИ ЛАБОРАТОРНОЙ РАБОТЫ

- 1) Название и цель работы.
- 2) Постановка задачи.
- 3) Этапы выполнения задачи в Gretl.
- 4) Выводы.

ЛАБОРАТОРНАЯ РАБОТА №6

АНАЛИЗ СИСТЕМ ОДНОВРЕМЕННЫХ ЭКОНОМЕТРИЧЕСКИХ УРАВНЕНИЙ В СРЕДЕ GRETL 1.7.1

1. ЦЕЛЬ РАБОТЫ

Целью данной работы является получение практических навыков описания сложных экономических процессов и объектов управления с помощью систем взаимосвязанных (одновременных) уравнений.

2. ТЕОРЕТИЧЕСКИЙ РАЗДЕЛ

Многие экономические взаимосвязи допускают моделирование одним уравнением. При этом предполагается, что между независимыми переменными $x_1 \dots x_m$ и зависимой переменной y существует только прямая связь: $x_i \rightarrow y, i=1 \dots m$. В такой ситуации зависимая переменная y не оказывает никакого влияния на переменные, входящие в правую часть модели, которые в свою очередь можно изменять независимо друг от друга. В большинстве случаев для оценки таких моделей используется метод 1МНК.

Однако, описание сложного экономического процесса предполагает использование системы взаимосвязанных (одновременных) уравнений – simultaneous equations, структурная форма которой представлена формулой (1).

$$\begin{cases} y_1 = b_{12} * y_2 + b_{13} * y_3 + \dots + b_{1n} * y_n + a_{11} * x_1 + a_{12} * x_2 + \dots + a_{1m} * x_m + e_1; \\ y_2 = b_{21} * y_1 + b_{23} * y_3 + \dots + b_{2n} * y_n + a_{21} * x_1 + a_{22} * x_2 + \dots + a_{2m} * x_m + e_2; \\ \dots \\ y_n = b_{n1} * y_1 + b_{n2} * y_2 + \dots + b_{nn-1} * y_{n-1} + a_{n1} * x_1 + a_{n2} * x_2 + \dots + a_{nm} * x_m + e_n, \end{cases} \Leftrightarrow BY + AX = E, (1)$$

где a и b – структурные параметры модели;

y_i – эндогенная (зависимая) переменная, определяемая внутри модели ($i=1 \dots n$);
 x_i – предопределённая переменная: экзогенная (независимая) переменная, определяемая вне системы, или лаговая (запаздывающая) эндогенная.

В данной системе эндогенные переменные взаимосвязаны, одни и те же эндогенные переменные в одних уравнениях входят в левую часть системы, а в других – в правую, поэтому каждое уравнение не может рассматриваться самостоятельно и для нахождения его параметров традиционный 1МНК неприменим.

Для оценивания параметров структурной модели (1) используются следующие методы:

- косвенный метод наименьших квадратов (КМНК);
- двухшаговый метод наименьших квадратов (2МНК – Two-Stage Least Squares, *tsls*);
- трехшаговый метод наименьших квадратов (3МНК – Three-Stage Least Squares, *3sls*);
- метод максимального правдоподобия с полной информацией (ММП);

- метод максимального правдоподобия при ограниченной информации (ММП).

С позиции идентифицируемости структурные модели можно подразделить на три вида:

1. Точно идентифицируемые (все структурные параметры определяются однозначно, единственным образом; для решения системы используется КМНК, 2МНК или 3МНК), $D+1=N$;
2. Неидентифицируемые (нерешаемы, т.к. один или более параметров не могут быть определены), $D+1 < N$;
3. Сверхидентифицируемые (все структурные параметры определяются, но некоторые из них могут принимать одновременно несколько значений; для решения системы используется 2МНК или 3МНК), $D+1 > N$.

N – число эндогенных переменных в уравнении, D – число предопределенных переменных, отсутствующих в уравнении, но присутствующих в системе.

Система линейных функций эндогенных переменных от всех предопределенных переменных системы является приведенной формой модели (2).

$$\begin{cases} y_1 = p_{11} \cdot x_1 + p_{12} \cdot x_2 + \dots + p_{1m} \cdot x_m + v_1 \\ \dots\dots\dots \\ y_n = p_{n1} \cdot x_1 + p_{n2} \cdot x_2 + \dots + p_{nm} \cdot x_m + v_2, \end{cases} \Leftrightarrow Y = P \cdot X + V \quad (2)$$

где p_{ij} - коэффициенты приведенной формы модели.

Двухшаговый метод наименьших квадратов (2МНК) является наиболее общим и широко распространенным методом решения системы одновременных уравнений, поэтому в ряде компьютерных программ для её решения рассматривается лишь двухшаговый метод наименьших квадратов; пакет Gretl 1.7.1 содержит 2МНК и 3МНК методы.

Основная идея 2МНК заключается в том, что на основе приведенной формы модели (2) получают методом 1МНК для каждого i -го (сверхидентифицируемого или идентифицируемого) уравнения системы (1) теоретические значения эндогенных переменных, содержащихся в правой части уравнения $\hat{Y}_i$, формула (3).

$$\hat{Y}_i = X \cdot \hat{P}_i, \quad (3)$$

где X - матрица значений всех предопределённых переменных системы;

$\hat{Y}_i$ - матрица оценок эндогенных переменных в правой части i -го уравнения;

$\hat{P}_i$ – подматрица матрицы $\hat{P}$ оценок параметров приведенной формы (2), соответствующих эндогенным переменным, включённым в правую часть i -го структурного уравнения ($\hat{P}$ получено применением 1МНК к системе (2));

Затем, подставив $\hat{Y}_i$ вместо фактических значений Y_i в правой части уравнения, можно применить 1МНК к каждому уравнению структурной формы (1). Т.е. строятся 1МНК оценки структурных параметров $\vec{\beta}_i, \vec{\alpha}_i$ в регрессии (4).

$$\bar{y}_i = \hat{Y}_i \vec{\beta}_i + X_i \vec{\alpha}_i + \vec{\varepsilon}_i \quad (4) \text{ Формула}$$

(4) отражает каждое уравнение системы (1) после того как фактические значения

эндогенных переменных Y_i в правой части были заменены на их теоретические значения (оценки) $\hat{Y}_i$.

3. ОПИСАНИЕ СРЕДСТВ АНАЛИЗА СИСТЕМ ОДНОВРЕМЕННЫХ ЭКОНОМЕТРИЧЕСКИХ УРАВНЕНИЙ ПАКЕТА GRETl

Рассмотрим процедуру построения системы одновременных уравнений с помощью двухшагового метода наименьших квадратов (2МНК) в пакете Gretl.

Первый способ реализации 2МНК в Gretl осуществляется путём выбора пункта меню **Model\Other Linear Models\Two-Stage Least Squares....**

Второй способ – путём выбора пункта **Simultaneous Equations** меню **Model**.

В первом случае оценивание каждого уравнения происходит по очереди за один шаг. Во втором – вся система оценивается за одну операцию. Помимо метода 2МНК (TSLS) **Simultaneous Equations** также предусматривает другие методы оценки, такие как 3МНК (3SLS) и т.д.

Рассмотрим пример структурной модели, представляющей взаимозависимость между объёмом производства и занятостью в химической промышленности Польши 1962-1985гг. (ежегодные данные), модель (5). Для оценивания данной взаимозависимости будем использовать фактор инвестиций.

$$\begin{cases} y_1 = b_{12} * y_2 + a_{13} * x_3 + a_{14} * x_4 + e_1; \\ y_2 = b_{21} * y_1 + a_{21} + a_{22} * x_2 + a_{23} * x_3 + e_2; \end{cases} \quad (5)$$

y_1 - объём производства в химической промышленности за год ($produkt_t$);

y_2 - занятость в химической промышленности ($zatrud_t$);

x_2 – объём инвестиций в химическую промышленность за год ($inwest_t$);

x_3 - объём инвестиций предыдущего периода ($inwest_{t-1}$);

x_4 - объём производства предыдущего периода ($produkt_{t-1}$);

Вид данной системы одновременных уравнений – сверхидентифицируемая, т.к. первое уравнение сверх идентифицируемое ($H=2 < (D+1=3)$), второе – точно идентифицируемое ($H=2 = (D+1=2)$), поэтому можем применить 2МНК для оценки её параметров.

Откроем набор исходных данных **przemysl_chemiczny.gdt** (рисунок 1), выбрав пункт меню **File\Open Data\Sample file** и дважды щёлкнув левой кнопкой мыши по названию файла на закладке KUFEL; или создадим **przemysl_chemiczny.gdt**, перенеся данные таблицы 1 в файл **lab6.xls** и импортировав его в пакет Gretl. Для этого в меню Gretl выберем пункт **File\Open Data\Import\Excel**, в появившемся окне укажем номер строки и столбца начала таблицы Excel и нажмём кнопку ОК, в следующем окне нажмём кнопку YES, затем выберем тип данных временной ряд (time series), нажмём кнопку FORWARD, отметим периодичность данных Annual, нажмём кнопку FORWARD, введём год начала сбора данных 1962, нажмём кнопку FORWARD и ОК.

Сохраним созданный набор данных (рисунок 1) **File\Save Data** как файл **przemysl_chemiczny.gdt** на рабочем столе. Данный файл также доступен на сайте www.kufel.torun.pl

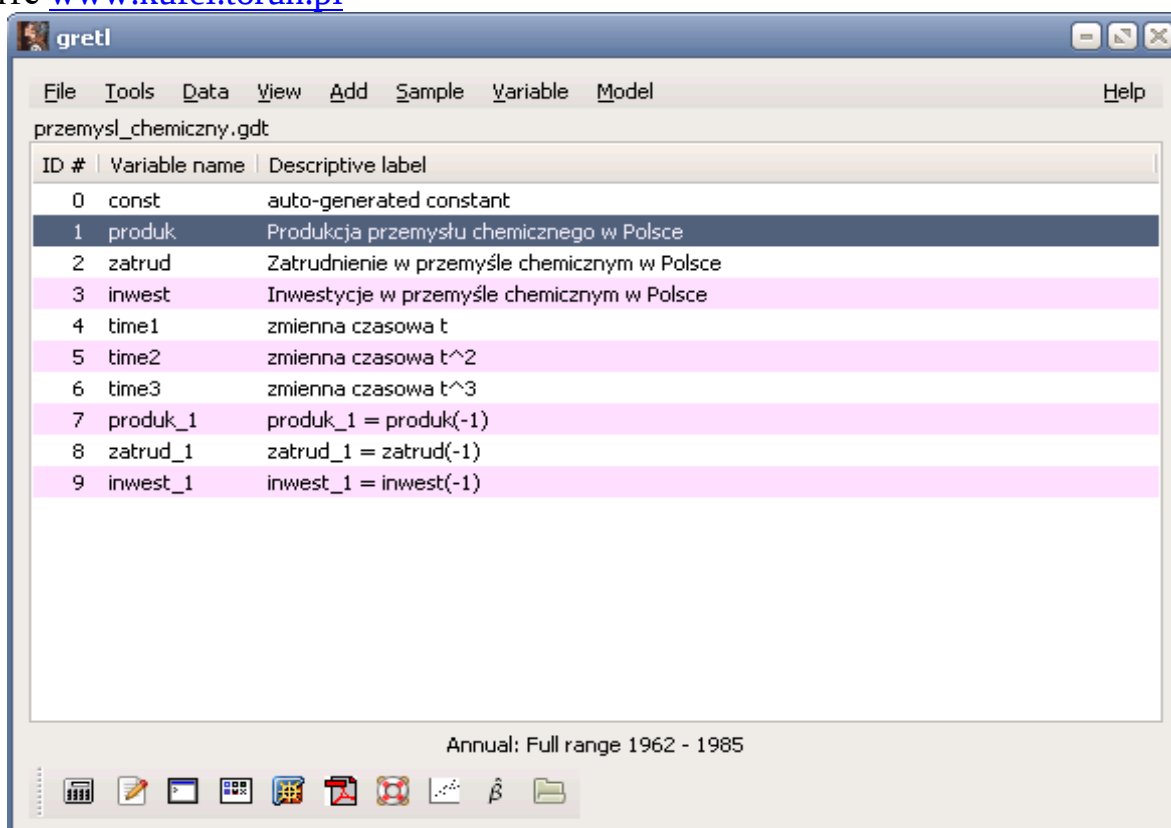


Рисунок 1– Набор данных przemysl_chemiczny.gdt

Таблица 1 – Показатели химической промышленности в Польше, 1962-1986гг.

Obs	produk	zatrud	inwest	produk_1	zatrud_1	inwest_1
1962	176	301	52	163	280	51
1963	185	317	52	176	301	52
1964	188	342	53	185	317	52
1965	198	356	53	188	342	53
1966	199	373	56	198	356	53
1967	201	381	59	199	373	56
1968	205	407	62	201	381	59
1969	211	432	66	205	407	62
1970	216	436	70	211	432	66
1971	214	438	70	216	436	70
1972	215	440	72	214	438	70
1973	220	444	73	215	440	72
1974	223	460	71	220	444	73
1975	224	480	65	223	460	71
1976	223	475	63	224	480	65
1977	224	483	57	223	475	63
1978	226	510	59	224	483	57
1979	225	501	60	226	510	59
1980	224	514	61	225	501	60
1981	226	520	61	224	514	61
1982	227	524	62	226	520	61
1983	230	538	65	227	524	62
1984	232	543	64	230	538	65
1985	235	560	63	232	543	64

Удерживая кнопку “Ctrl”, отметим щелчком левой кнопки мыши переменные $produk_t$ (объём производства), $zatrud_t$ (занятость) и $inwest_t$ (объём инвестиций). Затем щелчком правой кнопки мыши вызовем контекстное меню и выберем пункт **Time series plot** (рисунок 2), нажмём ОК. Характер динамики отмеченных процессов показывает график, представленный на рисунке 3.

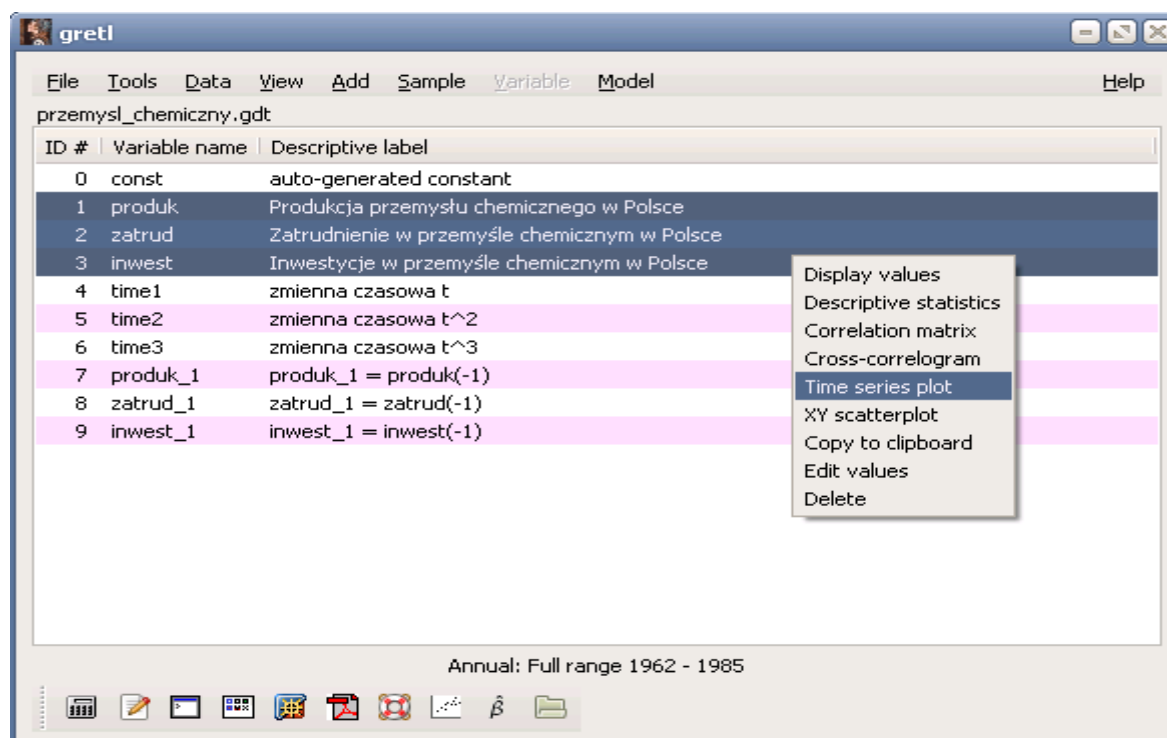


Рисунок 2 – Построение графика выбранных процессов

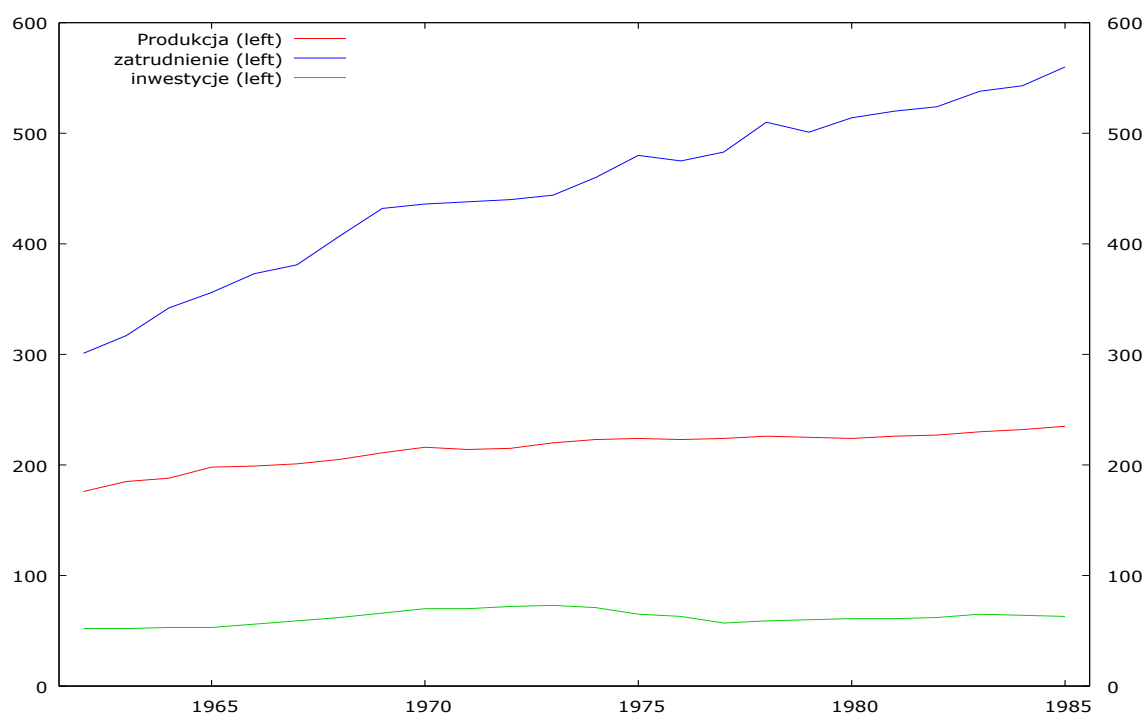


Рисунок 3 – Динамика показателей $produk_t$ (объём производства), $zatrud_t$ (занятость) и $inwest_t$ (объём инвестиций).

Как видно из графика за анализируемый период только изменение объёма производства показывает явно выраженную положительную динамику.

Выберем пункт меню **Model\Other Linear Models\Two-Stage Least Squares**..., что активирует окно спецификации одиночного уравнения. В трёх сегментах окна спецификации при помощи кнопок Choose и Add необходимо определить для отдельного уравнения структурной формы модели (1):

Dependent variable – (эндогенную) зависимую переменную (y) в левой части рассматриваемого уравнения;

Independent variables – все переменные в правой части рассматриваемого уравнения.

Instruments – все предопределённые (экзогенные и лаговые эндогенные) переменные всей системы;

Как показано на рисунке 4 для первого уравнения укажем:

Dependent variable – $y_1(\text{produk}_t)$;

Independent variables – $y_2(\text{zatrud}_t)$, $x_3(\text{inwest}_{t-1})$, $x_4(\text{produk}_{t-1})$;

Instruments – $x_2(\text{inwest}_t)$, $x_3(\text{inwest}_{t-1})$, $x_4(\text{produk}_{t-1})$ и const (поскольку x_1 присутствует в системе и равна 1).

Как показано на рисунке 5 для второго уравнения укажем:

Dependent variable – $y_2(\text{zatrud}_t)$;

Independent variables – $y_1(\text{produk}_t)$, $x_2(\text{inwest}_t)$, $x_3(\text{inwest}_{t-1})$ и const;

Instruments – $x_2(\text{inwest}_t)$, $x_3(\text{inwest}_{t-1})$, $x_4(\text{produk}_{t-1})$ и const.

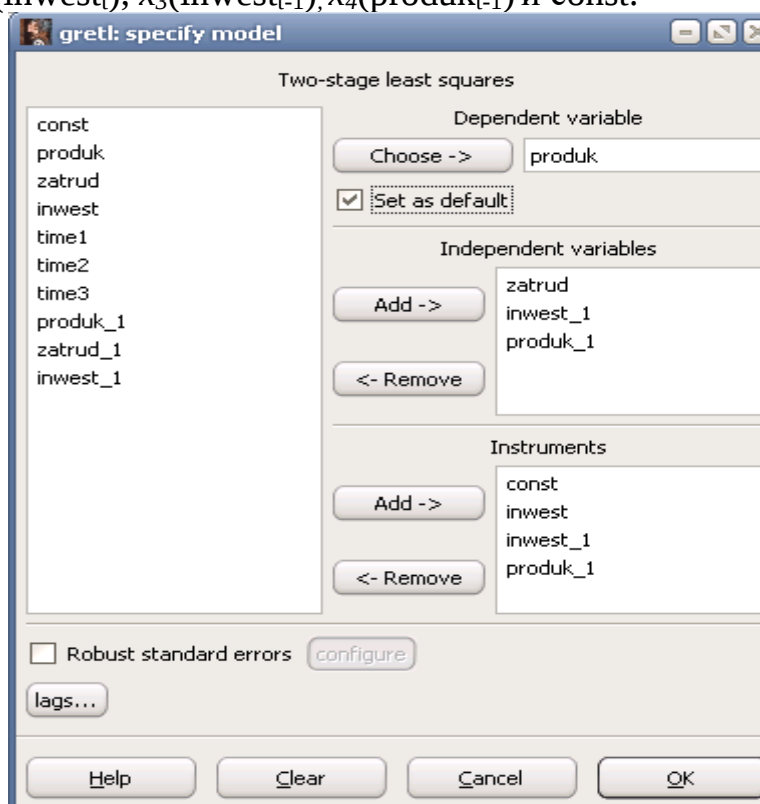


Рисунок 4 – Определение спецификации первого уравнения системы

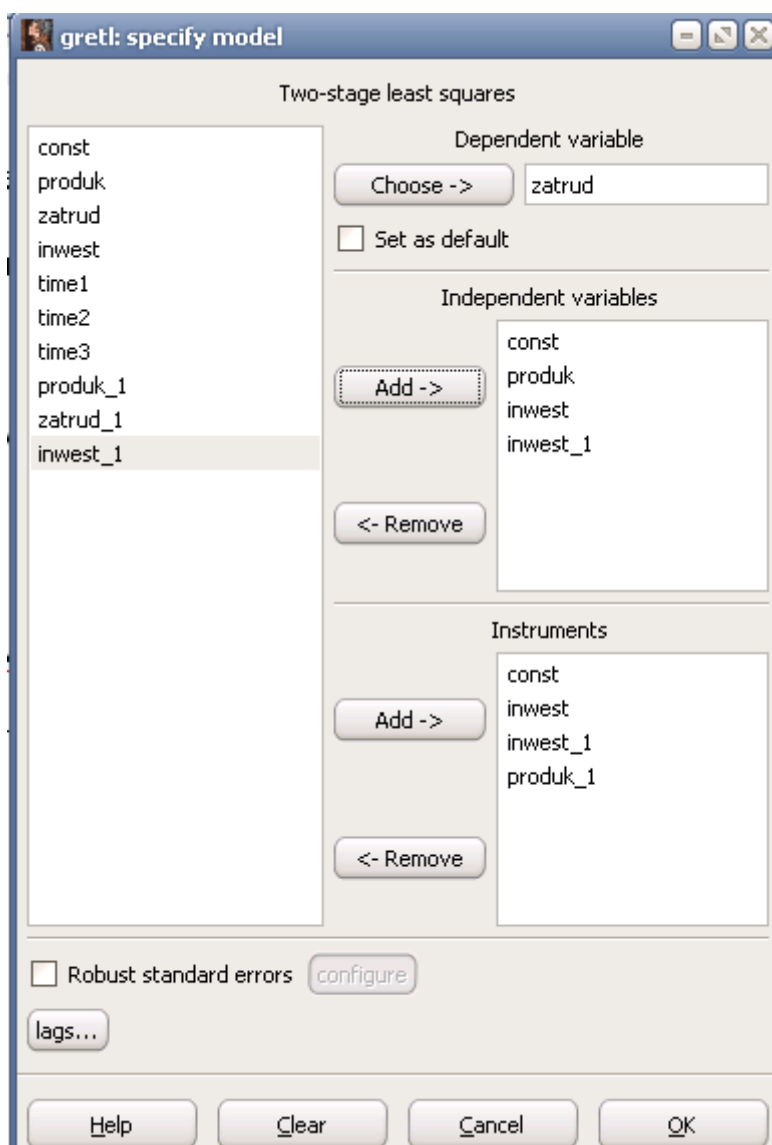


Рисунок 5 – Определение спецификации второго уравнения системы

Результаты оценивания рассматриваемой системы с применением 2МНК представлены на рисунках 6 и 7 для первого и второго уравнения соответственно. Структурная форма системы примет вид:

$$\begin{cases} y_1 = -0,097*y_2 - 0,135*x_3 + 1,26*x_4 + e_1; \\ y_2 = 5,137*y_1 - 522,135 + 0,194*x_2 - 2,32*x_3 + e_2; \end{cases}$$

или

$$\begin{cases} \text{produk}_t = -0,097* \text{zatrud}_t - 0,135* \text{inwest}_{t-1} + 1,26* \text{produk}_{t-1} + e_1; \\ \text{zatrud}_t = 5,137* \text{produk}_t - 522,135 + 0,194* \text{inwest}_t - 2,32* \text{inwest}_{t-1} + e_2; \end{cases}$$

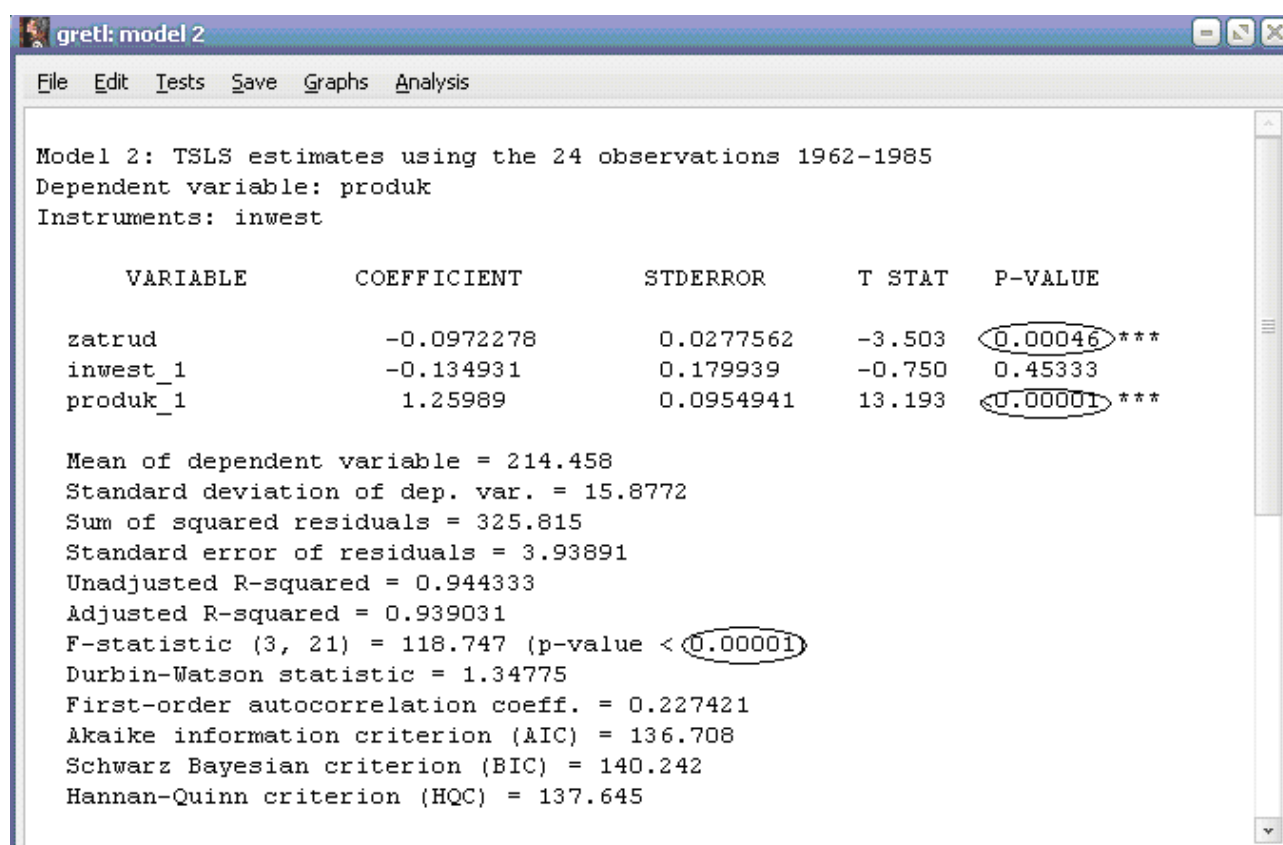


Рисунок 6 – Окно результатов моделирования с применением 2МНК для первого уравнения системы

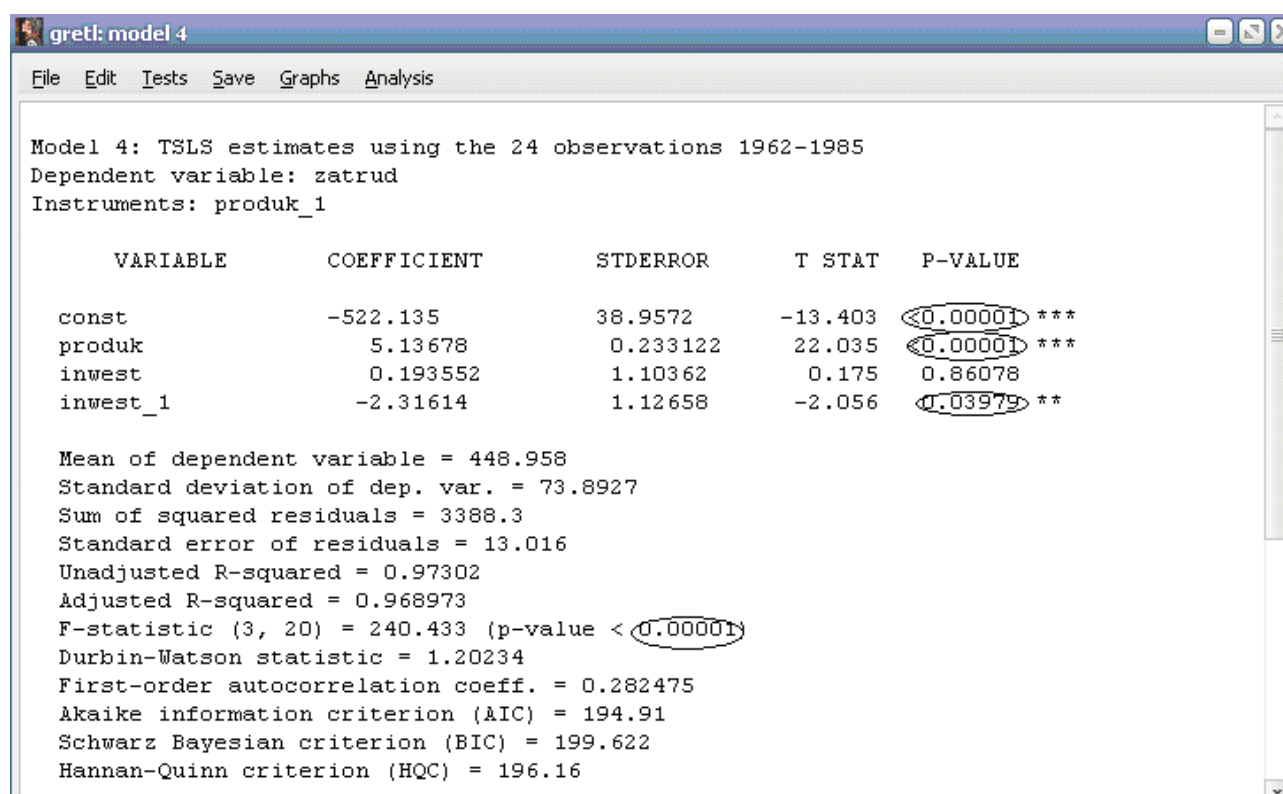


Рисунок 7 – Окно результатов моделирования с применением 2МНК для второго уравнения системы

Результаты оценивания первого уравнения свидетельствуют, что существенное влияние на объём производства ($produk_t$) оказывают занятость ($zatrud_t$) и значение объёма производства прошлого периода ($produk_{t-1}$), поскольку, согласно t-критерию Стьюдента, параметры при данных переменных являются статистически существенными при уровне значимости 1% (**), т.к. значения p-value 0,046% и 0,001% меньше 1%. В целом модель адекватна при уровне значимости 1%, поскольку для F-критерия Фишера p-value 0,001% меньше 1%.

Результаты оценивания второго уравнения показывают, что существенное влияние на занятость ($zatrud_t$) оказывают объём производства ($produk_t$) и объём инвестиций прошлого периода ($inwest_{t-1}$). Модель в целом также адекватна.

Исключим из первого уравнения переменную $x_3(inwest_{t-1})$, а из второго - $x_2(inwest_t)$, поскольку они не оказывают существенного влияния. Система примет вид (6):

$$\begin{cases} y_1 = b_{12} * y_2 + a_{14} * x_4 + e_1; \\ y_2 = b_{21} * y_1 + a_{21} + a_{23} * x_3 + e_2. \end{cases} \quad (6)$$

Преобразованная система также является сверхидентифицируемой: первое уравнение сверх идентифицируемое ($H=2 < (D+1=3)$), второе – точно идентифицируемое ($H=2 = (D+1=2)$).

Оценим параметры данной системы методом 2МНК вторым способом, обратившись к пункту **Simultaneous Equations** меню **Model**. Окно спецификации системы одновременных уравнений представлено на рисунке 8.

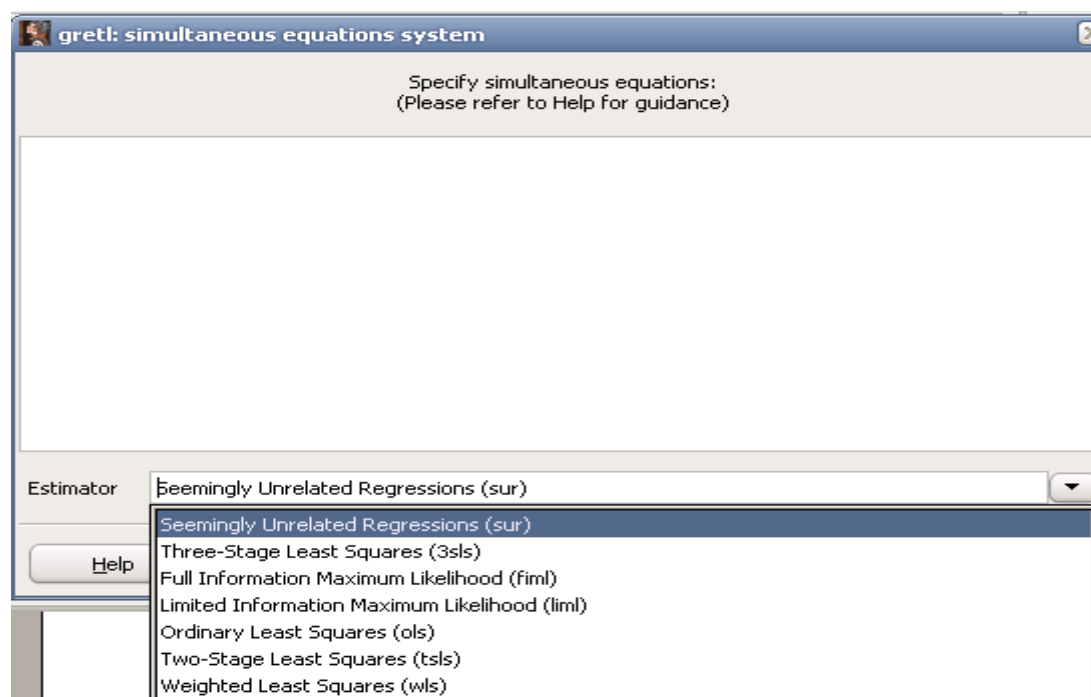


Рисунок 8 - Окно спецификации системы одновременных уравнений

Из открывающегося списка внизу окна выберем метод оценки параметров системы – Two-Stage Least Squares (tsls) и правой кнопкой мыши щёлкнем по окну для вызова контекстного меню (рисунок 9).

Выберем пункт контекстного меню Add equation для спецификации первого уравнения системы (6). Затем через пробел введём названия переменных в первом уравнении, щелчком мыши выбирая их из анализируемого набора данных (рисунок 1). Аналогичным образом второй строкой добавим второе уравнение системы. Отметим, что «0» используется для обозначения константы. Затем третьей строкой введём список эндогенных переменных системы (endog), вызвав контекстное меню и выбрав в нём пункт Add list of endogenous variables (рисунок 9).

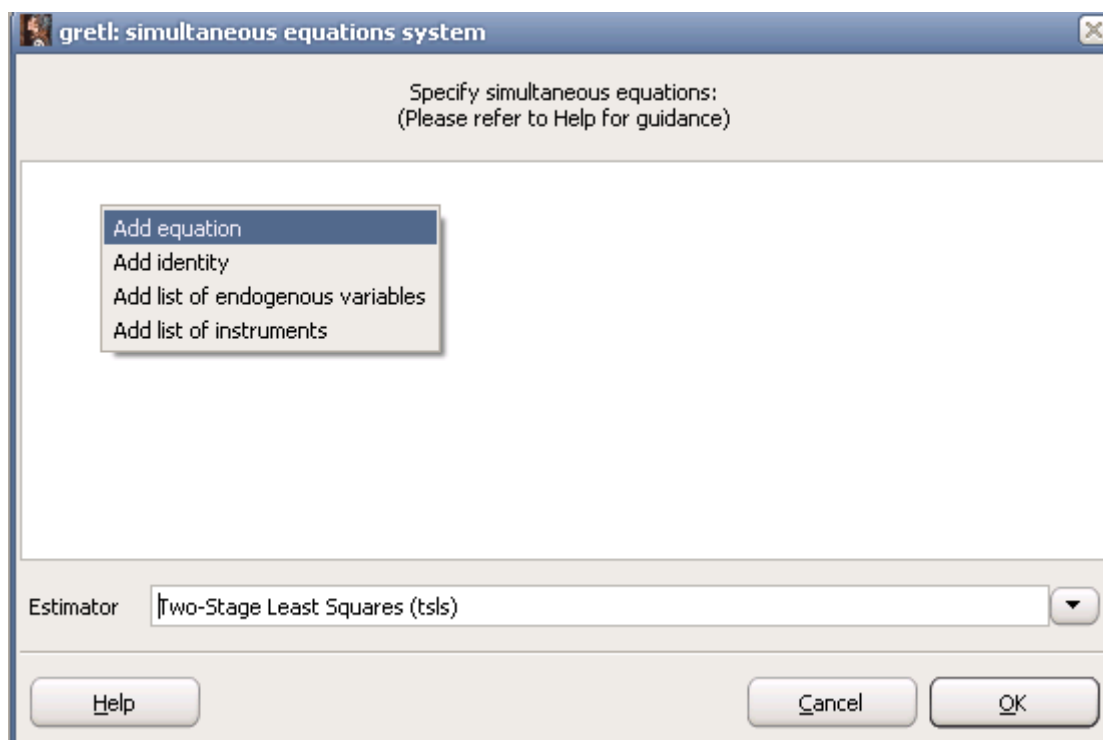


Рисунок 9 – Вызов контекстного меню

В результате для анализируемой системы (6) спецификация будет иметь вид, представленный на рисунке 10. Нажмём кнопку ОК.

Окно результатов моделирования (рисунок 11) показывает, что все параметры обоих уравнений скорректированной системы значимы по t-критерию Стьюдента при уровне значимости 1%, т.е. все переменные в правой части обоих уравнений оказывают существенное влияние на эндогенные переменные в левой части:

$$\begin{cases} \text{produk}_t = -0,088 * \text{zatrud}_t + 1,2 * \text{produk}_{t-1} + e_1; \\ \text{zatrud}_t = 5,134 * \text{produk}_t - 520,08 - 2,14 * \text{inwest}_{t-1} + e_2; \end{cases}$$

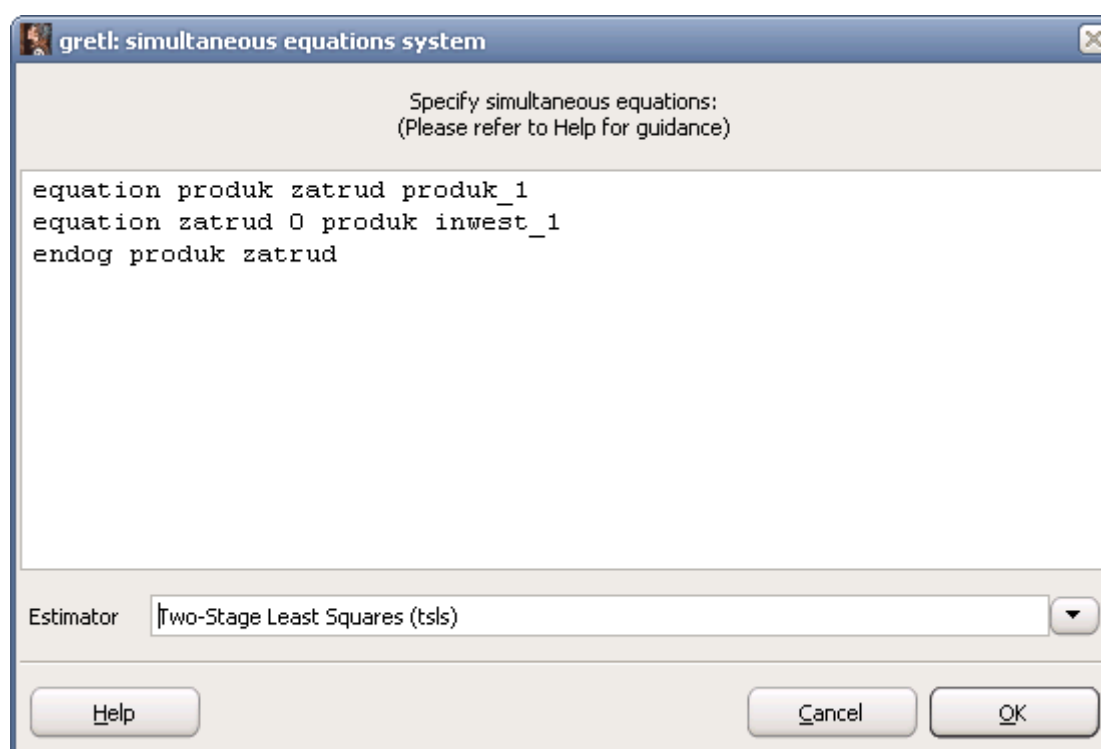


Рисунок 10 – Ввод данных анализируемой модели

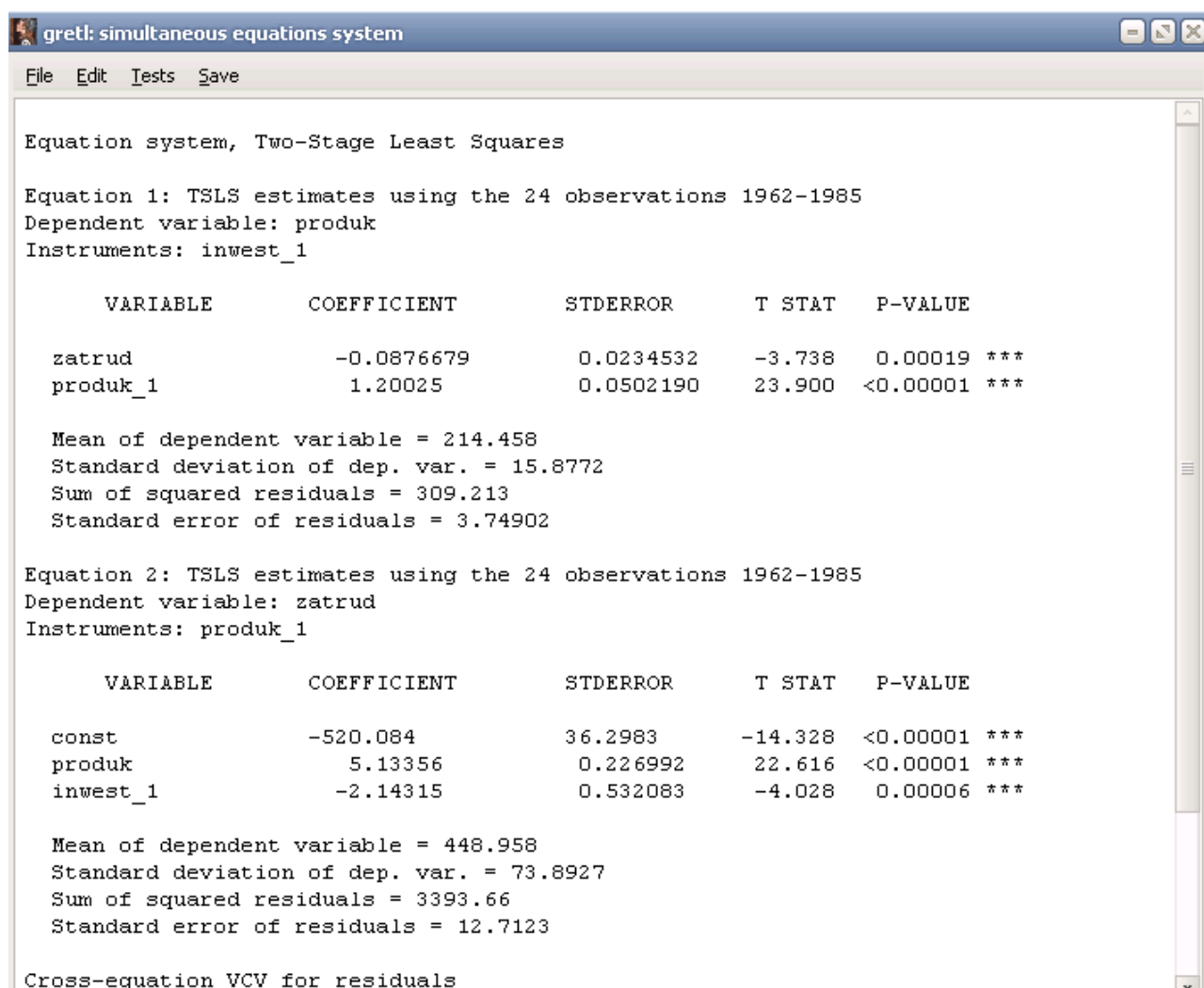


Рисунок 11 – Окно результатов моделирования методом 2 МНК (второй способ)

Выберем пункт меню **Tools\command log** (рисунок 12). В окне command log выводится перечень функций, зарегистрированных в меню команд языка Gretl в процессе оценивания модели. Т.о. существует возможность сохранения всей последовательности действий (показанных выше) по оцениванию рассмотренных систем в виде файла- скрипта формата *.inp. При новом запуске программы Gretl появится возможность открытия и выполнения сохранённого ранее скрипта для получения результатов моделирования, описанных выше.

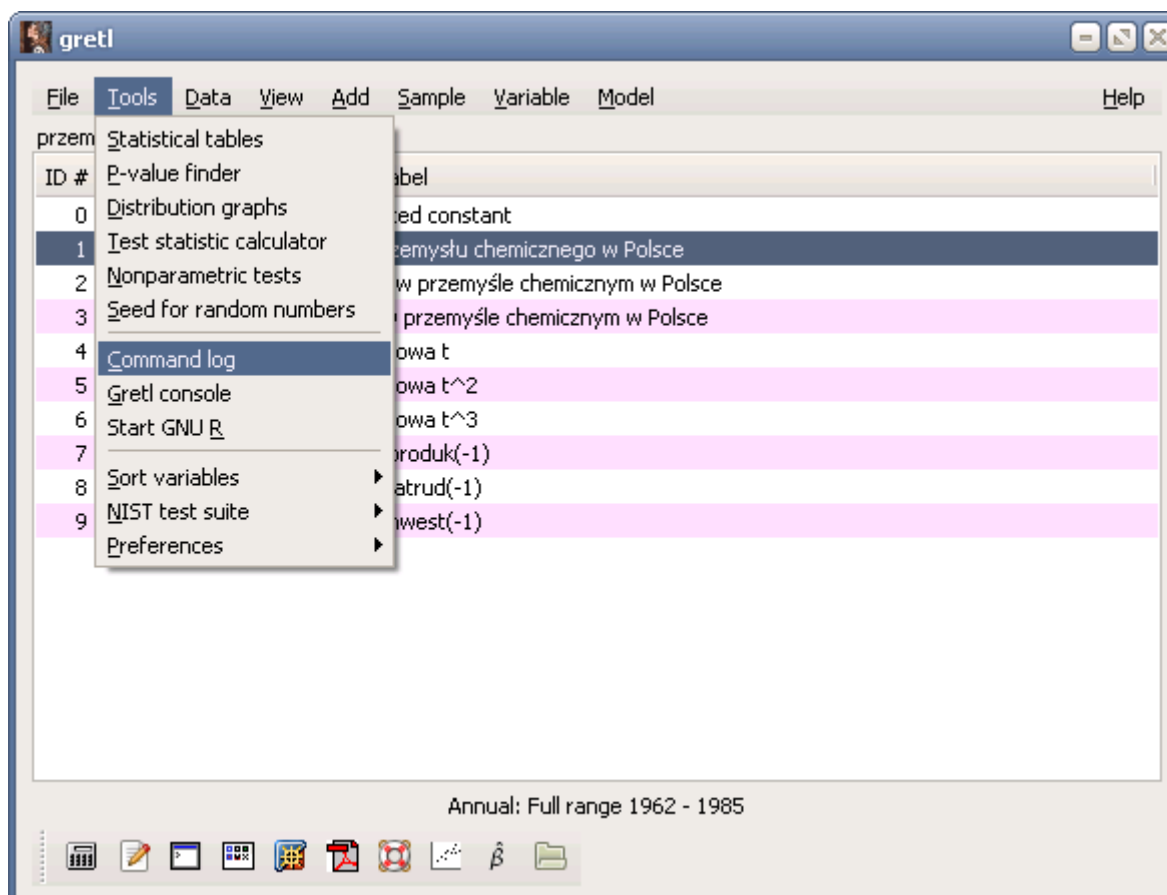


Рисунок 12 – Вывод перечня функций, зарегистрированных в меню команд языка Gretl

В открывшемся окне command log (рисунок 13) щёлкнем левой кнопкой мыши по иконке меню «сохранить как» (Save as) с изображением дискеты и введём в открывшемся окне название файла-скрипта lab6.inp, по умолчанию сохранив его в папку gretl.

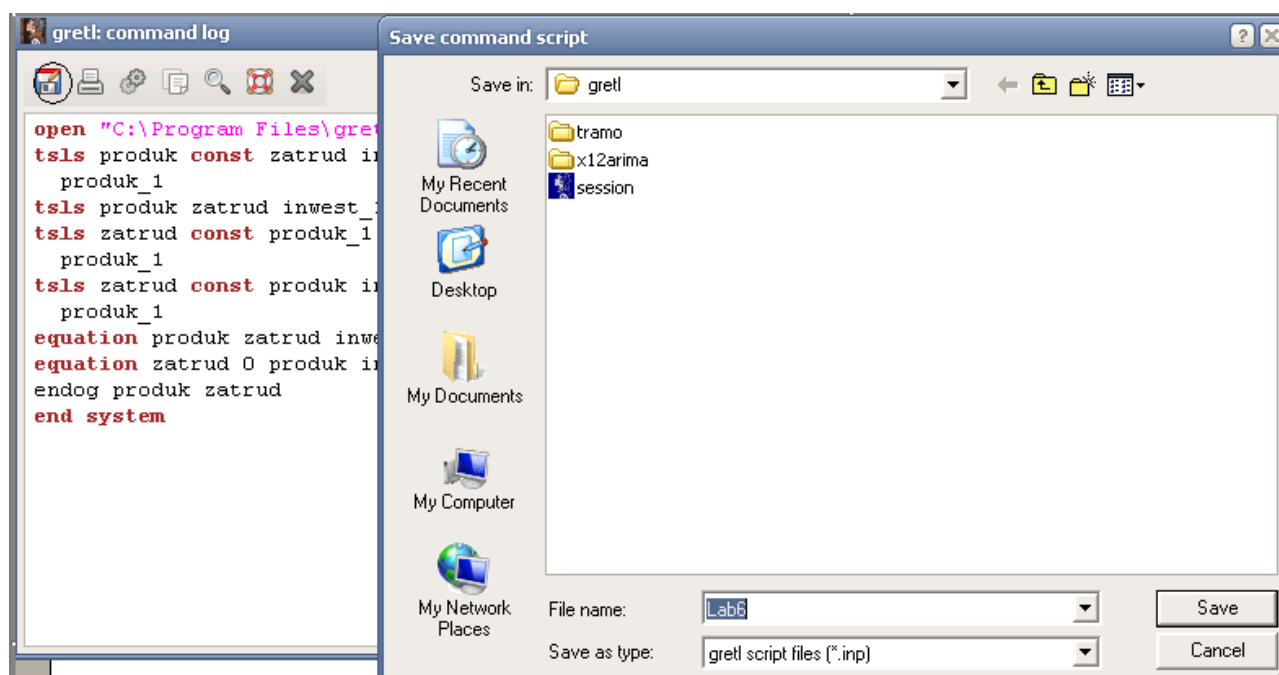


Рисунок 13 – Создание файла– скрипта Lab6.inp

Затем выйдем из программы Gretl и запустим её заново. Обратимся к пункту меню **File\Script files\User file...** (рисунок 14) и выберем сохранённый Lab6.inp.

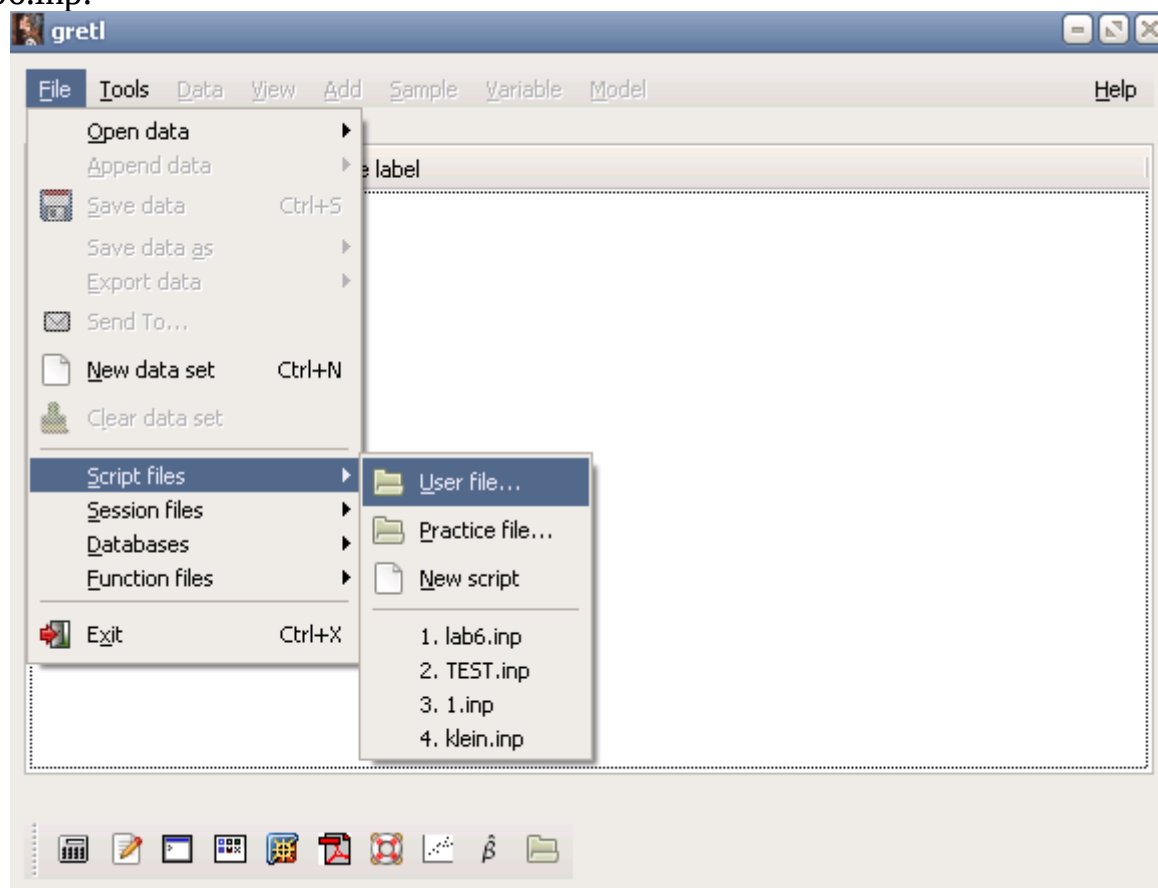
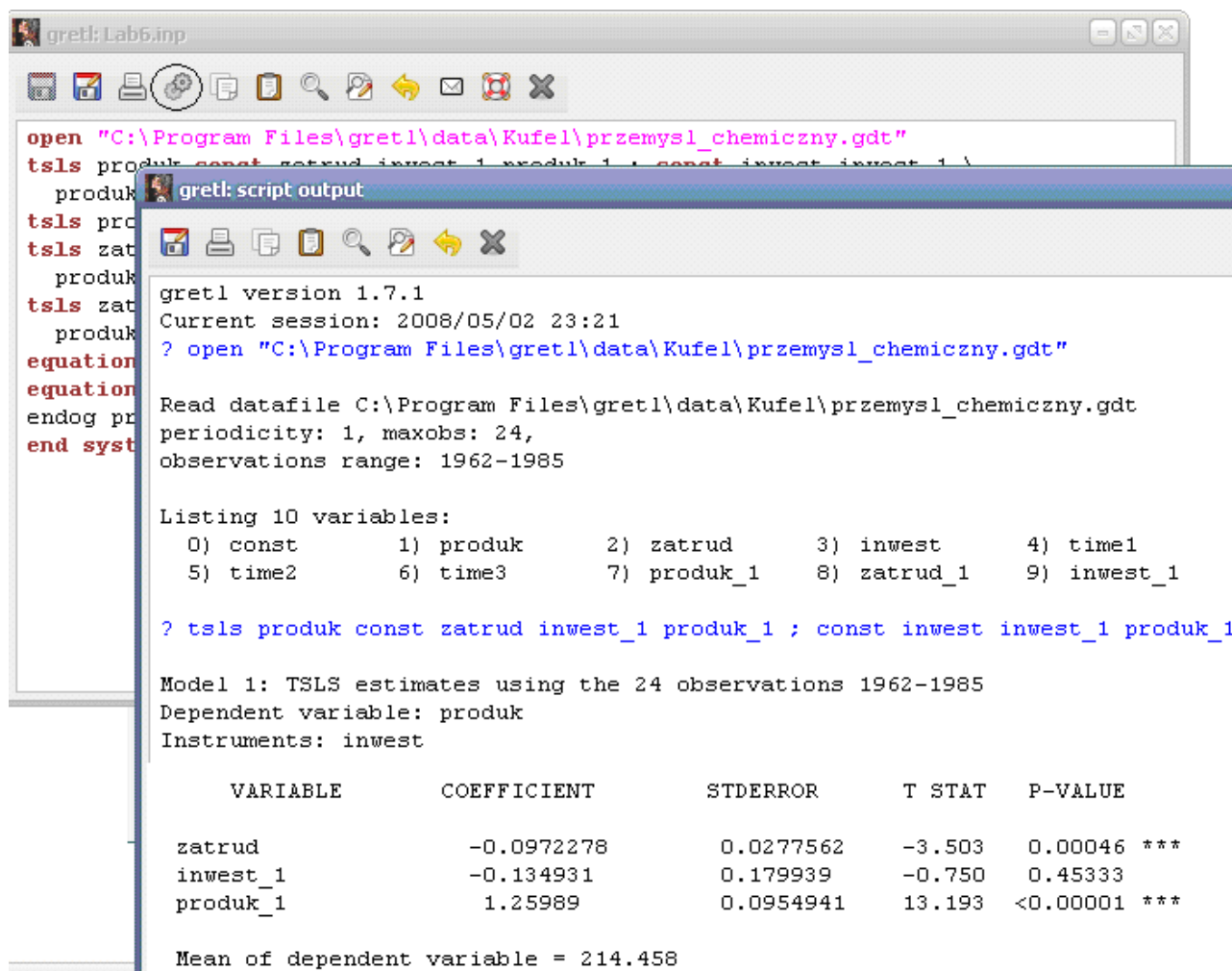


Рисунок 14 – Открытие файла-скрипта Lab6.inp

В открывшемся окне выберем иконку меню «выполнить» (run) скрипт (рисунок 15) и в окне script output получим результаты проведённого ранее моделирования.



```

open "C:\Program Files\gretl\data\Kufel\przemysl_chemiczny.gdt"
tsls produk const zatrud inwest_1 produk_1 ; const inwest inwest_1 produk_1
produk
gretl: script output
tsls pro
tsls zat
produk
tsls zat
produk
equation
equation
endog pr
end syst

```

gretl version 1.7.1
Current session: 2008/05/02 23:21
? open "C:\Program Files\gretl\data\Kufel\przemysl_chemiczny.gdt"

Read datafile C:\Program Files\gretl\data\Kufel\przemysl_chemiczny.gdt
periodicity: 1, maxobs: 24,
observations range: 1962-1985

Listing 10 variables:

0) const	1) produk	2) zatrud	3) inwest	4) time1
5) time2	6) time3	7) produk_1	8) zatrud_1	9) inwest_1

? tsls produk const zatrud inwest_1 produk_1 ; const inwest inwest_1 produk_1

Model 1: TSLS estimates using the 24 observations 1962-1985
Dependent variable: produk
Instruments: inwest

VARIABLE	COEFFICIENT	STDERROR	T STAT	P-VALUE
zatrud	-0.0972278	0.0277562	-3.503	0.00046 ***
inwest_1	-0.134931	0.179939	-0.750	0.45333
produk_1	1.25989	0.0954941	13.193	<0.00001 ***

Mean of dependent variable = 214.458

Рисунок 15 – Выполнение файла-скрипта Lab6.inp

4. ПОРЯДОК ВЫПОЛНЕНИЯ ЛАБОРАТОРНОЙ РАБОТЫ

1. Открыть или создать набор данных №1 и №2 согласно варианту (таблица 2), проанализировать переменные в их структуре. Задать состав эндогенных и предопределённых переменных и составить структурную форму системы одновременных уравнений по каждому набору данных.

2. Оценить параметры структурной формы каждой системы на основе 2МНК двумя способами.

3. Оценить адекватность регрессионных уравнений в целом и значимость их параметров для обеих систем.

4. Сохранить всю последовательность действий по оцениванию систем в виде файла-скрипта, запустить скрипт.

Таблица 2 – Варианты заданий

Вариант	Набор данных №1	Набор данных №2
1	Рисунок 16, таблица 3	Рисунок 17, таблица 4
2	Рисунок 16, таблица 3	Рисунок 18, таблица 5
3	Рисунок 16, таблица 3	Рисунок 19, таблица 6
4	Рисунок 16, таблица 3	Рисунок 20, таблица 7
5	Рисунок 16, таблица 3	Рисунок 1, таблица 1
6	Рисунок 17, таблица 4	Рисунок 18, таблица 5
7	Рисунок 17, таблица 4	Рисунок 19, таблица 6
8	Рисунок 17, таблица 4	Рисунок 20, таблица 7
9	Рисунок 17, таблица 4	Рисунок 1, таблица 1
10	Рисунок 18, таблица 5	Рисунок 19, таблица 6
11	Рисунок 18, таблица 5	Рисунок 20, таблица 7
12	Рисунок 18, таблица 5	Рисунок 1, таблица 1
13	Рисунок 19, таблица 6	Рисунок 20, таблица 7
14	Рисунок 19, таблица 6	Рисунок 1, таблица 1
15	Рисунок 20, таблица 7	Рисунок 1, таблица 1

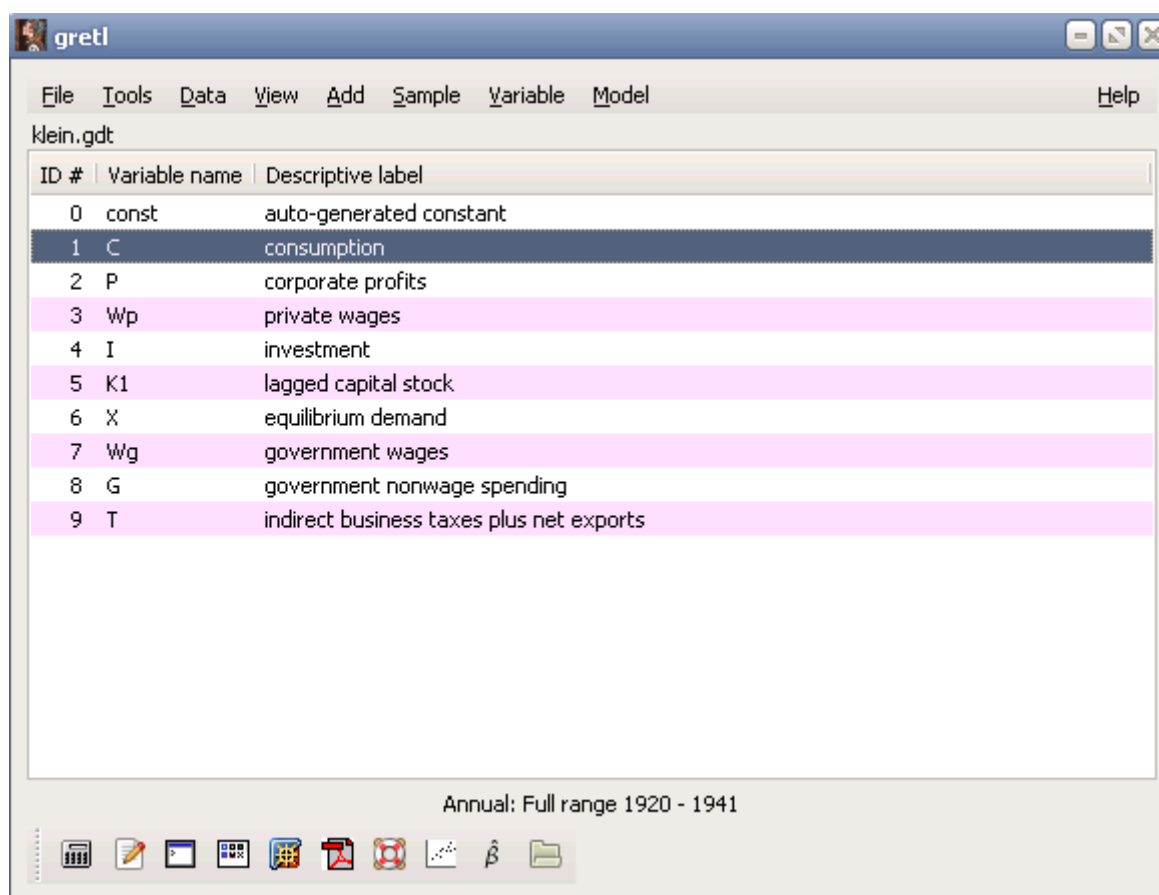


Рисунок 16 – Исходные данные (File\Open Data\ Sample File...GRETl:Klein.gdt)

Таблица 3 – Исходные данные GRETL:Klein.gdt

obs	C	P	Wp	I	K1	X	Wg	G	T
1920	39.8	12.7	28.8	2.7	180.1	44.9	2.2	2.4	3.4
1921	41.9	12.4	25.5	-0.2	182.8	45.6	2.7	3.9	7.7
1922	45	16.9	29.3	1.9	182.6	50.1	2.9	3.2	3.9
1923	49.2	18.4	34.1	5.2	184.5	57.2	2.9	2.8	4.7
1924	50.6	19.4	33.9	3	189.7	57.1	3.1	3.5	3.8
1925	52.6	20.1	35.4	5.1	192.7	61	3.2	3.3	5.5
1926	55.1	19.6	37.4	5.6	197.8	64	3.3	3.3	7
1927	56.2	19.8	37.9	4.2	203.4	64.4	3.6	4	6.7
1928	57.3	21.1	39.2	3	207.6	64.5	3.7	4.2	4.2
1929	57.8	21.7	41.3	5.1	210.6	67	4	4.1	4
1930	55	15.6	37.9	1	215.7	61.2	4.2	5.2	7.7
1931	50.9	11.4	34.5	-3.4	216.7	53.4	4.8	5.9	7.5
1932	45.6	7	29	-6.2	213.3	44.3	5.3	4.9	8.3
1933	46.5	11.2	28.5	-5.1	207.1	45.1	5.6	3.7	5.4
1934	48.7	12.3	30.6	-3	202	49.7	6	4	6.8
1935	51.3	14	33.2	-1.3	199	54.4	6.1	4.4	7.2
1936	57.7	17.6	36.8	2.1	197.7	62.7	7.4	2.9	8.3
1937	58.7	17.3	41	2	199.8	65	6.7	4.3	6.7
1938	57.5	15.3	38.2	-1.9	201.8	60.9	7.7	5.3	7.4
1939	61.6	19	41.6	1.3	199.9	69.5	7.8	6.6	8.9
1940	65	21.1	45	3.3	201.2	75.7	8	7.4	9.6
1941	69.7	23.5	53.3	4.9	204.5	88.4	8.5	13.8	11.6

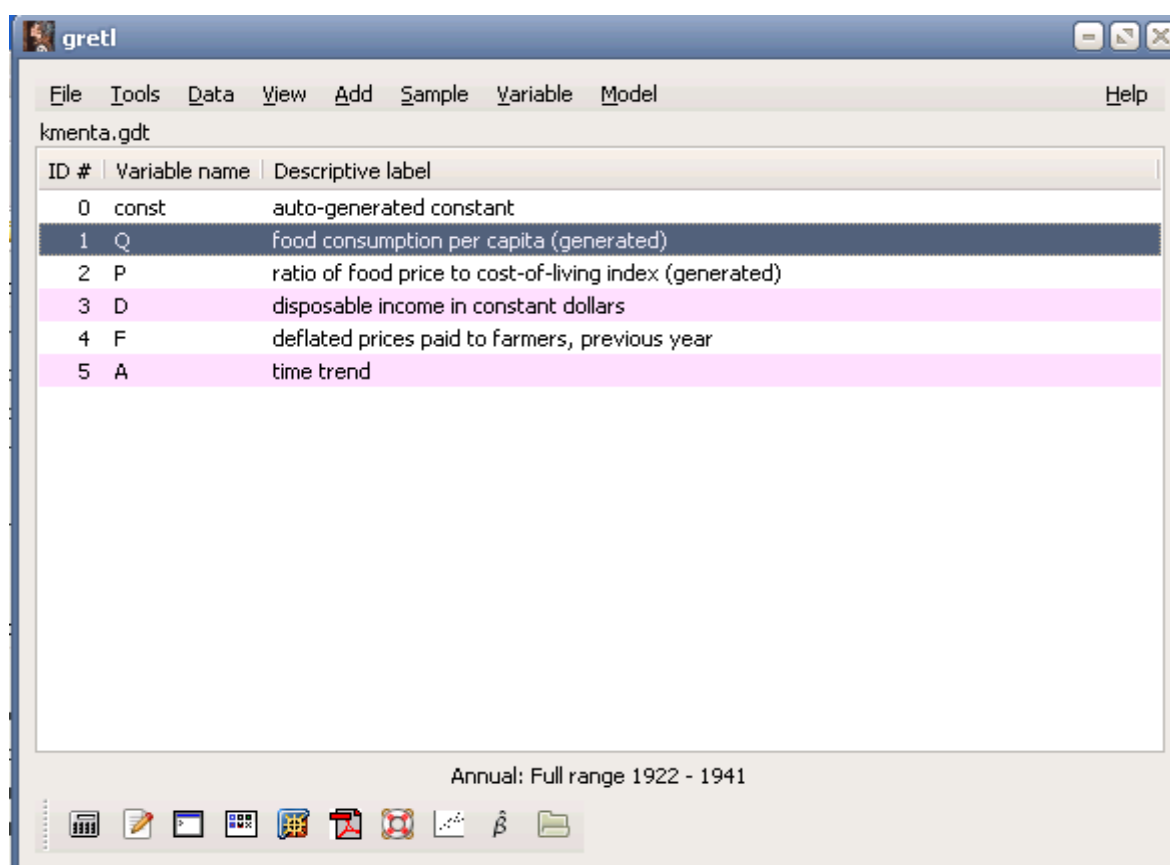


Рисунок 17 – Исходные данные (File\Open Data\ Sample File...GRETL:Kmenta.gdt)

Таблица 4 – Исходные данные GRETL:Kmenta.gdt

obs	Q	P	D	F	A
1922	98.485	100.323	87.4	98	1
1923	99.187	104.264	97.6	99.1	2
1924	102.163	103.435	96.7	99.1	3
1925	101.504	104.506	98.2	98.1	4
1926	104.24	98.001	99.8	110.8	5
1927	103.243	99.456	100.5	108.2	6
1928	103.993	101.066	103.2	105.6	7
1929	99.9	104.763	107.8	109.8	8
1930	100.35	96.446	96.6	108.7	9
1931	102.82	91.228	88.9	100.6	10
1932	95.435	93.085	75.1	81	11
1933	92.424	98.801	76.9	68.6	12
1934	94.535	102.908	84.6	70.9	13
1935	98.757	98.756	90.6	81.4	14
1936	105.797	95.119	103.1	102.3	15
1937	100.225	98.451	105.1	105	16
1938	103.522	86.498	96.4	110.5	17
1939	99.929	104.016	104.4	92.5	18
1940	105.223	105.769	110.7	89.3	19
1941	106.232	113.49	127.1	93	20

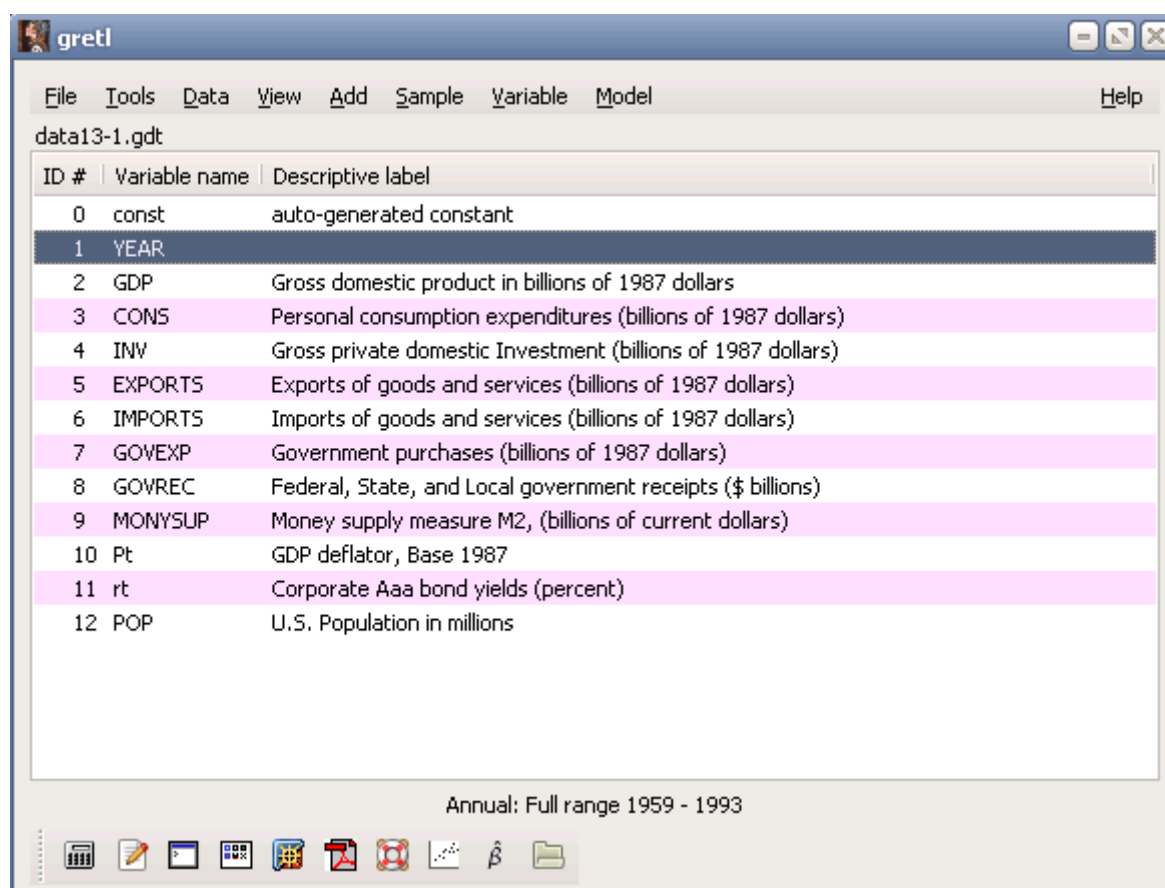
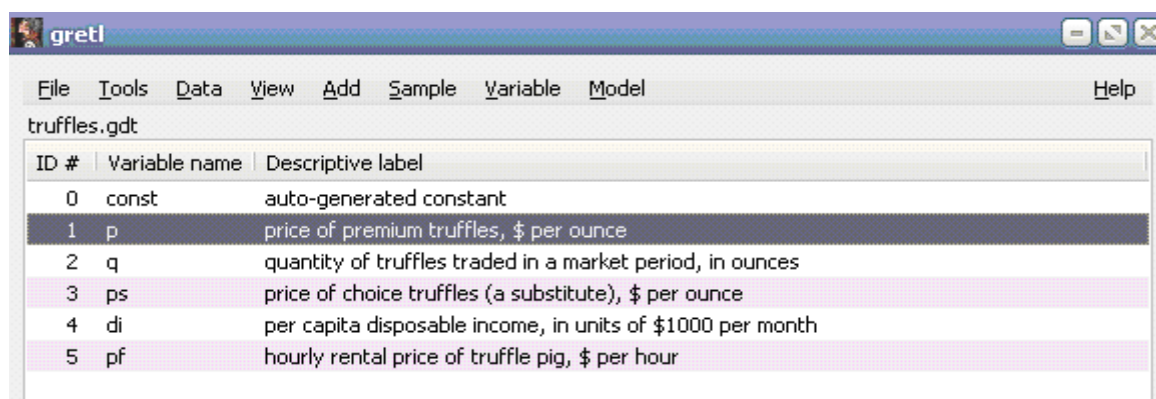


Рисунок 18 – Исходные данные (File\Open Data\ Sample File...RAMANATHAN:Data13-1.gdt)

Таблица 5 – Исходные данные RAMANATHAN:Data13-1.gdt

Obs	Gdp	Cons	Inv	Exports	Imports	Govexp	Govrec	Monysup	Pop
1959	1928.8	1178.9	296.4	73.8	95.6	475.3	128.8	297.8	177.83
1960	1970.8	1210.8	290.8	88.4	96.1	476.9	138.8	312.3	180.671
1961	2023.8	1238.4	289.4	89.9	95.3	501.5	144.1	335.5	183.691
1962	2128.1	1293.3	321.2	95	105.5	524.2	155.8	362.7	186.538
1963	2215.6	1341.9	343.3	101.8	107.7	536.3	167.5	393.2	189.242
1964	2340.6	1417.2	371.8	115.4	112.9	549.1	172.9	424.8	191.889
1965	2470.5	1497	413	118.1	124.5	566.9	187	459.3	194.303
1966	2616.2	1573.8	438	125.7	143.7	622.4	210.7	480	196.56
1967	2685.2	1622.4	418.6	130	153.7	667.9	226.4	524.3	198.712
1968	2796.9	1707.5	440.1	140.2	177.7	686.8	260.9	566.3	200.706
1969	2873	1771.2	461.3	147.8	189.2	682	294	589.5	202.677
1970	2873.9	1813.5	429.7	161.3	196.4	665.8	299.8	628	205.052
1971	2955.9	1873.7	475.7	161.9	207.8	652.4	318.9	712.6	207.661
1972	3107.1	1978.4	532.2	173.7	230.2	653	364.2	805.1	209.896
1973	3268.6	2066.7	591.7	210.3	244.4	644.2	408.5	860.9	211.909
1974	3248.1	2053.8	543	234.4	238.4	655.4	450.7	908.4	213.854
1975	3221.7	2097.5	437.6	232.9	209.8	663.5	465.8	1023.1	215.973
1976	3380.8	2207.3	520.6	243.4	249.7	659.2	532.6	1163.5	218.035
1977	3533.3	2296.6	600.4	246.9	274.7	664.1	598.4	1286.4	220.239
1978	3703.5	2391.8	664.6	270.2	300.1	677	673.2	1388.5	222.585
1979	3796.8	2448.4	669.7	293.5	304.1	689.3	754.7	1496.4	225.055
1980	3776.3	2447.1	594.4	320.5	289.9	704.2	825.7	1629.2	227.726
1981	3843.1	2476.9	631.1	326.1	304.1	713.2	941.9	1792.6	229.966
1982	3760.3	2503.7	540.5	296.7	304.1	723.6	960.5	1952.7	232.188
1983	3906.6	2619.4	599.5	285.9	342.1	743.8	1016.4	2186.5	234.307
1984	4148.5	2746.1	757.5	305.7	427.7	766.9	1123.6	2376	236.348
1985	4279.8	2865.8	745.9	309.2	454.6	813.4	1217	2572.4	238.466
1986	4404.5	2969.1	735.1	329.6	484.7	855.4	1290.8	2816.1	240.651
1987	4539.9	3052.2	749.3	364	507.1	881.5	1405.2	2917.2	242.804
1988	4718.6	3162.4	773.4	421.6	525.7	886.8	1492.4	3078.2	245.021
1989	4838	3223.3	784	471.8	545.4	904.4	1622.6	3233.3	247.342
1990	4897.3	3272.6	746.8	510.5	565.1	932.6	1709.1	3345.5	249.9
1991	4861.4	3258.6	675.7	543.4	562.5	946.3	1755.2	3445.8	252.671
1992	4986.3	3341.8	732.9	578	611.6	945.2	1849.4	3494.8	255.462
1993	5132.7	3452.5	820.9	596.4	675.7	938.6	1969.1	3551.7	258.233



ID #	Variable name	Descriptive label
0	const	auto-generated constant
1	p	price of premium truffles, \$ per ounce
2	q	quantity of truffles traded in a market period, in ounces
3	ps	price of choice truffles (a substitute), \$ per ounce
4	di	per capita disposable income, in units of \$1000 per month
5	pf	hourly rental price of truffle pig, \$ per hour

Рисунок 19 – Исходные данные ((File\Open Data\ Sample File...PoE: truffles.gdt)

Таблица 6 – Исходные данные PoE: truffles.gdt

p	q	ps	di	Pf
29.64	19.89	19.97	2.103	10.52
40.23	13.04	18.04	2.043	19.67
34.71	19.61	22.36	1.87	13.74
41.43	17.13	20.87	1.525	17.95
53.37	22.55	19.79	2.709	13.71
38.52	6.37	15.98	2.489	24.95
54.33	15.02	17.94	2.294	24.17
40.56	10.22	17.09	2.196	23.61
67.35	23.64	22.72	3.885	19.52
49.65	16.12	15.74	3.169	20.03
58.17	24.55	24.64	2.623	15.38
66.87	18.92	23.7	3.007	22.98
49.95	11.94	15.93	3.367	25.76
64.95	18.93	23.34	3.29	25.17
52.68	12.6	15.21	3.746	25.82
61.2	20.49	26.04	3.518	19.31
80.55	22.94	22.95	4.381	26.02
89.94	21.08	27.1	4.121	29.65
70.77	16.68	23.65	3.82	27.45
57.33	17.61	20.06	4.398	18
46.23	16.62	26.38	3.764	18.87
77.43	20.99	24.28	4.524	24.58
83.01	24.53	26.64	4.815	25.25
70.71	19.67	22.65	3.67	24.24
66.75	23.29	19.68	4.392	22.63
76.8	16.64	23.82	4.603	27.35
83.7	20.81	28.98	4.632	27.8
81	14.95	18.52	4.894	30.34
88.44	26.27	28.16	5.125	24.12
105.45	20.65	28.43	4.836	34.01

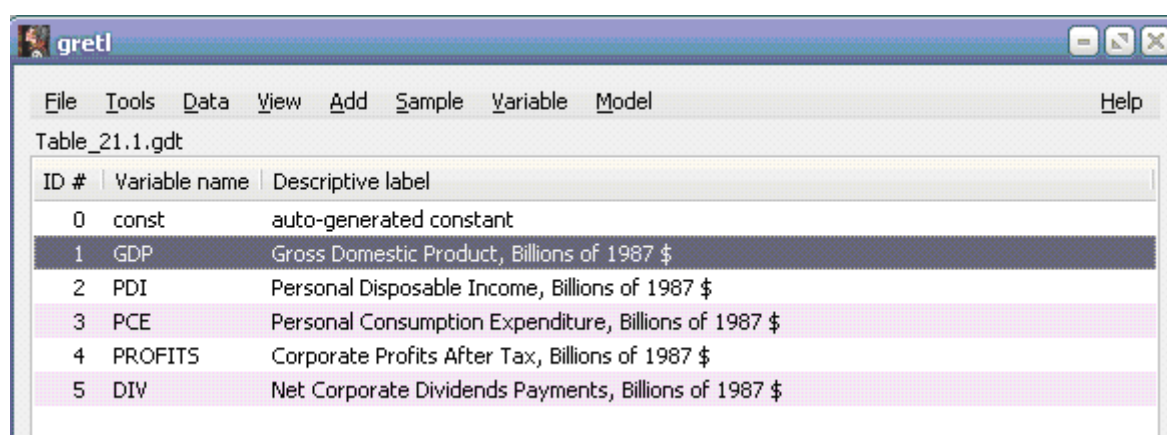


Рисунок 20 – Исходные данные (File\Open Data\ Sample File...Gujarati: table_21.1.gdt)

Таблица 7 – Исходные данные Gujarati: table_21.1.gdt

obs	GDP	PDI	PCE	PROFITS	DIV
1970Q1	2872.8	1990.6	1800.5	44.7	24.5
1970Q2	2860.3	2020.1	1807.5	44.4	23.9
1970Q3	2896.6	2045.3	1824.7	44.9	23.3
1970Q4	2873.7	2045.2	1821.2	42.1	23.1
1971Q1	2942.9	2073.9	1849.9	48.8	23.8
1971Q2	2947.4	2098	1863.5	50.7	23.7
1971Q3	2966	2106.6	1876.9	54.2	23.8
1971Q4	2980.8	2121.1	1904.6	55.7	23.7
1972Q1	3037.3	2129.7	1929.3	59.4	25
1972Q2	3089.7	2149.1	1963.3	60.1	25.5
1972Q3	3125.8	2193.9	1989.1	62.8	26.1
1972Q4	3175.5	2272	2032.1	68.3	26.5
1973Q1	3253.3	2300.7	2063.9	79.1	27
1973Q2	3267.6	2315.2	2062	81.2	27.8
1973Q3	3264.3	2337.9	2073.7	81.3	28.3
1973Q4	3289.1	2382.7	2067.4	85	29.4
1974Q1	3259.4	2334.7	2050.8	89	29.8
1974Q2	3267.6	2304.5	2059	91.2	30.4
1974Q3	3239.1	2315	2065.5	97.1	30.9
1974Q4	3226.4	2313.7	2039.9	86.8	30.5
1975Q1	3154	2282.5	2051.8	75.8	30
1975Q2	3190.4	2390.3	2086.9	81	29.7
1975Q3	3249.9	2354.4	2114.4	97.8	30.1
1975Q4	3292.5	2389.4	2137	103.4	30.6
1976Q1	3356.7	2424.5	2179.3	108.4	32.6
1976Q2	3369.2	2434.9	2194.7	109.2	35
1976Q3	3381	2444.7	2213	110	36.6
1976Q4	3416.3	2459.5	2242	110.3	38.3
1977Q1	3466.4	2463	2271.3	121.5	39.2
1977Q2	3525	2490.3	2280.8	129.7	40
1977Q3	3574.4	2541	2302.6	135.1	41.4
1977Q4	3567.2	2556.2	2331.6	134.8	42.4
1978Q1	3591.8	2587.3	2347.1	137.5	43.5
1978Q2	3707	2631.9	2394	154	44.5
1978Q3	3735.6	2653.2	2404.5	158	46.6
1978Q4	3779.6	2680.9	2421.6	167.8	48.9
1979Q1	3780.8	2699.2	2437.9	168.2	50.5
1979Q2	3784.3	2697.6	2435.4	174.1	51.8
1979Q3	3807.5	2715.3	2454.7	178.1	52.7
1979Q4	3814.6	2728.1	2465.4	173.4	54.5
1980Q1	3830.8	2742.9	2464.6	174.3	57.6
1980Q2	3732.6	2692	2414.2	144.5	58.7
1980Q3	3733.5	2722.5	2440.3	151	59.3
1980Q4	3808.5	2777	2469.2	154.6	60.5

5. СОДЕРЖАНИЕ ОТЧЕТА О ВЫПОЛНЕНИИ ЛАБОРАТОРНОЙ РАБОТЫ

- 1) Название и цель работы.
- 2) Постановка задачи.
- 3) Этапы выполнения задачи в Gretl.
- 4) Выводы.