
The purpose of this project is to familiarize you with macroeconomic data and with the measurement and basic properties of economic growth and cyclical fluctuations. Each group will work with data for a different country.

Partner assignments

<i>Partners for this project</i>		
Axcelle Bell	Mrijan Rimal	United Kingdom
Niall Murphy	Zakir Hakim	Mexico
Nicholas Wilson	Samantha Bruce	Chile
Nisma Elias	Ximeng Zhang	Finland
Mark Hintz	Gabriel Korzeniewicz	Japan
Antonio Friedman	Suraj Pant	Ireland

Note: It is possible that there may be some drops and adds during the first week. I'll do my best to stay on top of the situation day by day and reassign partners as necessary. If we end up with an odd number of students, one student will end up working alone.

In this project you will explore the behavior of real GDP for one country. It will introduce you to Stata, the statistical software package we will use periodically during the semester, and to some basic techniques of data analysis. This handout gives you some background on the software and statistical techniques you will be using. After this introduction, it is organized as a series of three exercises. Each of these exercises includes extensive explanatory text and descriptions of the statistical tasks that you are to perform using the data for your countries. To help you navigate the document, the sections that describe the assignment tasks are indented, printed in a sans serif typeface, and highlighted with a triple line along the left margin.

As with most other weekly projects in this course, you will work on this project as a team with one other student. Your team is to submit one report giving detailed responses to the questions and tasks described below. The report should be in the form of a Word (or, preferably, pdf) document, submitted electronically as an email attachment to the instructor prior to the 9:00am Wednesday deadline. (It would be helpful if you would append .doc to your Word file name or .pdf if it is a pdf and make sure that one or both of your names are in the title. It is confusing to get a bunch of files that are all titled "Project 1" or "Macro Assignment.") *All relevant results from your Stata output should be copied and pasted into the report file.* When you insert Stata text output into a Word file, you should use a fixed-width font such as **Courier 9pt** bold looks nice and fits in the margins, at least on my computer) so that the columns of your tables will align vertically. If you have problems pasting in graphics, you may submit an appendix in printed form, as long as each part is clearly identi-

fied and clearly referenced in the text. Although including computer output in your text is *not* appropriate for formal writing such as a thesis, it is useful for the instructor in resolving any questions about how you generated your results. Stata graphs pasted into Mac Word often can't be viewed in Word for Windows, so if you're using Mac Word, please convert your file to a pdf and send it in that format (or both formats) so that I can see the graphs.

This assignment is intended to give you a taste of how we begin to explore the behavior of economic time series. There is much more that can be done and you are more than welcome to explore the big wide world of macro-econometrics! If you encounter situations that do not seem to fit the procedures that I've outlined in the assignment, you are welcome to come and talk with me about how you might go beyond the basic assignment to explore some additional aspects of the data.

Exercise #1: Getting started with Stata and your data

Loading and running Stata

Before you can do any analysis, you must make sure that Stata is installed on the computer you are using. The current version is Stata 11, which should be available on the computers in the IRCs and in the Public Policy Workshop (E110). Access to the PPW is available for students using computers in HSS courses such as this one. You will need to fill out a user agreement, at which time your Reed ID card will be put on the list to allow you access. The user agreement can be downloaded from the class Web site, printed out, signed, and returned to the PPW office in E112.

Stata is a fairly user-friendly software package. Stata 11 has a menu-driven interface in addition to the traditional typed-command mode, so you don't need to memorize a lot of arcane command names. If you use the command interface, it is rather picky about how you type them. Commas are never used in lists, but rather separate the command proper from the options at the end. Moreover, Stata's variable names are "case-sensitive," so GDP, Gdp, and gdp are treated as three different variables.

This is a course in macroeconomics, not econometrics, so you are not expected to learn Stata beyond using it as a tool to perform these projects. I will try very hard to give you all the details about Stata commands that you need to know in order to do the projects. However, if you do not understand something or if you want to learn more, Stata has a pretty good online help system and there are umpteen volumes of printed manuals that are available in the PPW and on reserve in the Reed Library. These manuals are now also available from the Help menu within Stata in pdf form. (Earlier Stata manuals will work, but they will be missing some features and options that were added in version 11.)

- (a) Loading and checking your data.** Download the data file (.dta) from the server to a local disk and open it. Customize it by deleting the data for all countries other than yours. Look at the series in the Stata data editor to see whether any years are missing and to make sure that the overall data series seem reasonable.

The assignment page for this assignment on the course Web site contains links to individual country datasets drawn from version 6.3 of the Penn World Table (PWT) data set (<http://pwt.econ.upenn.edu/>). Each dataset contains data from 1950 to 2007 (or some subset of this period) for one country among the 190 countries included in the PWT. (A version of the entire PWT is also linked there.) The file for each country (with a .dta tag) has three variables in it: the year, per-capita real GDP measured in 2005 U.S. dollar equivalents, and population (in millions). To open the data set in Stata, either double-click on it directly or use the open command from inside Stata. Once you open it, the Variables window on the left of the screen should have three entries in it corresponding to the variables you read: year, rpcgdp, pop. (The Variables window may be “docked” along the edge of the window. If so, you can click where it says Variables to open it.)

Open the data editor by clicking the button that looks like a table grid (fourth from the right?) on the Stata toolbar just below the menus. Check to make sure that your data were read correctly. The overall database extends from 1950 to 2007, but there may be some missing values (indicated by dots) for your country.

(b) Transforming the data. Create new variables for real GDP, its log, and its annually and continuously compounded growth rates. Look at them in the data editor to verify that they are created correctly.

The PWT has variables for per-capita real GDP (rpcgdp in your data set) and population (pop), but none for total real GDP. The Stata command “generate” (or gen for short) allows you to create a new variable that is a transformation of existing variables. To create a new variable called gdp for total real GDP, we multiply per-capita GDP times the population by typing **generate gdp=rpcgdp*pop/1000** and hitting Enter. (In this assignment, the **upright, bold, sans serif typeface** indicates a command that you type into Stata. However, Stata commands never end with a comma or period, so any trailing punctuation should *not* be typed.) The population variable in your data set is in millions and the variable for per-capita GDP is in internationally comparable dollars per person. Dividing by 1000 scales the GDP variable to billions of dollars. (If you make a mistake in a generate command, you may not be able to simply re-type it because the variable may already exist—generate can only be used to create *new* variables. To replace the values of an existing variable with newly generated ones, use the replace command, which has the same format as the generate command but starts with the word **replace**. Or you can **drop** the variable that you mistakenly generated and use **generate** again.) The generate and replace commands are probably easiest to do by just entering them on the command line, but you can use the Data menu if you wish.

When you succeed in generating a new variable, it will be added to the Variables window. If there are any missing values in your data, you will get a warning message in the Results window telling you that values for the new variable could not be computed for some observations. The command you typed is logged in the Review window. To copy a previously typed command back into the input line, click it in the Review window. To enter a variable name at the Command window cursor, click it in the Variables window.

Now create the variable `lgdp` containing the natural log of real GDP. You can do this by typing **generate lgdp=log(gdp)** in the command window.

In order to do time-series operations like lags and differences, you need to tell Stata about the time-orientation of your data. You do this by entering the command **tsset year**, which tells Stata that the variable `year` is the “time indicator” in your data set. You can also do this in the Statistics menu under Time Series → Setup & Utilities → Declare Dataset to Be Time Series Data.

Once you have set the time indicator, the prefixes `l.` (lower-case L) and `d.` will take lags and differences. For example, `l.gdp` is the lagged (previous year’s) value of the variable `gdp`, `l2.gdp` is the value of `gdp` two periods before, and `d.gdp` is the change in `gdp` from the previous period (which is equal to $\text{gdp} - \text{l.gdp}$). Use the generate command to create the variables $\text{ggdp} = (\text{gdp} - \text{l.gdp})/\text{l.gdp}$ and $\text{dlgdp} = \text{d.gdp}$. These are the annually compounded and continuously compounded growth rates of `gdp`. (See Chapter 3, section B of the coursebook for more information about compounding of growth rates.)

You should now have seven variables in the data set. Return to the data editor and check to see that they seem correct and that there are no cells that are unexpectedly empty.

(c) Interpreting graphs of GDP. Make separate graphs with GDP and log GDP on the vertical axis and (for both graphs) the year on the horizontal axis. Copy the graphs to your report. What would a linear graph mean in each case? Based on your graphs, how reasonable is the assumption that the growth rate is stable over time? Are there major business cycles that stand out for your country? What are the dates of the apparent peaks and troughs?

Graphing in Stata uses the `graph` command. There are limitless options that you can include and it is fairly easy to tailor the appearance of your graph using the menu and windows in Stata 11. Choose Time Series Graphs → Line Plots from the Graphics menu. Once the graph window opens, you can select one or more variables to plot (select the Plots tab and click Create to add a variable) and control other options of the graph such as tick-marks and lines on the axes, titles, and more. To create the graph and leave the graphics window open (perhaps to refine it or do another one), click the Submit button. Clicking the OK button creates the graph and closes the graph window.

Stata usually takes a few seconds to create the graph then it opens a new window containing the graph. If you want to change the properties of the graph, add text, or edit the graph in some other way, you can click the Start Graph Editor button or select this option from the File menu. The graph editor gives you myriad options for customizing your graph. (Note: the graph editor is not available in Stata 9 or earlier versions.)

Business-cycle peaks and troughs are (at least for annual data) just what they sound like. A peak year is a high point followed by a decline in GDP. A trough is a low point after which GDP turns upward. Recessions start at peaks and end at troughs. Expansions generally start at troughs and end at peaks. However, two years of contraction separated by one year of modest growth might be interpreted as a long recession rather than two cycles.

(d) Interpreting GDP growth rate. Calculate and interpret the average growth rates of real GDP for your economy, both continuously and annually compounded. Compare the two average growth-rates (annually vs. continuously compounded). Which one is larger than the other and why? How variable is the growth rate for your country? What are the highest and lowest values.

To get summary statistics for a series (in this case, your two growth series) use **summarize dlgdp ggdp**. This will give you mean, standard deviation, maximum and minimum values. Standard deviation is a common measure of the variability of the series. For a series following the normal probability distribution, about 68% of the observations will lie within one standard deviation on either side of the mean (average). If you prefer to use the menu-driven interface, go to the Statistics menu and select Summaries, Tables & Tests → Summary Statistics → Summary Statistics.

If you are having trouble understanding the comparison of the annually and continuously compounded series, think about which would imply a greater increase in GDP: 4 percent growth compounded annually or 4 percent growth compounded continuously.

Exercise #2: Explaining GDP growth with a linear trend

Stationarity of a time-series variable is a technically complicated but intuitively simple concept in time-series analysis. Most simply, a variable is stationary if it tends to return to a fixed level eventually. Any shocks that cause a stationary variable to move away from its average level will eventually die out, with the variable tending to return to its fixed mean value. In growing economies, GDP and $\ln$ GDP are nonstationary, because growth is ongoing and GDP will never tend to return to the level it had at the beginning of the sample. However, the GDP *growth rate* could be stationary if the average rate of growth tends to be stable. There are formal tests of stationarity such as the Dickey-Fuller test, but these tend to be pretty unreliable with data sets as short as ours.

For most countries, GDP and its log follow nonstationary processes: They tend to increase steadily over time. There are two main kinds of nonstationary time series. *Difference-stationary* series are ones where the first difference of (*i.e.*, the change in) the series is stationary; *trend-stationary* series are ones that vary around a deterministic time trend. The most commonly used trend-stationary processes would express the log of GDP as a linear function of time plus a stationary disturbance term e_t :

$$\ln GDP_t = \alpha_0 + \alpha_1 t + e_t. \quad (1)$$

The key assumption for a series to be trend-stationary is that any deviation from the trend e_t must die out over time, so that the series always tends to return to the fixed trend line, regardless how far above or below the line it strays. For example, if U.S. GDP is trend stationary, then the existence of the Great Depression, when GDP was far below trend, should not have any permanent effect on GDP, since it eventually recovered to the unaffected trend line.

In contrast, a difference-stationary process might be characterized as

$$\ln GDP_t - \ln GDP_{t-1} = \alpha_1 + e_t,$$

or

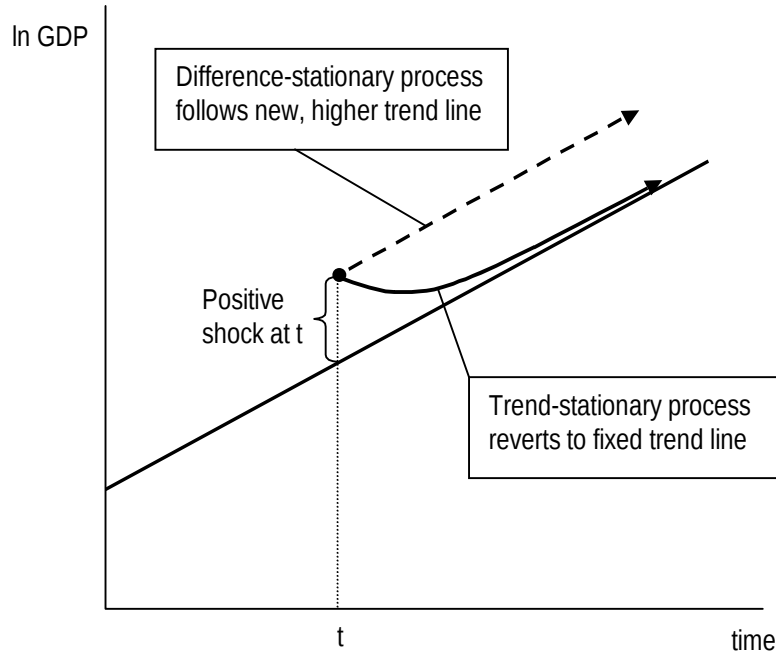
$$\ln GDP_t = \ln GDP_{t-1} + \alpha_1 + e_t. \quad (2)$$

Substituting backward in equation (2) yields

$$\ln GDP_t = \ln GDP_0 + \alpha_1 t + \sum_{s=1}^t e_s = \alpha_0 + \alpha_1 t + \sum_{s=1}^t e_s. \quad (3)$$

Comparing equation (1) with equation (3) shows that shocks never die out for a difference-stationary process as they do for a trend-stationary one. The deviation of current $\ln GDP$ from its trend in equation (3) is the sum of all the shocks that have occurred since time 0.

Figure 1



The distinction between a trend-stationary and a difference-stationary series is shown in Figure 1. At time t , a positive shock to GDP occurred, so GDP was above its original trend line. If GDP is trend-stationary, then the trend line itself will remain stable and GDP will tend to revert back toward it over time, as shown by the solid path. If GDP is difference-stationary, then the shock at t will permanently increase the level of GDP, shifting the trend line upward. The series would then tend to

grow along the higher, dotted path, with no tendency to revert to the original trend line. In this exercise, we will treat the variables as though they are trend-stationary.

The coefficients of equations like (1) can be estimated by the technique of *least-squares regression*. This technique fits a line through a scatter of data points with $\ln$ GDP on the vertical axis and time on the horizontal. The best-fit line is chosen to minimize the sum of the squared vertical distances of the data points from the line.

(a) Estimate a linear trend. Use Stata to estimate a trend line for $\ln$ GDP for your country using as long a time period as your data will allow. Save the “fitted values” from your regression and plot the actual and fitted values on the same graph against time. What is your estimate of the trend growth rate for your economy? How does this compare to the average (continuously compounded) growth rate? Why are they different? How much of the variation over time in $\ln$ GDP is due to the trend and how much is due to fluctuations around the trend?

The command for least-squares regression in Stata is **regress dvar ivars** where dvar is the dependent (left-hand) variable and ivars is a list of right-hand variables separated by blanks. In your application, the $\ln$ GDP variable is the dependent variable and year is the only independent variable. (A constant term will automatically be included.) If you prefer the menu-driven interface, the regression command can be found in the Statistics menu at Linear Regression and Related → Linear Regression. Type in or select the relevant variable names for the dependent and independent variables.

Recall that if a variable that grows at a constant rate, then its log will follow a straight-line path over time with the slope of the line equal to the continuously compounded growth rate. The slope of the trend line of $\ln$ GDP is therefore an estimate of the trend growth rate of GDP over the sample period. This value is reported in the Coef. column of the regression table on the line labeled “year.” The “R-square” statistic reported in your Stata output measures the fraction of the variance in the dependent variable that is explained by the independent variables in the regression. In this case, the R-square is the fraction explained by the trend and one minus R-square is the fraction due to variation around the trend.

(b) Testing for trend breaks. For many countries, a period of slower trend growth began about 1973. Does this seem to be the case for your country? Based on your results so far, choose a breakpoint that makes sense for your country (either 1973 or another) and estimate separate trends before and after the break. Compare the trend growth rates before and after, then perform a statistical test to examine whether the statistical evidence is strong enough to conclude that the growth rates differ. How does breaking the trend affect the share of variation explained by the trend vs. fluctuations around the trend? Why?

We can restrict Stata to use only part of the observations in the overall sample by including an “if clause” in the regress command. For example, to use only observations for which the year variable is less than 1973, you would type **regress dvar ivars if year < 1973**. The relational operator in the if

expression can be chosen from $\{<, >, <=, >=, ==, \text{ and } \sim=\}$. Note that you need two equal signs to test for equality.

A statistical test of equality of the slopes before and after the break can be performed in several ways. The following method is quite straightforward. In the exposition, I will assume that 1973 is the breakpoint. If you choose a different breakpoint, you will need to make the appropriate modifications.

Suppose that the trend line before the breakpoint is $\ln \text{GDP}_t = \alpha_0 + \alpha_1 t$. At the break, both of the coefficients may change. Suppose that the constant changes by an amount γ and the slope changes by δ . Then after the break, $\ln \text{GDP}_t = (\alpha_0 + \gamma) + (\alpha_1 + \delta) t$. We would like to know if $\delta = 0$.

To test whether $\delta = 0$, we define a “dummy” variable d_t that takes the value zero for the first part of the sample and one for the second part. You can create such a variable in Stata with the command **generate d = year>1972**, assuming that your data break after 1972. (The expression $\text{year} < 1972$ is a “Boolean” expression that takes the value one if the condition is true and zero otherwise.) With d_t defined in this way, we can rewrite the two trend lines as a single equation:

$$\ln \text{GDP}_t = \alpha_0 + \gamma d_t + (\alpha_1 + \delta d_t) t = \alpha_0 + \gamma d_t + \alpha_1 t + \delta d_t t. \quad (4)$$

We can run a regression on equation (4) if we construct the variable $d_t t$ by **generate dyear=d*year** using the d variable constructed above. Running **regress lgdp d year dyear** (using the full sample) would estimate the four coefficients of (4): α_0 , γ , α_1 , and δ . You can check to make sure you have done everything right by comparing your results to those of the two separate sub-sample regressions. The coefficients α_0 and α_1 on the constant and the year variable should be the same as the ones you estimated for the earlier sub-sample alone. The estimated coefficients for γ and δ (those on d and $dyear$) should equal the *difference between* the coefficients you estimated above for the two sample periods.

Since δ is the amount by which the early and late sample trend growth rates differ, we are interested in testing the hypothesis that $\delta = 0$. (If statistical hypothesis tests are unfamiliar to you, please see the box on the next page.) A statistical test of this hypothesis is printed out automatically in the Stata regression output. The column headed “t” in the regression table shows t statistics for the hypotheses that each coefficient is equal to zero. The entry for the row corresponding to the $dyear$ variable is a statistic testing $\delta = 0$. The “probability value” or p -value directly to the right of the t statistic tells you the result of the statistical test. If the p -value is less than 0.05, then we reject the hypothesis that $\delta = 0$ at the 5% level of significance. For our application, this would mean that the likelihood of seeing such a difference in estimated slopes if the underlying slopes were truly the same is less than 5%. If the p -value is greater than 0.05, then we accept the null hypothesis that the slopes are the same—we lack sufficient statistical evidence to reject it.

(If you are really ambitious and your data suggest doing so, you could test for multiple trend breaks. However, be careful because you need enough observations in each section to estimate a trend reliably.)

Statistical Tests of Hypothesis

Formal statistical analysis usually involves testing hypotheses. In a hypothesis test, we formulate a specific hypothesis, called the *null hypothesis* and often denoted H_0 , and then determine whether we have sufficient statistical evidence to conclude that our sample of data is *inconsistent* with the null hypothesis.

We first choose a statistical significance level, often 0.05 or 0.01. We then determine how likely it is that we would observe a sample whose characteristics differ from those expected under the null hypothesis, if in fact the null hypothesis were true. This probability is called the *p-value* and is often printed out by statistical software. If the p-value is less than our chosen significance level then we reject the null hypothesis, otherwise we must “accept” it.

For example, we often estimate a linear regression such as $y = \alpha + \beta x + e$, where y is the dependent variable, x is the independent variable, e is the error term, and α and β are coefficients to be estimated. We are usually interested in whether x has a significant effect on y . We formulate a test for this by specifying the null hypothesis as $H_0: \beta = 0$. For this example, we choose a significance level of 0.05.

If the null hypothesis is true, we would expect that our estimate of β would be near zero, but because of random variation in the sample we never expect our estimate to be *exactly* zero. The question we pose in a statistical test is “How likely is it that the observed deviation from the null hypothesis is due to random variation?” In this case, “How likely is it that we would observe this large a β (either positive or negative) if in fact $\beta = 0$?” This probability is the *p-value*. If the *p-value* is less than 0.05, then less than 5% of the time would we find such a large β estimate resulting from random chance in a world where the true β is zero. Thus, when the *p-value* is less than 0.05, we reject the null hypothesis and conclude that $\beta \neq 0$ and that x does affect y .

To summarize, the steps in a statistical hypothesis test are (1) specify the null hypothesis and significance level, (2) compute a test statistic that allows us to test the null hypothesis, (3) compute the *p-value* associated with that test statistic, and (4) form our conclusion to either accept or reject the null hypothesis depending on whether the *p-value* is greater than or less than our significance level.

(c) Calculate trend and cyclical components of ln GDP. Calculate and plot the fitted values from the trend regression to show the path of ln GDP along with its trend line. Use a split trend from part (b) or the single trend from part (a) depending on which is more appropriate. Calculate and plot the cyclical component of ln GDP as the difference between actual and trend ln GDP. Identify the periods during which ln GDP was significantly above and below its trend. Compare the business cycles emerging from the trend analysis to the turning points you identified earlier.

After you have run the regress command, Stata stores information about the regression in memory. You can then go back and use the predict command to create a data series containing the esti-

mated or fitted values based on the regression equation. (Predict uses the most recently estimated equation, so make sure that you re-do the estimation if you have estimated any other equations since you did the original regress operation. You will need to be especially careful if you are using a split sample.) For the trend regression, this series is the “trend component” of $\ln$ GDP. The format of the command is **predict tlgdp** where tlgdp is the name you choose for the variable that will contain the fitted (trend) values.

Plotting actual GDP and trend GDP against the year variable is similar to what you did in Exercise 1 except both variables are listed as yvars. With the menu interface, use Graphics → Time Series Graphs → Line Plots as before, but be sure to include both the actual (lgdp) and predicted (tlgdp) variables as Plot 1 and Plot 2.

Calculating the cyclical component uses the generate command described earlier: **generate clgdp = lgdp - tlgdp**. Use the same methods as above to plot this variable over time as a line plot. Plotting the cyclical component on the same graph as total GDP or the trend may be problematic because of scaling. The cyclical component will vary in a fairly narrow range above and below zero, whereas the others have very large values. You can use separate graphs for cyclical, but if you decide for some reason to put them on the same graph, you should use the option that allows a second, “right” y-axis in order to scale the cyclical series differently from the others.

Exercise #3: A more “flexible” trend

The constant underlying growth rate implied by a log-linear trend is not appropriate for many countries. Even a piece-wise log-linear trend such as the one you estimated in part (b) of Exercise #2 cannot capture smooth changes in the underlying trend growth rate. Economists have devised several more flexible ways of capturing the trend in a series that allow the trend growth rate to vary smoothly over time. One of the most popular is the Hodrick-Prescott (HP) filter.

Suppose that y is log GDP and let s be the underlying trend series. Linear regression calculates the trend series to be the straight line that minimizes the squared deviations of y from its prediction s .

Thus, the linear trend minimizes $\sum_{t=1}^T (y_t - s_t)^2$ subject to s_t being a linear function of time, which

means that the difference between s_t and s_{t-1} is constant over all t , or $(s_{t+1} - s_t) - (s_t - s_{t-1}) = 0$. This last equation can be thought of as the constraint that imposes linearity on the trend.

The idea of the HP filter is that we want to impose the linearity constraint “softly” rather than absolutely, allowing some deviations from perfect linearity of the trend but “penalizing” them in some way. The way that the HP filter estimates the flexible trend series s is to minimize

$$\sum_{t=1}^T (y_t - s_t)^2 + \lambda \sum_{t=2}^{T-1} ((s_{t+1} - s_t) - (s_t - s_{t-1}))^2. \quad (5)$$

The second summation in (5) will be exactly zero for a linear trend, but differs from zero if the slope of the flexible trend is not constant. The HP filter thus minimizes the sum of two sums of squares.

The first is the traditional squared deviations of the actual series from the fitted trend series. The second is the deviations of the trend series itself from linearity.

The weight λ is used to adjust the relative importance of these two criteria. The larger is λ the more tightly the HP trend will be constrained to be linear. If $\lambda = 0$, then we would minimize only the first summation and would do so by setting $s = y$ for all t . In this case “trend” series is just the series itself and there is no linearity (or trend in the traditional sense) at all. As $\lambda \rightarrow \infty$, the linearity constraint becomes perfectly binding and we minimize subject to a linearity constraint. In this case, the HP trend is identical to the linear trend estimated with standard regression.

For intermediate values of λ , we get a trend series that is, in smoothness, somewhere between the perfectly smooth linear trend and the perfectly “unsmooth” series itself. The larger the value of λ we choose, the closer the resulting series is to the linear trend. The smaller the value of λ , the more that the year-to-year variation in the series will be admitted into the trend.

Hodrick and Prescott recommend choosing $\lambda = 1600$ for quarterly series and argue that the value of λ should vary with the square of the frequency of the series. This implies λ of 100 for annual data and 14,400 for monthly data. However, Ravn and Uhlig have argued that λ should vary with the fourth power of the frequency, so $\lambda = 6.25$ is appropriate for annual data and 129,600 would be best for monthly data.

- ||| (a) Apply the Hodrick-Prescott filter to your log GDP series using $\lambda = 100$, then again using $\lambda = 6.25$. Plot each of the trended series on a graph with the log GDP series (two plots, two series on each plot). Comment on how the two HP trends compare with each other and with the linear trend series you computed in Exercise #2.

The Hodrick-Prescott filter command **hprescott** is an add-on that must be downloaded from the Stata Web site before use. If the Angels of CUS have performed their anticipated miracles on our behalf, this should be completed on both the IRC and PPW computers. If not, you may have to try to do this yourself. If you need to do it, go to the Stata help menu and select Search.... Click the button that says “Search all” and type **hprescott** in the box. You should see a line that says “hprescott from <http://fmwww.bc.edu/RePEc/bocode/h>.” Click this line and you should be able to download and install the **hprescott** procedure. (However, permissions may be an issue on campus computers, so you may need help from CUS if it gets to that point.)

The format of the **hprescott** command is **hprescott lgdp , smooth(100) stub(H1)**. This will calculate an HP trend series based on **lgdp** with $\lambda = 100$ and place the smoothed series in **H1_lgdp_sm_1** and the cyclical series in **H1_lgdp_1**, two newly created variables. Changing the λ value to 6.25 is accomplished by replacing the value of the **smooth** parameter, but you’ll also need to change the **stub** parameter to **H6** or something else in order to have different names for these series.

- ||| (b) Examine the two cyclical series created by the HP filter. Comment on how the cyclical series compare to each other and to those implied by the linear trend. Are the implied cyclical fluctuations (booms and recessions) in the same place for all three series? Are they of the same severity?

You may want to plot the three cyclical series together on the same graph in order to facilitate your comparison.

- ||| (c) Which method of trend estimation seems to you to capture best the underlying long-run movement in the log GDP series for your country? Why? (Are there important changes in the trend growth rate for your country? Are they smooth or sudden?) Write a short paragraph in layman's language describing your conclusions about your country's growth path during the sample and the cyclical fluctuations around that path.