

Chapter 6

Relaxing Assumptions

Flow of study

For this chapter, we are going to relaxing some assumptions we imposed since we estimate the first model. These topics will be covered: (1) multicollinearity (2) heteroscedasticity (3) autocorrelation.

For each section, we will explore

- The nature of the assumption or problem
- Effects on the estimated coefficients and variance, also standard error.
- How to detect the problem.
- What are remedial measure(s).

Further reading can be found in Gujarati and Porter, Chapter 10-13.

(1) *Nature of the assumption*

This is a simplified version, compared to the book, of multicollinearity problem. Let λ_i be a constant, consider the following argument when X_{2i} and X_{3i} are linearly correlated perfectly, or **perfect collinearity**.

$$\circ \lambda_2 X_{2i} + \lambda_3 X_{3i} = 0$$

when $\lambda_i \neq 0$ for all i simultaneously. Another relation that describe a non-perfect collinearity, serious collinearity or **multicollinearity**, is

$$\circ \lambda_2 X_{2i} + \lambda_3 X_{3i} + v_i = 0$$

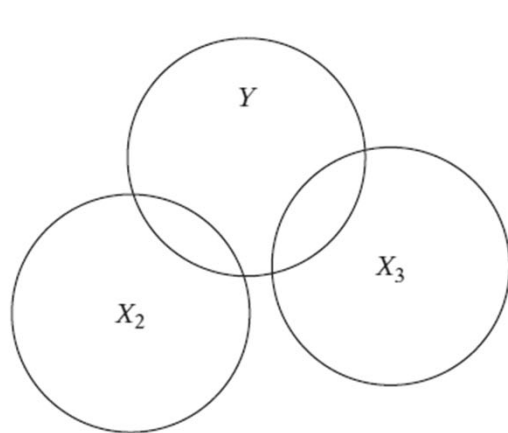
where v_i is a stochastic error term, then,

$$\circ X_{2i} = -\frac{\lambda_3}{\lambda_2} X_{3i} \text{ for the first equation.}$$

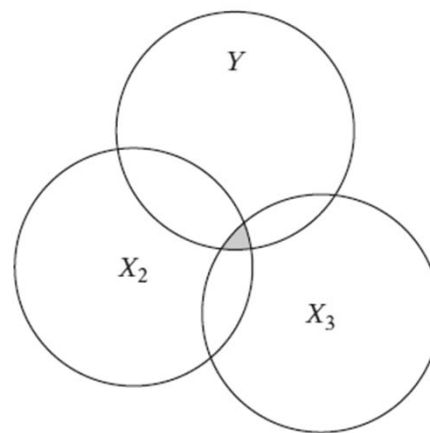
$$\circ X_{2i} = -\frac{\lambda_3}{\lambda_2} X_{3i} - v_i \text{ for the second equation.}$$

6.2 Multicollinearity

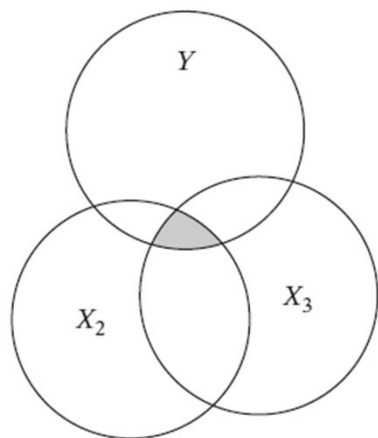
(1) Nature of the assumption



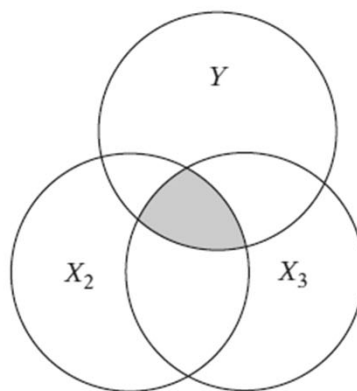
(a) No collinearity



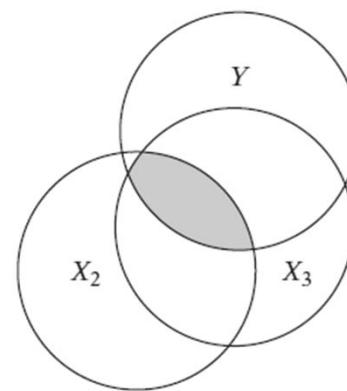
(b) Low collinearity



(c) Moderate collinearity



(d) High collinearity



(e) Very high collinearity

(1) *Nature of the assumption*

Precisely taken from the book, here are some causes of multicollinearity.

- **Data collection method** may limit range of values taken in the independent variables.
- **Constraints on the model or in the population being sampled.** i.e. income and house size tend to be correlated.
- **Model specification.** i.e. including a polynomial term especially when the range of X is small.
- **Overdetermined model** or when k is larger than n .
- **Common trend.** i.e. a time-series data consist of consumption expenditure, income, wealth, and number of population.

(1) *Nature of the assumption*

Looking from another perspective apart from stated above, multicollinearity is seen particularly as sampling problem, not a problem on a population, since when we postulate population regression, X variables included in a model have a separate or independent influence on Y .

Meaning that, as Goldberger coined the term, this may be considered as **micronumerosity** problem when our sampling may not be “rich” enough to capture X variability.

By the way, micronumerosity refers to the problem of small sample size.

Another important note is that multicollinearity is **much more common** in cross-sectional data compared to time-series data, in which autocorrelation is much more common.

(2) *Effects on estimation*

1. Perfect collinearity

Recall that the estimated coefficients are

$$\circ \hat{\beta}_2 = \frac{(\sum y_i x_{2i})(\sum x_{3i}^2) - (\sum y_i x_{3i})(\sum x_{2i} x_{3i})}{(\sum x_{2i}^2)(\sum x_{3i}^2) - (\sum x_{2i} x_{3i})^2}$$

$$\circ \hat{\beta}_3 = \frac{(\sum y_i x_{3i})(\sum x_{2i}^2) - (\sum y_i x_{2i})(\sum x_{2i} x_{3i})}{(\sum x_{2i}^2)(\sum x_{3i}^2) - (\sum x_{2i} x_{3i})^2}$$

If we assumed that $X_{3i} = \lambda X_{2i}$, replacing this into $\hat{\beta}_2$, we get,

$$\circ \hat{\beta}_2 = \frac{(\sum y_i x_{2i})(\lambda^2 \sum x_{2i}^2) - (\lambda \sum y_i x_{2i})(\lambda \sum x_{2i}^2)}{(\sum x_{2i}^2)(\lambda^2 \sum x_{2i}^2) - \lambda^2 (\sum x_{2i}^2)^2} = \frac{0}{0}$$

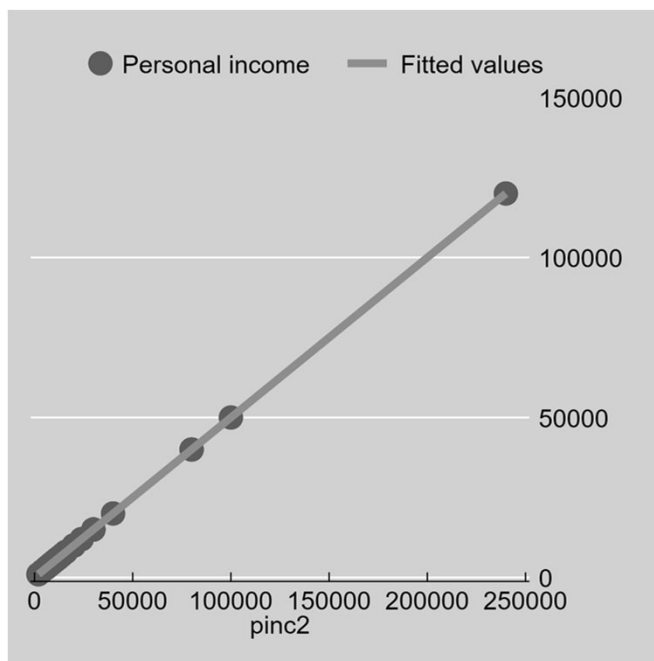
6.2 Multicollinearity

(2) Effects on estimation

Example: consider the standard income-consumption model with another added variable intentionally ($pinc2_i$) which is personal income x 2, the model becomes

$$pinc_i = \beta_1 + \beta_2 pinc_i + \beta_3 pinc2_i + u_i$$

This is the plot and correlation matrix of $pinc_i$ and $pinc2_i$



	pinc	pinc2
pinc	1.0000	
pinc2	1.0000	1.0000

6.2 Multicollinearity

(2) *Effects on estimation*

Throwing this model into STATA, we have the regression result as follows. We can see that STATA automatically rejects, or omits, one of these variables immediately because it cannot be estimated.

Source	SS	df	MS	Number of obs	=	51
Model	1.0334e+09	1	1.0334e+09	F(1, 49)	=	40.85
Residual	1.2396e+09	49	25297763	Prob > F	=	0.0000
Total	2.2730e+09	50	45460031.4	R-squared	=	0.4546
				Adj R-squared	=	0.4435
				Root MSE	=	5029.7

pexp	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
pinc	0	(omitted)				
pinc2	.1257751	.0196788	6.39	0.000	.086229	.1653211
_cons	3886.008	845.2106	4.60	0.000	2187.493	5584.522

(2) Effects on estimation

2. Multicollinearity

Given that $X_{3i} = \lambda X_{2i} + v_i$, replacing this into $\hat{\beta}_2$, we get,

$$\circ \hat{\beta}_2 = \frac{(\sum y_i x_{2i})(\lambda^2 \sum x_{2i}^2 + \sum v_i^2) - (\lambda \sum y_i x_{2i} + \sum y_i v_i)(\lambda \sum x_{2i}^2)}{(\sum x_{2i}^2)(\lambda^2 \sum x_{2i}^2 + \sum v_i^2) - \lambda^2 (\sum x_{2i}^2)^2}$$

If v_i is approaching zero, the more this will be closer to perfect multicollinearity.

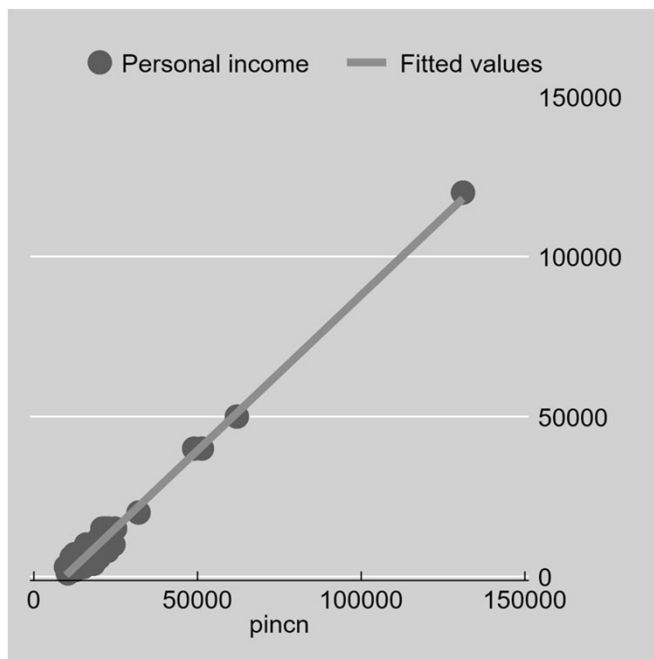
Example: consider the same model but now with another added variable intentionally ($pincn_i$) which is personal income + a randomly normally generated value with 10,000 mean and 2,000 standard deviation, the model becomes

$$\circ pinc_i = \beta_1 + \beta_2 pinc_i + \beta_3 pincn_i + u_i$$

6.2 Multicollinearity

(2) Effects on estimation

This is the plot and correlation matrix of $pinc_i$ and $pincn_i$.



	pinc	pincn
pinc	1.0000	
pincn	0.9922	1.0000

6.2 Multicollinearity

(2) Effects on estimation

Now compare the result between normal model and the model with multicollinearity.

Source	SS	df	MS	Number of obs	=	51
				F(1, 49)	=	40.85
Model	1.0334e+09	1	1.0334e+09	Prob > F	=	0.0000
Residual	1.2396e+09	49	25297763	R-squared	=	0.4546
				Adj R-squared	=	0.4435
Total	2.2730e+09	50	45460031.4	Root MSE	=	5029.7

pexp	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
pinc	.2515502	.0393576	6.39	0.000	.172458	.3306423
_cons	3886.008	845.2106	4.60	0.000	2187.493	5584.522

Source	SS	df	MS	Number of obs	=	51
				F(2, 48)	=	20.03
Model	1.0340e+09	2	517012377	Prob > F	=	0.0000
Residual	1.2390e+09	48	25812017	R-squared	=	0.4549
				Adj R-squared	=	0.4322
Total	2.2730e+09	50	45460031.4	Root MSE	=	5080.6

pexp	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
pinc	.3003663	.3191086	0.94	0.351	-.3412446	.9419772
pincn	-.0481394	.3122331	-0.15	0.878	-.6759262	.5796473
_cons	4347.126	3110.289	1.40	0.169	-1906.529	10600.78

(2) *Effects on estimation*

1. Coefficients estimated are still BLUE.

2. Large variances and covariances of OLS estimators.

Recall that the variance of estimated coefficients can be written into this form.

$$\circ \text{var}(\hat{\beta}_2) = \frac{\sigma^2}{\sum x_{2i}^2 (1 - r_{23}^2)}$$

$$\circ \text{var}(\hat{\beta}_3) = \frac{\sigma^2}{\sum x_{3i}^2 (1 - r_{23}^2)}$$

where r_{23}^2 is the coefficient of correlation between X_2 and X_3 . The value is between 0 and 1, as an absolute value.

We can see clearly that when the correlation between X_2 and X_3 gets **higher**, the denominator will be **smaller**, leading to **higher** variance.

(2) *Effects on estimation*

If we separate a part of $\text{var}(\hat{\beta}_2)$ like this,

$$\circ \text{var}(\hat{\beta}_2) = \frac{\sigma^2}{\sum x_{2i}^2} \cdot \frac{1}{(1-r_{23}^2)}$$

we can define the latter part as **variance-inflating factor** or VIF

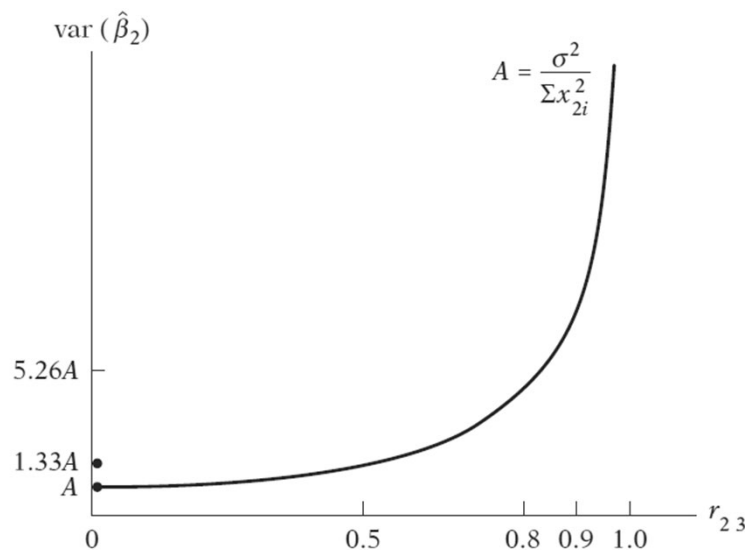
$$\circ \text{VIF} = \frac{1}{(1-r_{23}^2)}$$

The higher r_{23}^2 , the higher it is for VIF. We can also define the inverse of VIF as **tolerance** or TOL as

$$\circ \text{TOL} = \frac{1}{\text{VIF}} = (1 - r_{23}^2)$$

Note that these definition can be generalized for any pair of regressors.

6.2 Multicollinearity

(2) Effects on estimation

Since the variance is used to construct confidence interval, CI is stretched outwards and the t-curve is flatter.

Then, the acceptance region also becomes larger when variance is high. It is more likely that we would accept the null hypothesis when we should reject, causing more-likely type-II error.

(2) *Effects on estimation*

3. High R^2 but few significant t ratios

The coefficient of determination or R^2 from a model with multicollinearity is likely to be high, also the F-stat of overall model test, since the estimation is 'tricked' to have similar regressors with more explanatory power.

We can also see from the previous example that when we intentionally add a colinear variable into the regression, R^2 is higher.

However, this is not due to more explanatory power of regressors, but multicollinearity. Therefore, each coefficient is not very likely to be significant.

4. Sensitivity of coefficient due to small changes in data.

You can read for this example in page 331. In conclusion, a very slight change in data will affect the value of coefficient tremendously. Some of the direction of coefficient can be different from theoretical speculation.

It would be beneficial to read an illustrative example from page 332 to 337.

(3) *Detecting multicollinearity*

There are several ways to detect multicollinearity. Some of them are mentioned earlier. Some of them from the book are skipped.

Firstly, the phrase “**Rule of Thumb**” should be introduced. It refers to a specific level of criterion that is usually and mutually accepted as a threshold. More illustrative examples later here.

1. *Conflicting test*

This is already mentioned that when our estimation reports high R^2 or F value, but rarely coefficients are significant, we should suspect that there might be multicollinearity problem.

2. *Pair-wise correlation among regressors and scatter plot*

Another easy method to detect multicollinearity is to perform a pair-wise correlation on all regressors. (Easy when using STATA) The rule of thumb suggests that coefficient of correlation **exceeding 0.8** can be problematic and researcher may seek a remedial approach.

(3) Detecting multicollinearity

3. Auxiliary regressions

Regressing X_i on other X variables and obtain R_i^2 from the estimation then calculate

$$F_i = \frac{R_{xi \cdot x_2 x_3 \dots x_k}^2 / (k-2)}{(1 - R_{xi \cdot x_2 x_3 \dots x_k}^2) / (n-k+1)}$$

where k is the number of independent variables including intercept.

If F_i exceeds critical value from chosen level of significant, it means that X_i is collinear with other X .

Instead of testing all the auxiliary R_i^2 , **Klein's rule of thumb** suggests that multicollinearity is troublesome if the R_i^2 is greater than R^2 from another model that we regress Y on these X_i and other X .

4. VIF and TOL

Again, we follow the rule of thumb that VIF should not exceed 10, which will happen when $r_{23}^2 = 0.8$, while TOL should be closer to 1 rather than 0.

6.2 Multicollinearity

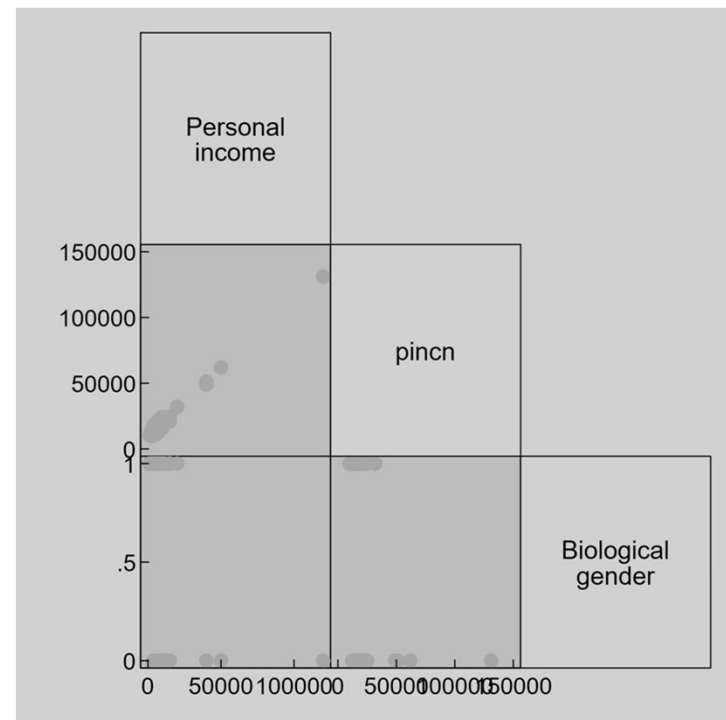
(3) Detecting multicollinearity

Example: using the basic income-consumption model with some additional variables.

$$pinc_i = \beta_1 + \beta_2 pinc_i + \beta_3 pincn_i + \beta_4 sex_i + u_i$$

We can summon scatter plot and correlation matrix pre-estimation.

	pinc	pincn	sex
pinc	1.0000		
pincn	0.9922	1.0000	
sex	-0.4230	-0.4193	1.0000



6.2 Multicollinearity

(3) Detecting multicollinearity

Or we can find VIF and tolerance post-estimation.

Source	SS	df	MS	Number of obs	=	51
				F(3, 47)	=	15.19
Model	1.1191e+09	3	373035915	Prob > F	=	0.0000
Residual	1.1539e+09	47	24550932.4	R-squared	=	0.4923
				Adj R-squared	=	0.4599
Total	2.2730e+09	50	45460031.4	Root MSE	=	4954.9

pexp	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
pinc	.2643351	.311817	0.85	0.401	-.3629599	.89163
pincn	-.0458366	.3045128	-0.15	0.881	-.6584374	.5667641
sex						
Female	-3194.168	1715.815	-1.86	0.069	-6645.942	257.6059
_cons	7042.464	3361.184	2.10	0.042	280.6343	13804.29

Variable	VIF	1/VIF
pinc	64.68	0.015461
pincn	64.43	0.015521
1.sex	1.22	0.821046
Mean VIF	43.44	

(4) Remedial measures

1. Do nothing

If and only if when we do not any other choice than using deficient data set. Also, we may not be able to draw any meaningful insight from the regression.

2. *Priori* information

Supposed we have an income- consumption model as such,

$$\circ Y_i = \beta_1 + \beta_2 \text{income}_i + \beta_3 \text{wealth}_i + u_i$$

we know that income and wealth are highly colinear. If we know that from previous empirical work

$$\circ \beta_3 = 0.1\beta_2 \text{ then}$$

$$\circ Y_i = \beta_1 + \beta_2 \text{income}_i + 0.1\beta_2 \text{wealth}_i + u_i \text{ and}$$

$$\circ Y_i = \beta_1 + \beta_2 X_{2i} + u_i$$

where $X_{2i} = \text{income}_i + 0.1\text{wealth}_i$, we can eliminate one of the variables.

(4) Remedial measures

3. Combining cross-sectional and time-series data

The most completed data would be panel data, repeated samples over time. If that is not possible to obtain, we may use '**pooled-data**', combining multiple waves of data into one large data set.

4. Dropping a variable(s)

This is the easiest method, and probably the best if possible. If we are sure which variable should be dropped, according to wrong specification of a model, dropping one of them is very easy and efficient.

Consider dropping $hein_i$ from the show size model, our results would be as follows.

(4) Remedial measures

5) Adding more observations

When possible, more observations lead to more variability in X and therefore, may lead to reduction of severity of multicollinearity problem.

6) Variable transformation

Transforming variables can be complicated, we can either perform

- **Ratio transformation** which will lead us to another problem of heteroscedasticity or
- **First difference form** of variable which is popular in time-series data.

6.3 Heteroscedasticity

(1) Nature of the assumption

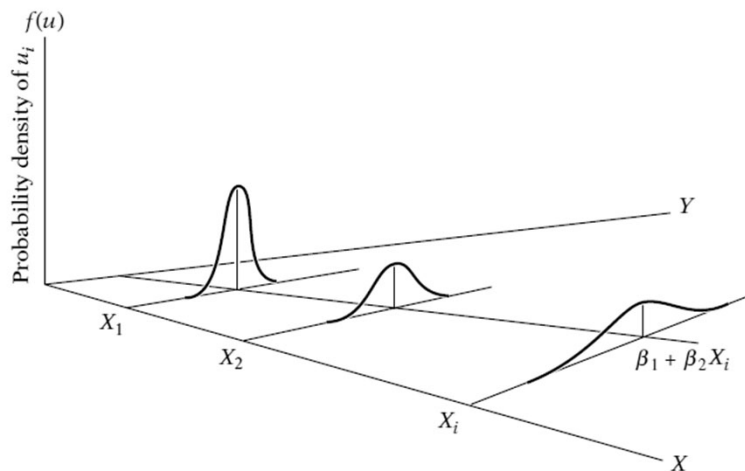
Recall the assumption of homoscedasticity or

$$\circ E(u_i^2|X_i) = \sigma^2$$

when this assumption is relaxed, we have

$$\circ E(u_i^2|X_i) = \sigma_i^2$$

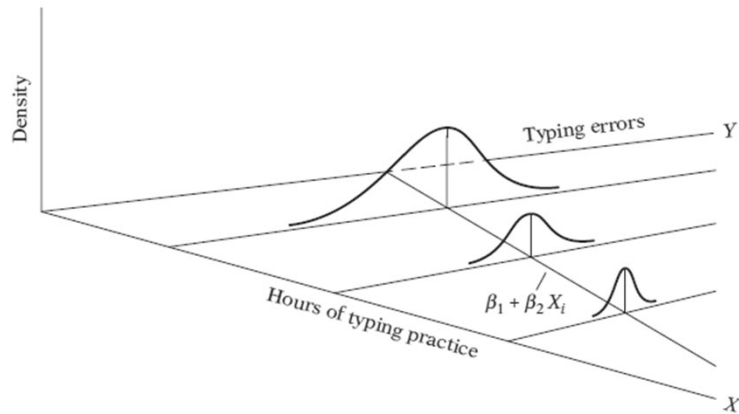
which means that we allow the error term scattered around each X_i to be different. Two classic examples are income-saving model and error-learning model as follows.



As people get richer, they have more choices over their consumption-saving, leading to a larger variance of saving (Y_i) on larger income (X_i).

6.3 Heteroscedasticity

(1) Nature of the assumption



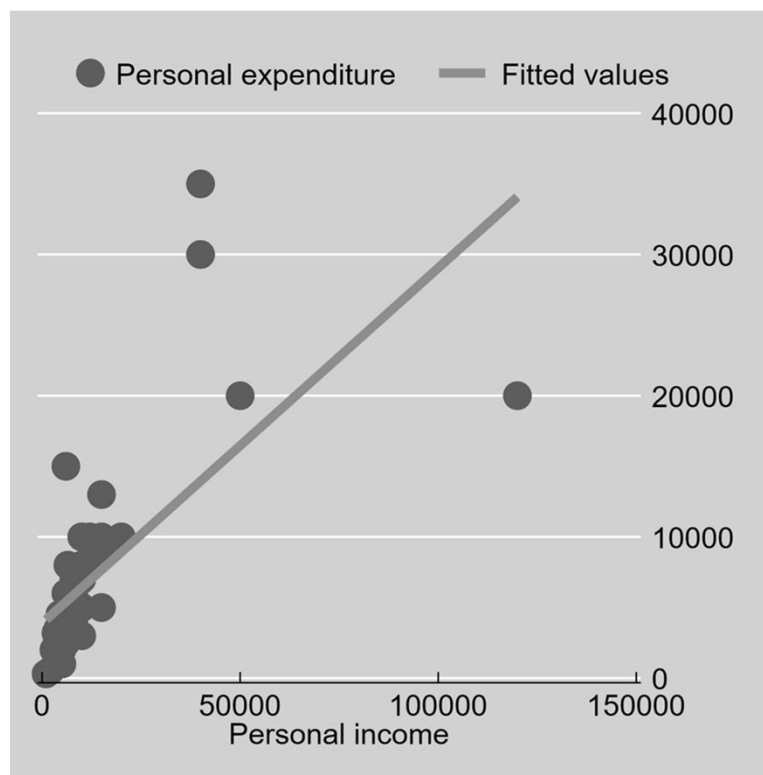
Similarly, people practicing more hours on typing leads to lower typing errors. However, some people maybe pretty good at typing at first and some can be pretty bad due to their familiarity to keyboard layout or language.

Apart from the nature of data, there are some other causes for heteroscedasticity.

(1) *Nature of the assumption*

1. *Presence of outliers*

An outlier is an observation that is much different from the rest, either little or largely different. Inclusion and exclusion of an outlier can alter the result of a regression substantially.



(1) *Nature of the assumption*

2. *Specification error*

For example, a price demand model can be heteroscedastic if we do not include price of complementary or competing commodities. Other commodities' price may be the source of scattered quantity demanded at some level.

3. *Skewness of distribution*

For example, plotting wealth and education level can show this problem because there are fewer people with high wealth, leading to lower variance in education level, while larger groups of population with lower wealth.

4. *Incorrect data transformation and functional form.*

Details for this topic will be skipped on this point.

Heteroscedasticity is more likely to be found in cross-sectional data rather than time-series since they deal with members of a population at a given point of time.

(2) *Effects on estimation*

We begin our examination on simple linear regression, recall that with homoscedasticity assumption yields the variance of the estimator as

$$\circ \text{var}(\hat{\beta}_2) = \frac{\sigma^2}{\sum x_i^2}$$

Relaxing the assumption, the variance becomes

$$\circ \text{var}(\hat{\beta}_2) = \frac{\sum x_i^2 \sigma_i^2}{(\sum x_i^2)^2}$$

$\hat{\beta}_2$ is not BLUE, with heteroscedasticity. It is still linear and unbiased but it is not efficient anymore, comparing to deriving estimator from another method called **generalized least squares (GLS)**.

(Note that an estimator not being efficient meaning that $\text{var}(\hat{\beta}_i)$ is not the lowest.)

(2) *Effects on estimation*

To estimate with GLS, we start from the basic model of

- $Y_i = \beta_1 + \beta_2 X_i + u_i$

Then we transform these variables by dividing by σ_i , if this standard error is known.

- $\frac{Y_i}{\sigma_i} = \frac{\beta_1}{\sigma_i} + \frac{\beta_2 X_i}{\sigma_i} + \frac{u_i}{\sigma_i}$

Then figure out the variance, we get

- $\text{var} \left(\frac{u_i}{\sigma_i} \right) = E \left(\frac{u_i}{\sigma_i} \right)^2 = \frac{1}{\sigma_i^2} E(u_i^2)$ and if σ_i is known

- $\text{var} \left(\frac{u_i}{\sigma_i} \right) = \frac{1}{\sigma_i^2} (\sigma_i^2) = 1$

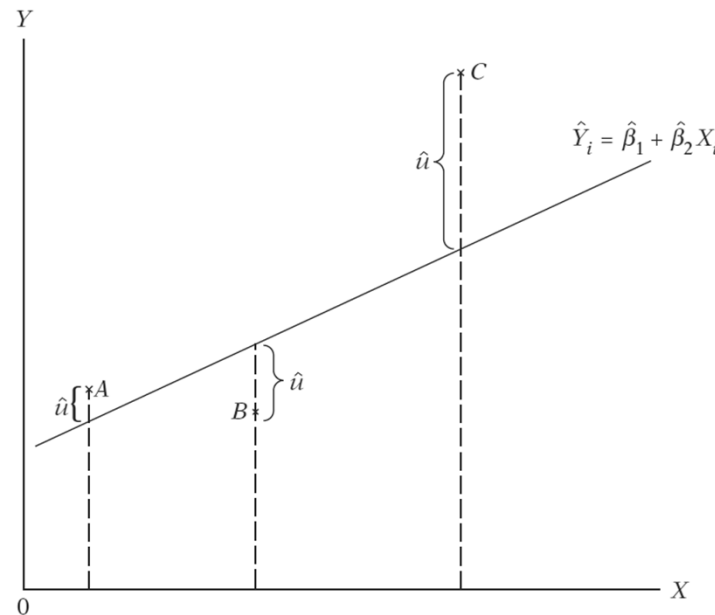
which is actually homoscedastic. To retrieve the estimators, we follow least squares method as usual.

6.3 Heteroscedasticity

(2) *Effects on estimation*

$$\circ \min_{\hat{\beta}_1, \hat{\beta}_2} \sum \left(\frac{u_i}{\sigma_i} \right)^2 = \sum \left(\frac{Y_i}{\sigma_i} - \frac{\hat{\beta}_1}{\sigma_i} - \frac{\hat{\beta}_2 X_i}{\sigma_i} \right)^2$$

This method is specifically called **weight least squares (WLS)**, weighing each term especially the error term with the error term itself. Estimators derived from this estimation are called **WLS estimators**. WLS is a class of GLS.



(2) *Effects on estimation*

WLS estimators derived will have less variance compared to ordinary OLS. Therefore, estimators from OLS is not with the least variance anymore, losing the quality of being efficient.

Consequences of relying on OLS are as follows.

1) OLS estimation allowing heteroscedasticity

- Using OLS while assuming σ_i^2 are known, $\text{var}(\hat{\beta}_2)$ is larger compared to the variance from WLS.
- Larger CI and t value is small, leading to insignificant conclusion.

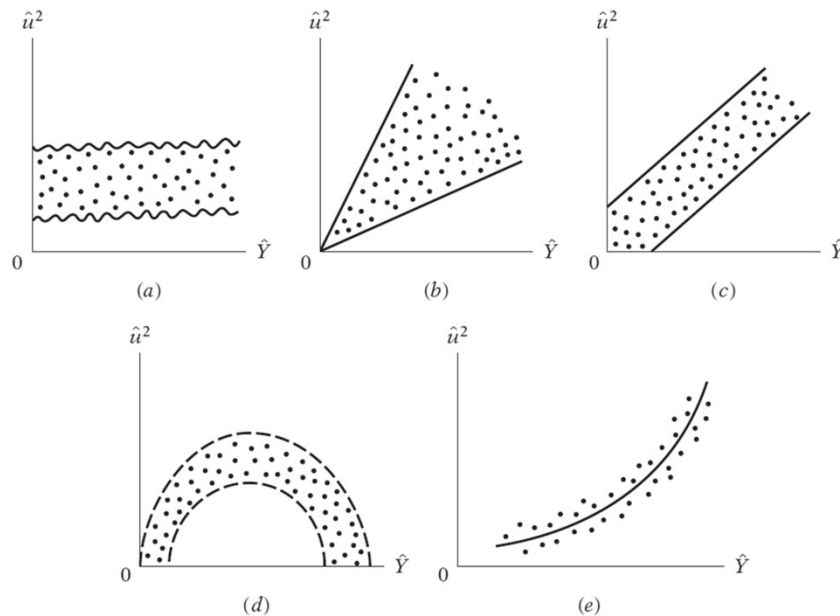
2) OLS estimation disregarding heteroscedasticity

- Using OLS while assuming homoscedasticity (σ^2) when heteroscedasticity is present, $\text{var}(\hat{\beta}_2)$ will be biased.
- We do not know whether the bias is positive (overestimate: actual variance is lower) or negative (underestimate: actual variance is higher), depending on the relationship between σ_i^2 and X_i .
- Conclusion or inference drawn from hypothesis tests may be misleading.

(3) Detecting heteroscedasticity

1. Graphical method

- **Step 1:** Estimate coefficients with OLS.
- **Step 2:** Retrieve $\hat{u}_i^2$. (In STATA, look for predicting residual)
- **Step 3:** Plot $\hat{u}_i^2$ with X_i or $\hat{Y}_i$. There might be multiple X_i in our regression function, so we can rely on $\hat{Y}_i$ as well.



Which of these plots that heteroscedasticity is present?

(3) Detecting heteroscedasticity

2. Park Test

The intuition of Park test assumes that when heteroscedasticity is present, σ_i^2 is a kind of function of X_i . Steps are as follows.

Step 1: Estimate coefficients with OLS and retrieve $\ln \hat{u}_i^2$

Step 2: Estimate $\ln \hat{u}_i^2 = \alpha + \beta \ln X_i + v_i$

Step 3: Test the significance from zero of β . If it is, it would suggest that heteroscedasticity is present.

Example: Revisit the income-consumption model with only one explanatory variable as follows.

- $pexp_i = \beta_1 + \beta_2 pinc_i + u_i$

6.3 Heteroscedasticity

(3) Detecting heteroscedasticity

○ Original model

Source	SS	df	MS	Number of obs	=	51
				F(1, 49)	=	40.85
Model	1.0334e+09	1	1.0334e+09	Prob > F	=	0.0000
Residual	1.2396e+09	49	25297763	R-squared	=	0.4546
				Adj R-squared	=	0.4435
Total	2.2730e+09	50	45460031.4	Root MSE	=	5029.7

pexp	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
pinc	.2515502	.0393576	6.39	0.000	.172458	.3306423
_cons	3886.008	845.2106	4.60	0.000	2187.493	5584.522

○ Park test

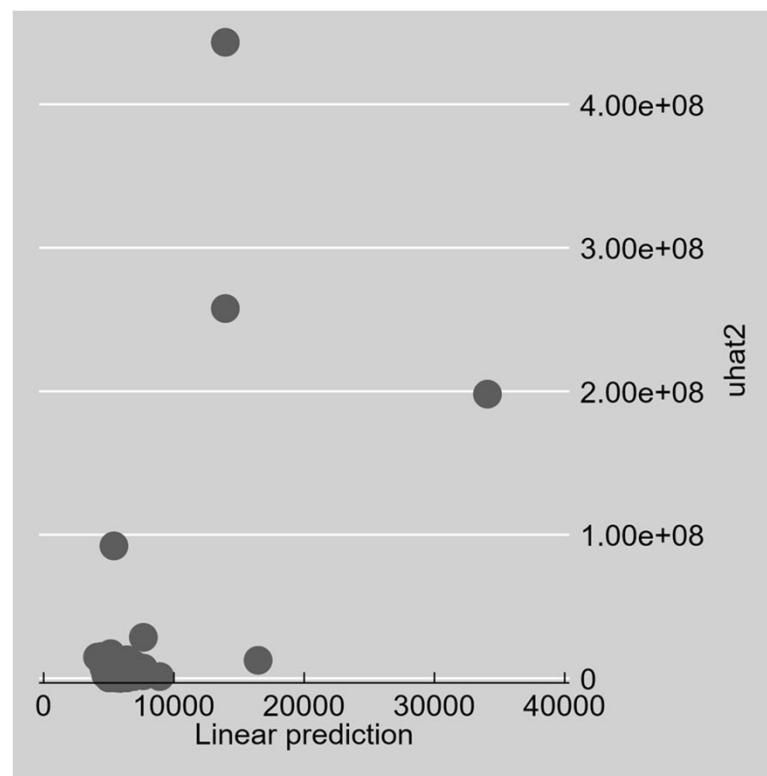
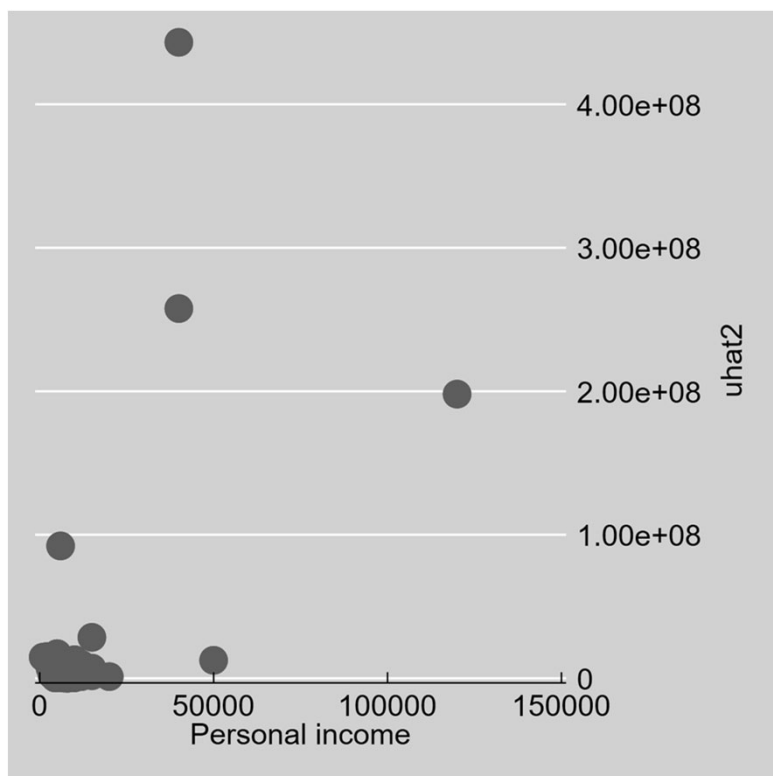
Source	SS	df	MS	Number of obs	=	51
				F(1, 49)	=	4.56
Model	15.7724964	1	15.7724964	Prob > F	=	0.0378
Residual	169.625547	49	3.46174585	R-squared	=	0.0851
				Adj R-squared	=	0.0664
Total	185.398043	50	3.70796086	Root MSE	=	1.8606

luhat2	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
lninc	.6966395	.3263664	2.13	0.038	.0407816	1.352497
_cons	8.848344	2.934702	3.02	0.004	2.950841	14.74585

6.3 Heteroscedasticity

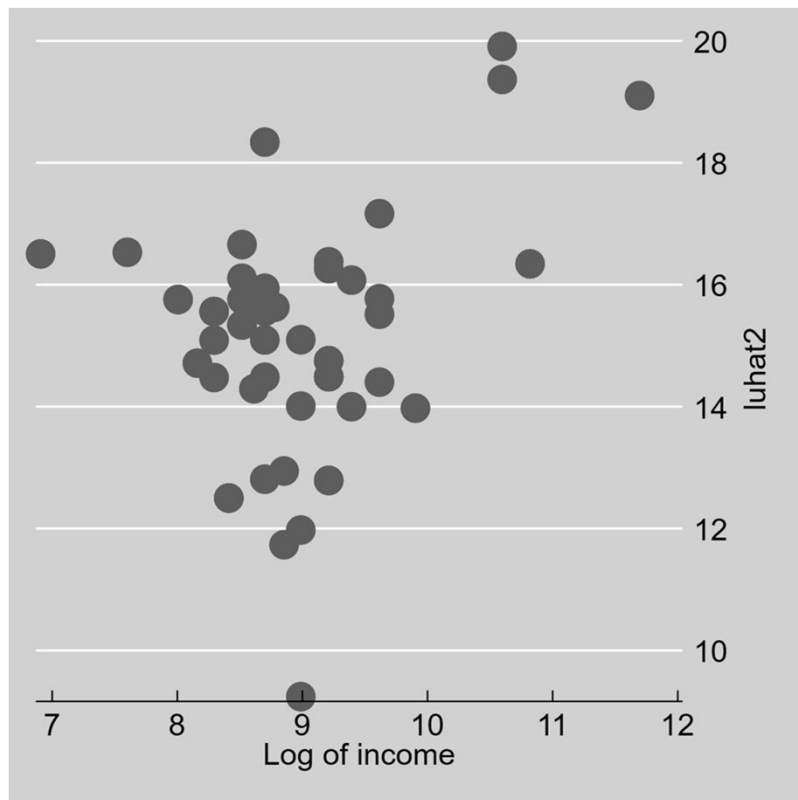
(3) Detecting heteroscedasticity

Graphing $\hat{u}_i^2$ with X_i and $\hat{Y}_i$, respectively, shows upward trend and they are exactly the same.



6.3 Heteroscedasticity

(3) Detecting heteroscedasticity



Another useful plot is $\ln \hat{u}_i^2$ versus $\ln X_i$. If heteroscedasticity is not present, the plots should scatter around zero.

(3) Detecting heteroscedasticity

3. Breusch-Pagan (BP) Test

This is a general case for other tests that follow Breusch and Pagan's idea (such as the Breusch-Pagan-Godfrey: BPG test in the textbook). This test is taken from Wooldridge page 270.

Step 1: Estimate coefficients with OLS and retrieve $\hat{u}_i^2$.

Step 2: Regress $\hat{u}_i^2 = \delta_1 + \delta_2 X_{2i} + \dots + \delta_k X_{ki} + v_i$ and retrieve $R_{\hat{u}_i^2}^2$.

Step 3: Calculate F-stat by

$$\circ F_{cal} = \frac{R_{\hat{u}_i^2}^2 / (k)}{(1 - R_{\hat{u}_i^2}^2) / (n - k - 1)}$$

Step 4: Test the null hypothesis of homoscedasticity. If we can reject the null hypothesis ($F_{cal} > F_{cri}$) at the selected significant level, heteroscedasticity is present in our model.

6.3 Heteroscedasticity

(3) Detecting heteroscedasticity

Example: Here is the result of BP-test ($\hat{u}_i^2 = \delta_1 + \delta_2 pinc_i + v_i$)

Source	SS	df	MS	Number of obs	=	51
Model	9.8138e+16	1	9.8138e+16	F(1, 49)	=	25.99
Residual	1.8505e+17	49	3.7765e+15	Prob > F	=	0.0000
				R-squared	=	0.3466
				Adj R-squared	=	0.3332
Total	2.8318e+17	50	5.6637e+15	Root MSE	=	6.1e+07

uhat2	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
pinc	2451.363	480.8728	5.10	0.000	1485.013	3417.713
_cons	-4798232	1.03e+07	-0.46	0.644	-2.56e+07	1.60e+07

Now let's try calculate F ,

$$\circ F_{cal} = \frac{R_{\hat{u}_i^2}^2 / (k)}{(1 - R_{\hat{u}_i^2}^2) / (n - k - 1)} =$$

Then, find the $F_{cri} =$

(3) *Detecting heteroscedasticity*

4. **White's test**

White's test is also very similar to Breusch-Pagan. The differences are the residual model and test statistics.

The steps here applies for 2 explanatory variables, but it is extendable. White's test has an advantage over BP test as it is not sensitive to the assumption of normality.

Step 1: Estimate coefficients with OLS and retrieve $\hat{u}_i^2$.

Step 2: Regress

$$\circ \hat{u}_i^2 = \delta_1 + \delta_2 X_{2i} + \delta_3 X_{3i} + \delta_4 X_{2i}^2 + \delta_5 X_{3i}^2 + \delta_6 X_{2i} X_{3i} + v_i$$

and retrieve $R_{\hat{u}_i^2}^2$.

Additions of higher power and cross product imply that the error variance is functionally related to regressors, their squares, and their cross product.

(3) *Detecting heteroscedasticity*

Step 3: Calculate the test stat, which in this case, we use **Lagrange Multiplier** (LM) stat

$$\circ LM_{cal} = n \cdot R_{\hat{u}_i^2}^2 \sim \chi_{k-1}^2.$$

LM is very useful in many cases due to its basic calculation and **asymptotically** distributed as Chi-square with k d.f.

Step 4: Test the null hypothesis of homoscedasticity. If we can reject the null hypothesis ($LM_{cal} > \chi_{k-1}^2$) at the selected significant level, heteroscedasticity is present in our model.

6.3 Heteroscedasticity

(3) Detecting heteroscedasticity

Example: Here is the result of White's test ($\hat{u}_i^2 = \delta_1 + \delta_2 \text{pinc}_i + \delta_2 \text{pinc}_i^2 + v_i$)

Source	SS	df	MS	Number of obs	=	51
				F(2, 48)	=	19.62
Model	1.2737e+17	2	6.3684e+16	Prob > F	=	0.0000
Residual	1.5582e+17	48	3.2462e+15	R-squared	=	0.4498
				Adj R-squared	=	0.4268
Total	2.8318e+17	50	5.6637e+15	Root MSE	=	5.7e+07

uhat2	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
pinc	6276.491	1350.466	4.65	0.000	3561.196	8991.785
pinc2	-.0358717	.0119545	-3.00	0.004	-.0599079	-.0118356
_cons	-3.37e+07	1.36e+07	-2.48	0.017	-6.10e+07	-6377575

Now let's try calculate LM stat

$$\circ LM_{cal} = n \cdot R_{\hat{u}_i^2}^2 =$$

Then, find the critical value of $\chi_{k-1.}^2 =$

(3) *Detecting heteroscedasticity*

Note that in STATA, the tests can be found in a package (post-estimation) and far simpler.

Breusch-Pagan/Cook-Weisberg test for heteroskedasticity

Assumption: Normal error terms

Variable: Fitted values of pexp

H0: Constant variance

```
chi2(1) = 83.06
Prob > chi2 = 0.0000
```

White's test

H0: Homoskedasticity

Ha: Unrestricted heteroskedasticity

```
chi2(2) = 22.94
Prob > chi2 = 0.0000
```

(4) Remedial measures

1. Weighted Least Squares (WLS)

It can be estimated when σ_i^2 is known.

2. Data transformation: selecting a class of GLS

Advantage of this method is that we do not need asymptotic property. However, a major drawback is we need to speculate relationship between σ_i^2 and X_i .

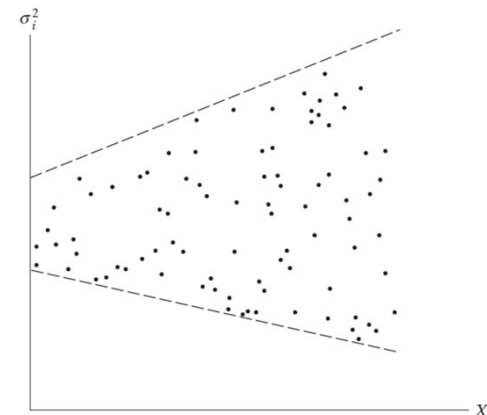
For example, if we assume that σ_i^2 is proportional to X_i so that

$$\circ E(u_i^2) = \sigma^2 X_i \text{ then } \sigma^2 = \frac{E(u_i^2)}{X_i} = \frac{E(u_i)}{\sqrt{X_i}}$$

we can transform our model into

$$\circ \frac{Y_i}{\sqrt{X_i}} = \frac{\beta_1}{\sqrt{X_i}} + \beta_2 \sqrt{X_i} + \frac{u_i}{\sqrt{X_i}} \text{ where } X_i \text{ must be } > 0$$

$$E\left(\frac{u_i}{\sqrt{X_i}}\right) = \sigma^2 \text{ or homoscedastic.}$$



(4) *Remedial measures*

However, we can see another drawback of this method is that the interpretation of coefficients are now different.

Moreover, there are multiple forms of transformation, due to relationship between σ_i^2 and X_i .

3. *White's robust standard errors*

This method assumes asymptotic property of our data, which is quite common in national cross-sectional data.

We do not cover how it is derived, even in the book Gujarati does not as well, but we compared the result of normal estimation with using White's robust standard errors.

6.3 Heteroscedasticity

(4) Remedial measures

○ Original model

Source	SS	df	MS	Number of obs	=	51
Model	1.0334e+09	1	1.0334e+09	F(1, 49)	=	40.85
Residual	1.2396e+09	49	25297763	Prob > F	=	0.0000
Total	2.2730e+09	50	45460031.4	R-squared	=	0.4546
				Adj R-squared	=	0.4435
				Root MSE	=	5029.7

pexp	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
pinc	.2515502	.0393576	6.39	0.000	.172458	.3306423
_cons	3886.008	845.2106	4.60	0.000	2187.493	5584.522

○ Original model with robust standard error

Linear regression				Number of obs	=	51
				F(1, 49)	=	5.59
				Prob > F	=	0.0221
				R-squared	=	0.4546
				Root MSE	=	5029.7

pexp	Coefficient	Robust std. err.	t	P> t	[95% conf. interval]	
pinc	.2515502	.1063896	2.36	0.022	.0377522	.4653481
_cons	3886.008	950.5687	4.09	0.000	1975.768	5796.247

(1) *Nature of the assumption*

Recall the assumption of no autocorrelation or serial correlation

$$\circ \text{cov}(u_i, u_j | x_i, x_j) = E(u_i u_j) = 0 \text{ where } i \neq j$$

Put simply, the term means “correlation between members of series of observations ordered in time (as in time series data) or space (as in cross-sectional data).” If the problem exists,

$$\circ E(u_i u_j) \neq 0 \text{ where } i \neq j$$

There can be multiple sources of autocorrelation.

1. Inertia

For example, when an economy is in recovery, there is a ‘momentum’ built into policies driving economic outcome from previous period.

(1) *Nature of the assumption*

2. *Lags*

For instance, a consumption-income model usually includes previous consumption expenditure

$$\circ \text{ } cons_t = \beta_1 + \beta_2 income_t + \beta_3 cons_{t-1} + u_i$$

This is because either psychologically, technologically, or institutionally, people do not change consumption pattern rapidly across periods.

3. *Cobweb phenomenon*

Most likely to occur with agricultural products, supply cannot adjust with price instantaneously.



(1) *Nature of the assumption*

4. Non-stationarity

Stationarity refers to time-invariant characteristics such as mean, variance, covariance, which is quite rare in time-series data. For example, GDP of an economy is always growing with non-systematic shocks that makes the characteristics time-variant.

5. *Specification bias: incorrect functional form*

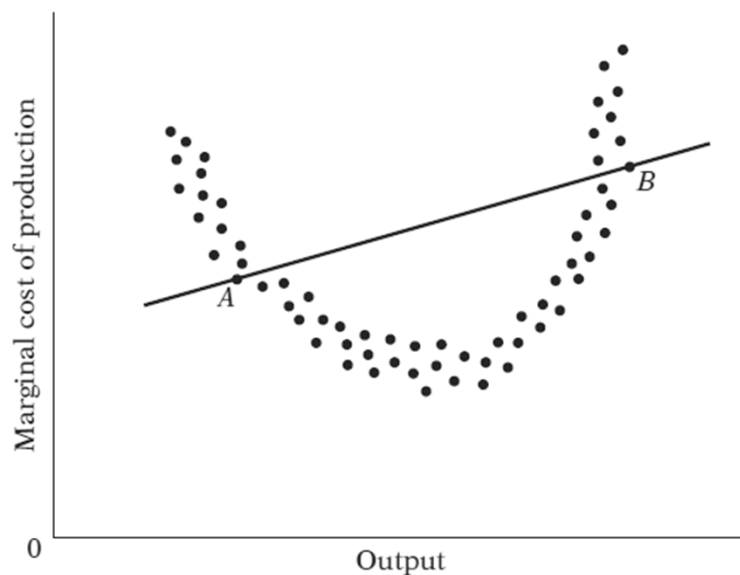
For example, a curved marginal cost which the ‘true’ model is

$$\circ MC_i = \beta_1 + \beta_2 Q_i + \beta_3 Q_i^2 + u_i$$

but if we try to fit our data with a linear model instead

$$\circ MC_i = \beta_1 + \beta_2 Q_i + u_i$$

(1) Nature of the assumption



Between point A and B, linear model underestimates marginal cost while before point A and beyond point B, the model overestimates marginal cost.

If we correlate u_i, u_j we can see clearly that there is a pattern following the curve, hence autocorrelation is present in the linear model.

However, this problem does not surface when we fit the data with polynomial model.

(2) *Effects on estimation*

From this point on, we will focus on a time-series model, varying Y_t and X_t through time, not across groups of observation in the same period, the model becomes

- $Y_t = \beta_1 + \beta_2 X_t + u_t$

If the error terms are correlated, assumed linearly

- $u_t = \rho u_{t-1} + \varepsilon_t$

This equation is called **first-order autoregressive scheme** or shortly **AR(1)**. If the error term t also correlates with two period back, the model is called **AR(2)**, and so on.

- $u_t = \rho_1 u_{t-1} + \rho_2 u_{t-2} + \varepsilon_t$

(2) *Effects on estimation*

Normally, OLS with no autocorrelation will yield the variance of an estimator as

$$\circ \text{var}(\hat{\beta}_2) = \frac{\sigma^2}{\sum x_t^2}$$

If u_t follows AR(1), the variance becomes

$$\circ \text{var}(\hat{\beta}_2)_{AR(1)} = \frac{\sigma^2}{\sum x_t^2} \left[1 + 2\rho \frac{\sum x_t x_{t-1}}{\sum x_t^2} + 2\rho^2 \frac{\sum x_t x_{t-2}}{\sum x_t^2} + \dots + 2\rho^{n-1} \frac{\sum x_t x_n}{\sum x_t^2} \right]$$

We cannot say for sure whether $\text{var}(\hat{\beta}_2)$ is more or less than $\text{var}(\hat{\beta}_2)_{AR(1)}$.

Though $\hat{\beta}_2$ is still linear and unbiased, variance is not minimum, or not being efficient. See this proof of GLS on page 422.

(2) *Effects on estimation*

If we run the regression, disregarding autocorrelation, we will find that

- $\hat{\sigma}^2 = \frac{\sum \hat{u}_t^2}{(n-2)}$ is likely to underestimate the true σ^2 .
- Overestimation of R^2 .
- If σ^2 is not underestimated, $\text{var}(\hat{\beta}_2)$ may still underestimate $\text{var}(\hat{\beta}_2)_{AR(1)}$, and the latter is still inefficient compared to $\text{var}(\hat{\beta}_2)_{GLS}$.
- Both t and F tests are no longer valid.

(3) *Detecting autocorrelation*

1. *Graphical method*

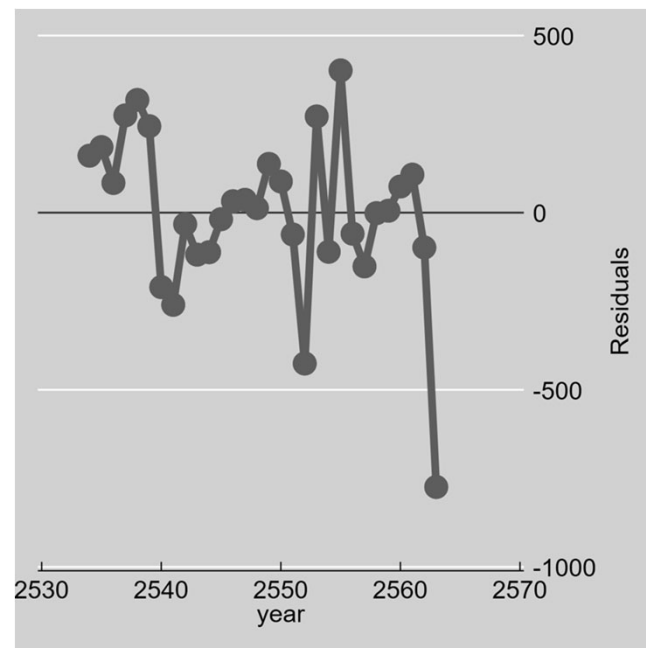
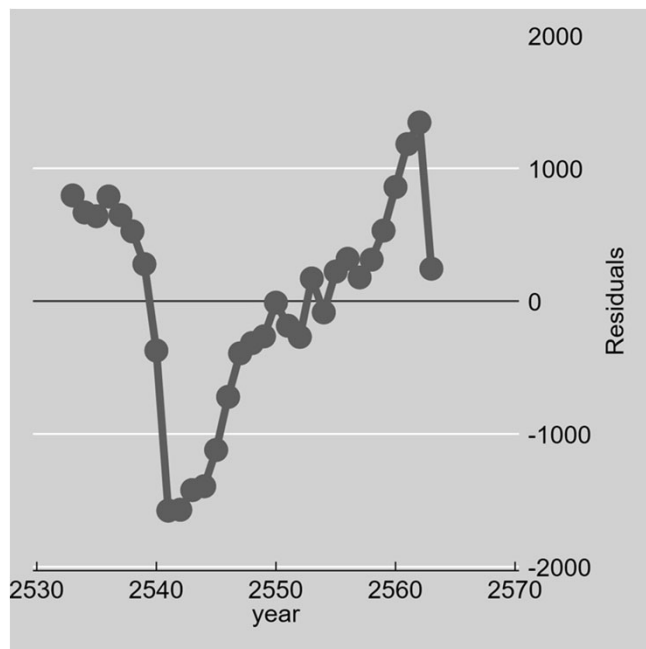
The first one, once again, is not informal but visually telling. If we plot the residuals with time period, we might see interconnection between time period.

Example: Data of GDP (Y_t), in CVM, and core CPI (X_t) in Thailand from 1990 to 2020 are taken from the Bank of Thailand statistics page. We model it as usual.

$$\circ Y_t = \beta_1 + \beta_2 X_t + u_t$$

Now see the result of the estimation on the right-hand side. After that, we can predict $\hat{u}_t$ and plot them over time.

(3) *Detecting autocorrelation*



On the left-hand side, this is a plot from the model residual. On the other hand, the right-hand side shows a plot from another model that autocorrelation is already resolved.

If autocorrelation is not present, residuals should be randomly distributed around 0 and has no obvious interconnection with the previous period.

6.4 Autocorrelation

(3) Detecting autocorrelation

Source	SS	df	MS	Number of obs	=	31
				F(1, 29)	=	210.05
Model	135778013	1	135778013	Prob > F	=	0.0000
Residual	18745422.4	29	646393.875	R-squared	=	0.8787
				Adj R-squared	=	0.8745
Total	154523435	30	5150781.18	Root MSE	=	803.99

cvm	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
core	152.3648	10.51281	14.49	0.000	130.8637	173.8659
_cons	-5649.429	884.7684	-6.39	0.000	-7458.983	-3839.874

We can also see that R^2 , t and F are very high, leading to a very significant coefficient, but this is due autocorrelation.

(3) *Detecting autocorrelation*

2. **Durbin-Watson d Test**

Durbin-Watson d statistics is defined as

$$\circ \quad d = \frac{\sum_{t=2}^n (\hat{u}_t - \hat{u}_{t-1})^2}{\sum_{t=1}^n \hat{u}_t^2}$$

implies sum of squared differences to the RSS.

For Durbin-Watson test, we assume

- Regression model includes intercept term.
- Regressors X are stochastic or fixed in repeated sampling.
- u_t are generated by AR(1) and normally distributed.
- The model does not include lagged variable(s) of the dependent variable, such as

$$Y_t = \beta_1 + \beta_2 X_t + \gamma Y_{t-1} + u_t$$

- No missing observations in the data.

(3) Detecting autocorrelation

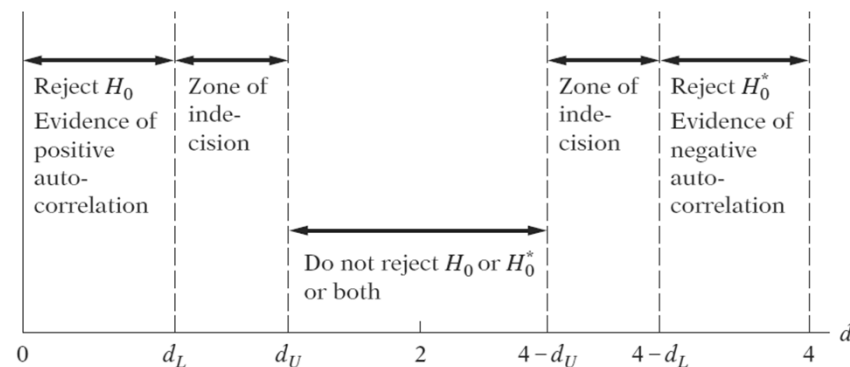
We can also define 'approximate' version of the d stat as

○ $d \approx 2\left(1 - \frac{\sum \hat{u}_t \hat{u}_{t-1}}{\sum \hat{u}_t^2}\right)$ now let's define

○ $\hat{\rho} = \frac{\sum \hat{u}_t \hat{u}_{t-1}}{\sum \hat{u}_t^2}$ then

○ $d \approx 2(1 - \hat{\rho})$

Since $-1 \leq \hat{\rho} \leq 1$, therefore $0 \leq d \leq 4$



Legend

H_0 : No positive autocorrelation

H_0^* : No negative autocorrelation

(3) *Detecting autocorrelation*

Example: using the same dataset, we follow the steps here.

Step 1: State the hypothesis

Since we are testing for general autocorrelation, without any speculation of positive or negative correlation, we are going to set hypothesis for both.

- H_0 : No autocorrelation ($d_U < d < 4 - d_U$)
- H_a : Positive autocorrelation ($0 < d < d_L$) or
negative autocorrelation ($4 - d_L < d < 4$)

Do not forget that we can also reach the inconclusive answer for this test.

Step 2: Run the OLS regression and obtain the predicted residuals.

○ $d_L =$ and $4 - d_L =$

$$\circ \quad d_{II} = \quad \text{and } 4 - d_{II} =$$

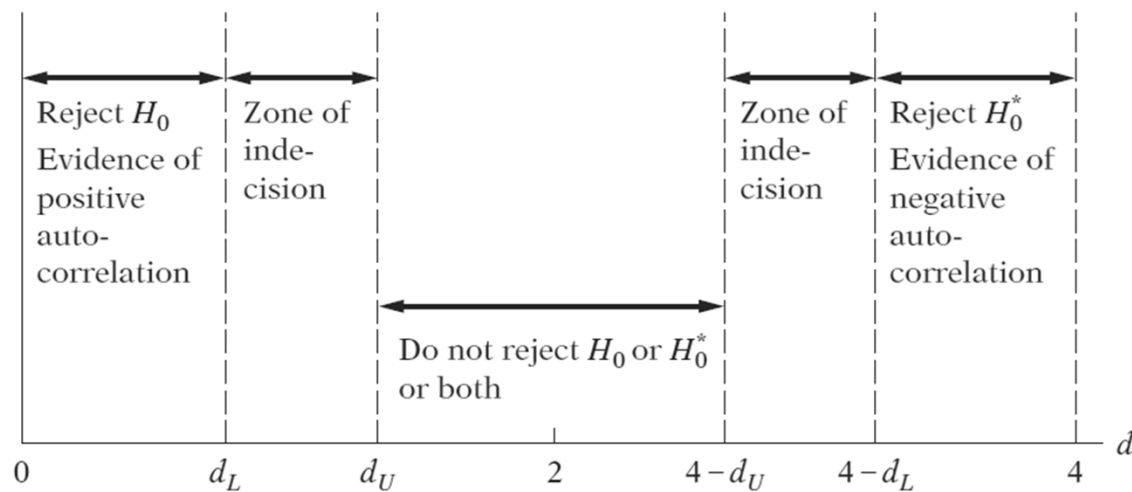
Source	SS	df	MS	Number of obs	=	31
				F(1, 29)	=	210.05
Model	135778013	1	135778013	Prob > F	=	0.0000
Residual	18745422.4	29	646393.875	R-squared	=	0.8787
				Adj R-squared	=	0.8745
Total	154523435	30	5150781.18	Root MSE	=	803.99

cvm	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
core	152.3648	10.51281	14.49	0.000	130.8637	173.8659
_cons	-5649.429	884.7684	-6.39	0.000	-7458.983	-3839.874

```
Durbin-Watson d-statistic( 2, 31) = .2338122
```

(3) Detecting autocorrelation

Step 4: Conclude the test.



Legend

H_0 : No positive autocorrelation

H_0^* : No negative autocorrelation

(3) *Detecting autocorrelation*

3. General test of autocorrelation: The Breusch-Godfrey test (BG)

Assume a model of

$$\circ Y_t = \beta_1 + \beta_2 X_t + u_t$$

Now if the error term u_t follows the p^{th} -order autoregressive AR(P) as follows

$$\circ u_t = \rho_1 u_{t-1} + \rho_2 u_{t-2} + \cdots + \rho_p u_{t-p} + \varepsilon_t$$

The concept is to test these coefficients simultaneously by following the steps.

Step 1: State the hypothesis

- $\circ H_0$: No autocorrelation ($\rho_1 = \rho_2 = \cdots = \rho_p = 0$)
- $\circ H_a$: Autocorrelation (ρ are not simultaneously zero)

Step 2: Run the OLS regression and obtain the residuals.

(3) *Detecting autocorrelation*

Step 3: Estimate this equation, also including regressor(s) into this one.

$$\circ \hat{u}_t = \alpha_1 + \alpha_2 X_t + \hat{\rho}_1 \hat{u}_{t-1} + \hat{\rho}_2 \hat{u}_{t-2} + \cdots + \hat{\rho}_p \hat{u}_{t-p} + \varepsilon_t$$

Then obtain the R^2 from this model

Step 4: Calculate LM-statistics, if the sample is large, BG have shown that

$$\circ LM = (n - p)R^2 \sim \chi_p^2$$

Step 5: Look for the critical value in chi-square table and reject the null hypothesis if the LM exceeds the critical value.

Note that, in STATA, the default lag is 1 but there is an option to include more lags.

6.4 Autocorrelation

(3) Detecting autocorrelation

Example: reconsider the model of GDP (Y_t), in CVM, and core CPI (X_t) once again.

Source	SS	df	MS	Number of obs	=	31
Model	135778013	1	135778013	F(1, 29)	=	210.05
Residual	18745422.4	29	646393.875	Prob > F	=	0.0000
Total	154523435	30	5150781.18	R-squared	=	0.8787
				Adj R-squared	=	0.8745
				Root MSE	=	803.99

cvm	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
core	152.3648	10.51281	14.49	0.000	130.8637	173.8659
_cons	-5649.429	884.7684	-6.39	0.000	-7458.983	-3839.874

Breusch-Godfrey LM test for autocorrelation

lags(p)	chi2	df	Prob > chi2
1	23.257	1	0.0000

H0: no serial correlation

(4) Remedial measures

First of all, make sure that the model is correctly specified (we are going to deal with autocorrelation only). Most of the time we do not know the relationship between u_t and u_{t-1} . In other words, we do not know the value of ρ .

1. First-difference method

The first difference equation takes the form of

- $Y_t - Y_{t-1} = \beta_2(X_t - X_{t-1}) + (u_t - u_{t-1})$ or
- $\Delta Y_t = \beta_2 \Delta X_t + \varepsilon_t$ where $\varepsilon_t = \Delta u_t$

The rule of thumb is that we can use this equation to estimate when $d < R^2$. Note that this model has no intercept term. If included, the intercept is interpreted as **time trend**.

6.4 Autocorrelation

(4) Remedial measures

As we can see from the result, when we use the first-difference model, we cannot reject the null hypothesis of BG test any longer because ε_t is a stationary white-noise error term. However, note that we lose 1 observation due to using difference term.

Source	SS	df	MS	Number of obs	=	30
Model	442238.296	1	442238.296	F(1, 28)	=	7.95
Residual	1557421.97	28	55622.2132	Prob > F	=	0.0087
				R-squared	=	0.2212
				Adj R-squared	=	0.1933
Total	1999660.27	29	68953.8023	Root MSE	=	235.84

D.cvm	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
core						
D1.	-100.0344	35.47687	-2.82	0.009	-172.7054	-27.36332
_cons	392.616	72.05022	5.45	0.000	245.0278	540.2042

Breusch-Godfrey LM test for autocorrelation

lags(p)	chi2	df	Prob > chi2
1	0.608	1	0.4356

H0: no serial correlation

(4) Remedial measures

Now we can try excluding the constant term.

Source	SS	df	MS	Number of obs	=	30
				F(1, 29)	=	3.38
Model	373835.398	1	373835.398	Prob > F	=	0.0763
Residual	3209055.02	29	110657.07	R-squared	=	0.1043
				Adj R-squared	=	0.0735
Total	3582890.42	30	119429.681	Root MSE	=	332.65

D.cvm	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
core						
D1.	54.96542	29.90466	1.84	0.076	-6.196484	116.1273

(4) Remedial measures

2. Estimating ρ

The reason why we want to know the value of ρ is that we can use this value to transform variables, then use the transformed ones in the GLS estimation.

Assume that the error term follows AR(1) scheme,

$$\circ u_t = \rho u_{t-1} + \varepsilon_t \text{ where } 1 < \rho < 1$$

we transform $t - 1$ equation by multiply ρ

$$\circ \rho Y_{t-1} = \rho \beta_1 + \rho \beta_2 X_{t-1} + \rho u_{t-1}$$

Then, subtract the t with the equation above to remove the coexistence of the element that are the same between time

$$\circ (Y_t - \rho Y_{t-1}) = \beta_1(1 - \rho) + \beta_2(X_t - \rho X_{t-1}) + \varepsilon_t \text{ or}$$

$$\circ Y_t^* = \beta_1^* + \beta_2^* X_t^* + \varepsilon_t$$

There are several methods that we can estimate this value of ρ .

(4) Remedial measures

2.1 ρ based on Durbin-Watson d statistics.

From the Durbin-Watson test, we can derive that

$$\circ \hat{\rho} \approx 1 - \frac{d}{2}$$

2.2 ρ estimated from the residuals

We can also estimate another model postestimation,

$$\circ \hat{u}_t = \rho \hat{u}_{t-1} + v_t$$

2.3 Cochrane-Orcutt iterative procedure

Iterative method estimates ρ by starting at some value, mostly 0, then successively approximate multiple times until the value of ρ is stable. Then, ρ can be put into the transformation.

(4) Remedial measures

2.4 Prais-Winsten transformation

Using the same concept of iterative ρ , Prais-Winsten fixed losing 1 observation from the first-difference method because of no antecedent by adding

- $Y_1\sqrt{1-\rho^2}$ and $X_1\sqrt{1-\rho^2}$

Note that Prais-Winsten transformation will retain original number of observation, which might be very important especially for a small sample dataset.

Compare the results from the original model to the third and the fourth method in the next page.

6.4 Autocorrelation

(4) Remedial measures

Original model

Source	SS	df	MS	Number of obs	=	31
Model	135778013	1	135778013	F(1, 29)	=	210.05
Residual	18745422.4	29	646393.875	Prob > F	=	0.0000
				R-squared	=	0.8787
				Adj R-squared	=	0.8745
Total	154523435	30	5150781.18	Root MSE	=	803.99

cvm	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
core	152.3648	10.51281	14.49	0.000	130.8637	173.8659
_cons	-5649.429	884.7684	-6.39	0.000	-7458.983	-3839.874

Cochrane-Orcutt iterative procedure

Source	SS	df	MS	Number of obs	=	30
Model	677268.662	1	677268.662	F(1, 28)	=	14.58
Residual	1300323.16	28	46440.1127	Prob > F	=	0.0007
				R-squared	=	0.3425
				Adj R-squared	=	0.3190
Total	1977591.82	29	68192.8213	Root MSE	=	215.5

cvm	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
core	-139.9205	36.63932	-3.82	0.001	-214.9727	-64.86822
_cons	35910.77	5508.175	6.52	0.000	24627.78	47193.75
rho	.9738418					

Durbin-Watson statistic (original) = 0.233812

Durbin-Watson statistic (transformed) = 1.588503

6.4 Autocorrelation

(4) Remedial measures

Original model

Source	SS	df	MS	Number of obs	=	31
Model	135778013	1	135778013	F(1, 29)	=	210.05
Residual	18745422.4	29	646393.875	Prob > F	=	0.0000
				R-squared	=	0.8787
				Adj R-squared	=	0.8745
Total	154523435	30	5150781.18	Root MSE	=	803.99

cvm	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
core	152.3648	10.51281	14.49	0.000	130.8637	173.8659
_cons	-5649.429	884.7684	-6.39	0.000	-7458.983	-3839.874

Prais-Winsten transformation

Source	SS	df	MS	Number of obs	=	31
Model	0	1	0	F(1, 29)	=	0.00
Residual	3356405.19	29	115738.11	Prob > F	=	1.0000
				R-squared	=	.
				Adj R-squared	=	.
Total	2124497.99	30	70816.5998	Root MSE	=	340.2

cvm	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
core	76.13285	26.97612	2.82	0.009	20.96049	131.3052
_cons	795.4023	2459.377	0.32	0.749	-4234.588	5825.392
rho	.9710003					

Durbin-Watson statistic (original) = 0.233812

Durbin-Watson statistic (transformed) = 1.167809

(4) Remedial measures

3. Newey-West method

This method does not deal with autocorrelation directly, instead, it is very much like White’s robust standard error.

The corrected standard errors are known as **HAC standard error**. (heteroscedasticity and autocorrelation).

Newey-West approach is strictly speaking valid in large samples.

Regression with Newey-West standard errors
Maximum lag = 1

Number of obs = 31
F(1, 29) = 140.34
Prob > F = 0.0000

cvm	Coefficient	Newey-West std. err.	t	P> t	[95% conf. interval]	
core	152.3648	12.86162	11.85	0.000	126.0599	178.6698
_cons	-5649.429	1101.437	-5.13	0.000	-7902.121	-3396.737