

Week 6: Linear Regression with Two Regressors

Brandon Stewart¹

Princeton

October 10–12, 2022

¹Many of these slides are adapted from Matt Blackwell, Adam Glynn, Jens Hainmueller and Erin Hartman.

Where We've Been and Where We're Going...

Where We've Been and Where We're Going...

- Last Week
 - ▶ mechanics of OLS with one variable
 - ▶ properties of OLS

Where We've Been and Where We're Going...

- Last Week

- ▶ mechanics of OLS with one variable
- ▶ properties of OLS

- This Week

- ▶ adding a second variable
- ▶ new considerations for estimation and inference
- ▶ interactions and non-linearity

Where We've Been and Where We're Going...

- Last Week
 - ▶ mechanics of OLS with one variable
 - ▶ properties of OLS
- This Week
 - ▶ adding a second variable
 - ▶ new considerations for estimation and inference
 - ▶ interactions and non-linearity
- Next Week
 - ▶ multiple regression
- Long Run
 - ▶ probability $\rightarrow$ inference $\rightarrow$ regression $\rightarrow$ causal inference

A Brief Review

A Brief Review

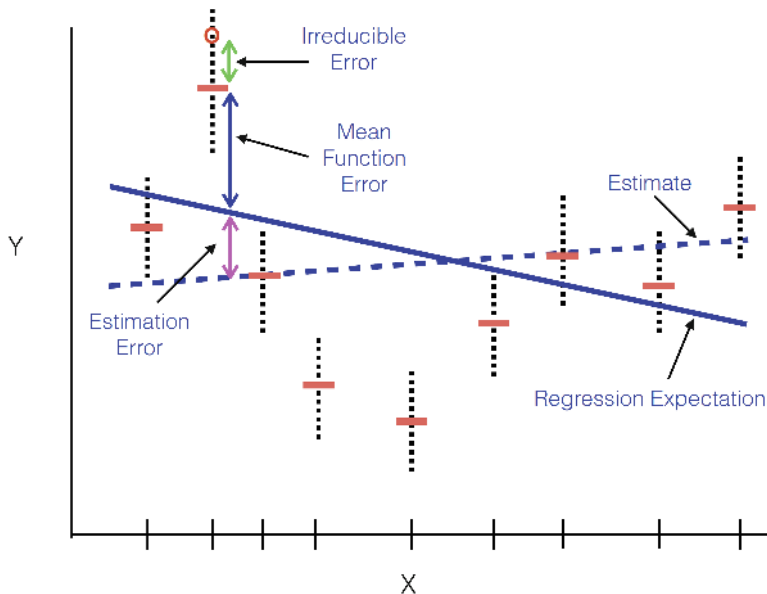


Image Credit: Berk (2020) *Statistical learning from a regression perspective*

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

Why Do We Want More Than One Predictor?

Why Do We Want More Than One Predictor?

- Summarize more information for descriptive inference

Why Do We Want More Than One Predictor?

- Summarize more information for descriptive inference
- Improve the fit and predictive power of our model

Why Do We Want More Than One Predictor?

- Summarize more information for descriptive inference
- Improve the fit and predictive power of our model
- Control for confounding factors for causal inference

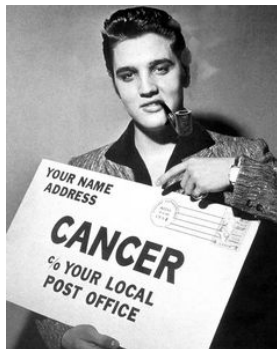
Why Do We Want More Than One Predictor?

- Summarize more information for descriptive inference
- Improve the fit and predictive power of our model
- Control for confounding factors for causal inference
- Model non-linearities (e.g. $Y = \beta_0 + \beta_1 X + \beta_2 X^2$)

Why Do We Want More Than One Predictor?

- Summarize more information for descriptive inference
- Improve the fit and predictive power of our model
- Control for confounding factors for causal inference
- Model non-linearities (e.g. $Y = \beta_0 + \beta_1 X + \beta_2 X^2$)
- Model interactive effects (e.g. $Y = \beta_0 + \beta_1 X + \beta_2 X_2 + \beta_3 X_1 X_2$)

Example: Cigarette Smokers and Pipe Smokers



Example: Cigarette Smokers and Pipe Smokers

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.
- X_1 : 0 if person is pipe smoker; 1 if person is cigarette smoker

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.
- X_1 : 0 if person is pipe smoker; 1 if person is cigarette smoker

We fit the regression and find:

$$\widehat{\text{Death Rate}} = 17 - 4 \text{ Cigarette Smoker}$$

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.
- X_1 : 0 if person is pipe smoker; 1 if person is cigarette smoker

We fit the regression and find:

$$\widehat{\text{Death Rate}} = 17 - 4 \text{ Cigarette Smoker}$$

What do we conclude?

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.
- X_1 : 0 if person is pipe smoker; 1 if person is cigarette smoker

We fit the regression and find:

$$\widehat{\text{Death Rate}} = 17 - 4 \text{ Cigarette Smoker}$$

What do we conclude?

- The average death rate is 17 deaths per 1,000 person-years for pipe smokers and 13 ($17 - 4$) for cigarette smokers.

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.
- X_1 : 0 if person is pipe smoker; 1 if person is cigarette smoker

We fit the regression and find:

$$\widehat{\text{Death Rate}} = 17 - 4 \text{ Cigarette Smoker}$$

What do we conclude?

- The average death rate is 17 deaths per 1,000 person-years for pipe smokers and 13 ($17 - 4$) for cigarette smokers.
- So cigarette smoking appears to lower the death rate by 4 deaths per 1,000 person years (relative to pipe smoking).

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.
- X_1 : 0 if person is pipe smoker; 1 if person is cigarette smoker

We fit the regression and find:

$$\widehat{\text{Death Rate}} = 17 - 4 \text{ Cigarette Smoker}$$

What do we conclude?

- The average death rate is 17 deaths per 1,000 person-years for pipe smokers and 13 ($17 - 4$) for cigarette smokers.
- So cigarette smoking appears to lower the death rate by 4 deaths per 1,000 person years (relative to pipe smoking).

When we “control” for age (in years) we find:

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.
- X_1 : 0 if person is pipe smoker; 1 if person is cigarette smoker

We fit the regression and find:

$$\widehat{\text{Death Rate}} = 17 - 4 \text{ Cigarette Smoker}$$

What do we conclude?

- The average death rate is 17 deaths per 1,000 person-years for pipe smokers and 13 ($17 - 4$) for cigarette smokers.
- So cigarette smoking appears to lower the death rate by 4 deaths per 1,000 person years (relative to pipe smoking).

When we “control” for age (in years) we find:

$$\widehat{\text{Death Rate}} = 14 + 4 \text{ Cigarette Smoker} + 10 \text{ Age}$$

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.
- X_1 : 0 if person is pipe smoker; 1 if person is cigarette smoker

We fit the regression and find:

$$\widehat{\text{Death Rate}} = 17 - 4 \text{ Cigarette Smoker}$$

What do we conclude?

- The average death rate is 17 deaths per 1,000 person-years for pipe smokers and 13 ($17 - 4$) for cigarette smokers.
- So cigarette smoking appears to lower the death rate by 4 deaths per 1,000 person years (relative to pipe smoking).

When we “control” for age (in years) we find:

$$\widehat{\text{Death Rate}} = 14 + 4 \text{ Cigarette Smoker} + 10 \text{ Age}$$

Why did the sign switch?

Example: Cigarette Smokers and Pipe Smokers

Consider the following example from Cochran (1968). We have a random sample of 20,000 smokers and run a regression using:

- Y : Deaths per 1,000 Person-Years.
- X_1 : 0 if person is pipe smoker; 1 if person is cigarette smoker

We fit the regression and find:

$$\widehat{\text{Death Rate}} = 17 - 4 \text{ Cigarette Smoker}$$

What do we conclude?

- The average death rate is 17 deaths per 1,000 person-years for pipe smokers and 13 ($17 - 4$) for cigarette smokers.
- So cigarette smoking appears to lower the death rate by 4 deaths per 1,000 person years (relative to pipe smoking).

When we “control” for age (in years) we find:

$$\widehat{\text{Death Rate}} = 14 + 4 \text{ Cigarette Smoker} + 10 \text{ Age}$$

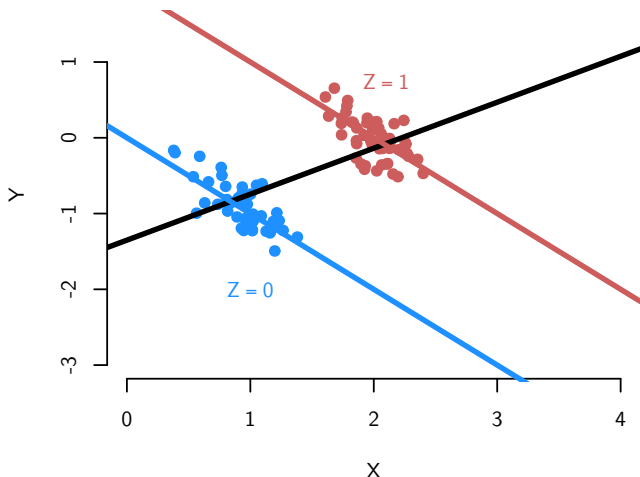
Why did the sign switch? Which estimate is more useful?

Simpson's Paradox

The smoking pattern is an instance of **Simpson's Paradox**.

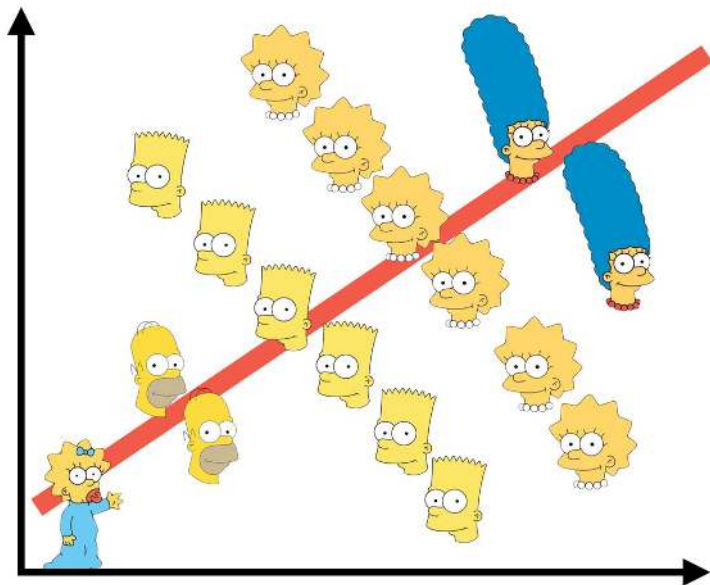
Simpson's Paradox

The smoking pattern is an instance of **Simpson's Paradox**.



Core idea: a relationship in one direction between Y_i and X_i but the opposite relationship within strata defined by Z_i .

Simpson's Paradox



Paradoxes Are Rarely Really Paradoxes

Paradoxes Are Rarely Really Paradoxes

- The paradox here is really just a case of **imprecision of language**.

Paradoxes Are Rarely Really Paradoxes

- The paradox here is really just a case of **imprecision of language**.
- We observe the statement

$$P(\text{Death}|\text{Cigarette Smoking}) < P(\text{Death}|\text{Pipe Smoking})$$

and translate it to 'cigarette smoking **lowers** the death rate.'

Paradoxes Are Rarely Really Paradoxes

- The paradox here is really just a case of **imprecision of language**.
- We observe the statement

$$P(\text{Death}|\text{Cigarette Smoking}) < P(\text{Death}|\text{Pipe Smoking})$$

and translate it to 'cigarette smoking **lowers** the death rate.'

- This is in tension with then the subsequent finding

$$P(\text{Death}|\text{Cigarette Smoking, Age}) > P(\text{Death}|\text{Pipe Smoking, Age})$$

which we translate to 'cigarette smoking **increases** the death rate for each age group.'

Paradoxes Are Rarely Really Paradoxes

- The paradox here is really just a case of **imprecision of language**.
- We observe the statement

$$P(\text{Death}|\text{Cigarette Smoking}) < P(\text{Death}|\text{Pipe Smoking})$$

and translate it to 'cigarette smoking **lowers** the death rate.'

- This is in tension with then the subsequent finding

$$P(\text{Death}|\text{Cigarette Smoking, Age}) > P(\text{Death}|\text{Pipe Smoking, Age})$$

which we translate to 'cigarette smoking **increases** the death rate for each age group.'

- Both text translations cannot be true, but the math **does not imply** the causal interpretation given in the text.

Paradoxes Are Rarely Really Paradoxes

- The paradox here is really just a case of **imprecision of language**.
- We observe the statement

$$P(\text{Death}|\text{Cigarette Smoking}) < P(\text{Death}|\text{Pipe Smoking})$$

and translate it to 'cigarette smoking **lowers** the death rate.'

- This is in tension with then the subsequent finding

$$P(\text{Death}|\text{Cigarette Smoking, Age}) > P(\text{Death}|\text{Pipe Smoking, Age})$$

which we translate to 'cigarette smoking **increases** the death rate for each age group.'

- Both text translations cannot be true, but the math **does not imply** the causal interpretation given in the text.
- Conditioning is just a way of looking at subgroups—we will see later that this plays a key role in making **causal inferences** but it requires careful assumptions.

Simpson's Paradox

Simpson's Paradox

- Simpson's paradox arises in many contexts- particularly where there is **selection**.

Simpson's Paradox

- Simpson's paradox arises in many contexts- particularly where there is **selection**.
- It is a particular problem in medical or demographic contexts, e.g. kidney stones, low-birth weight paradox.

Simpson's Paradox

- Simpson's paradox arises in many contexts- particularly where there is **selection**.
- It is a particular problem in medical or demographic contexts, e.g. kidney stones, low-birth weight paradox.
- It isn't clear that one version (the marginal or conditional) is necessarily the **right** way to examine the data. They just have different **meanings**.

Simpson's Paradox

- Simpson's paradox arises in many contexts- particularly where there is **selection**.
- It is a particular problem in medical or demographic contexts, e.g. kidney stones, low-birth weight paradox.
- It isn't clear that one version (the marginal or conditional) is necessarily the **right** way to examine the data. They just have different **meanings**.
- This is often an issue of not being clear what we want.

Simpson's Paradox

- Simpson's paradox arises in many contexts- particularly where there is **selection**.
- It is a particular problem in medical or demographic contexts, e.g. kidney stones, low-birth weight paradox.
- It isn't clear that one version (the marginal or conditional) is necessarily the **right** way to examine the data. They just have different **meanings**.
- This is often an issue of not being clear what we want.

Instance of a more general problem called the **ecological inference fallacy**.

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

Basic Idea of Two Variable Regressions

- Old goal: estimate the mean of Y as a function of one independent variable, X :

$$E[Y|X]$$

Basic Idea of Two Variable Regressions

- Old goal: estimate the mean of Y as a function of one independent variable, X :

$$E[Y|X]$$

- We modeled the CEF/regression function with a line:

$$Y_i = \beta_0 + \beta_1 X_i + u_i$$

Basic Idea of Two Variable Regressions

- Old goal: estimate the mean of Y as a function of one independent variable, X :

$$E[Y|X]$$

- We modeled the CEF/regression function with a line:

$$Y_i = \beta_0 + \beta_1 X_i + u_i$$

- New goal: estimate the relationship of two variables, Y_i and X_i , conditional on a third variable, Z_i :

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

Basic Idea of Two Variable Regressions

- Old goal: estimate the mean of Y as a function of one independent variable, X :

$$E[Y|X]$$

- We modeled the CEF/regression function with a line:

$$Y_i = \beta_0 + \beta_1 X_i + u_i$$

- New goal: estimate the relationship of two variables, Y_i and X_i , conditional on a third variable, Z_i :

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

- β 's are the population parameters we want to estimate

Regression with Two Explanatory Variables

Example: data from Fish (2002) “Islam and Authoritarianism.” *World Politics*. 55: 4-37. Data from 157 countries.

Regression with Two Explanatory Variables

Example: data from Fish (2002) “Islam and Authoritarianism.” *World Politics*. 55: 4-37. Data from 157 countries.

- Variables of interest:
 - ▶ Y: Level of democracy, measured as the 10-year average of Freedom House ratings

Regression with Two Explanatory Variables

Example: data from Fish (2002) “Islam and Authoritarianism.” *World Politics*. 55: 4-37. Data from 157 countries.

- Variables of interest:
 - ▶ Y : Level of democracy, measured as the 10-year average of Freedom House ratings
 - ▶ X_1 : Country income, measured as $\log(\text{GDP per capita in \$1000s})$

Regression with Two Explanatory Variables

Example: data from Fish (2002) “Islam and Authoritarianism.” *World Politics*. 55: 4-37. Data from 157 countries.

- Variables of interest:
 - ▶ Y : Level of democracy, measured as the 10-year average of Freedom House ratings
 - ▶ X_1 : Country income, measured as $\log(\text{GDP per capita in \$1000s})$
 - ▶ X_2 : Ethnic heterogeneity (continuous) or British colonial heritage (binary)

Regression with Two Explanatory Variables

Example: data from Fish (2002) “Islam and Authoritarianism.” *World Politics*. 55: 4-37. Data from 157 countries.

- Variables of interest:
 - ▶ Y : Level of democracy, measured as the 10-year average of Freedom House ratings
 - ▶ X_1 : Country income, measured as $\log(\text{GDP per capita in \$1000s})$
 - ▶ X_2 : Ethnic heterogeneity (continuous) or British colonial heritage (binary)
- With one predictor we ask: Does income (X_1) predict the level of democracy (Y)?

Regression with Two Explanatory Variables

Example: data from Fish (2002) “Islam and Authoritarianism.” *World Politics*. 55: 4-37. Data from 157 countries.

- Variables of interest:
 - ▶ Y : Level of democracy, measured as the 10-year average of Freedom House ratings
 - ▶ X_1 : Country income, measured as $\log(\text{GDP per capita in \$1000s})$
 - ▶ X_2 : Ethnic heterogeneity (continuous) or British colonial heritage (binary)
- With one predictor we ask: Does income (X_1) predict the level of democracy (Y)?

Regression with Two Explanatory Variables

Example: data from Fish (2002) “Islam and Authoritarianism.” *World Politics*. 55: 4-37. Data from 157 countries.

- Variables of interest:
 - ▶ Y : Level of democracy, measured as the 10-year average of Freedom House ratings
 - ▶ X_1 : Country income, measured as $\log(\text{GDP per capita in \$1000s})$
 - ▶ X_2 : Ethnic heterogeneity (continuous) or British colonial heritage (binary)
- With one predictor we ask: Does income (X_1) predict the level of democracy (Y)?
- With two predictors we ask questions like: Does income (X_1) predict the level of democracy (Y), once we “control” for ethnic heterogeneity or British colonial heritage (X_2)?

Regression with Two Explanatory Variables

Example: data from Fish (2002) “Islam and Authoritarianism.” *World Politics*. 55: 4-37. Data from 157 countries.

- Variables of interest:
 - ▶ Y : Level of democracy, measured as the 10-year average of Freedom House ratings
 - ▶ X_1 : Country income, measured as $\log(\text{GDP per capita in \$1000s})$
 - ▶ X_2 : Ethnic heterogeneity (continuous) or British colonial heritage (binary)
- With one predictor we ask: Does income (X_1) predict the level of democracy (Y)?
- With two predictors we ask questions like: Does income (X_1) predict the level of democracy (Y), once we “control” for ethnic heterogeneity or British colonial heritage (X_2)?
- Let’s talk about what is meant by “controlling for another variable” with linear regression.

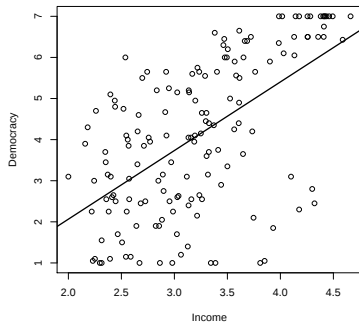
Simple Regression of Democracy on Income

Simple Regression of Democracy on Income

- Let's look at the bivariate regression of Democracy on Income:

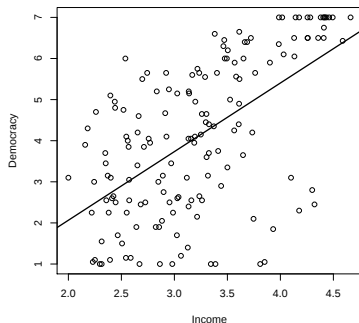
$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_i$$

$$\widehat{Demo} = -1.26 + 1.6 \text{ Log}(GDP)$$



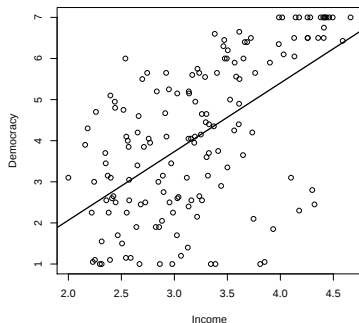
Simple Regression of Democracy on Income

- But we can use more information in our prediction equation.



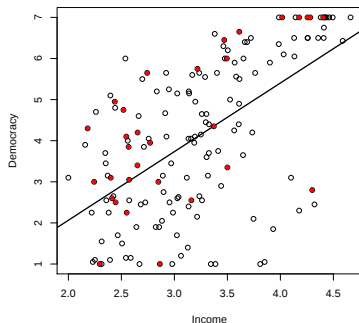
Simple Regression of Democracy on Income

- But we can use more information in our prediction equation.
- For example, some countries were originally British colonies and others were not:



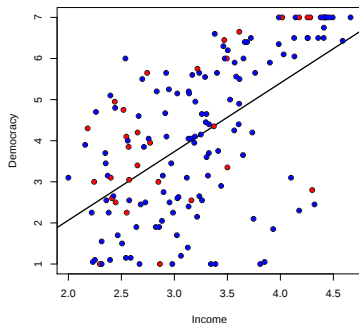
Simple Regression of Democracy on Income

- But we can use more information in our prediction equation.
- For example, some countries were originally British colonies and others were not:
 - ▶ Former British colonies tend to have higher levels of democracy



Simple Regression of Democracy on Income

- But we can use more information in our prediction equation.
- For example, some countries were originally British colonies and others were not:
 - ▶ Former British colonies tend to have higher levels of democracy
 - ▶ Non-colony countries tend to have lower levels of democracy



Adding a Covariate

Adding a Covariate

How do we do this?

Adding a Covariate

How do we do this? We can generalize the prediction equation:

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_{1i} + \hat{\beta}_2 x_{2i}$$

This implies that we want to predict y using the information we have about x_1 and x_2 , and we are assuming a linear functional form.

Adding a Covariate

How do we do this? We can generalize the prediction equation:

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_{1i} + \hat{\beta}_2 x_{2i}$$

This implies that we want to predict y using the information we have about x_1 and x_2 , and we are assuming a linear functional form.

Notice that now we write X_{ji} where:

- $j = 1, \dots, k$ is the index for the explanatory variables

Adding a Covariate

How do we do this? We can generalize the prediction equation:

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_{1i} + \hat{\beta}_2 x_{2i}$$

This implies that we want to predict y using the information we have about x_1 and x_2 , and we are assuming a linear functional form.

Notice that now we write X_{ji} where:

- $j = 1, \dots, k$ is the index for the explanatory variables
- $i = 1, \dots, n$ is the index for the observation

Adding a Covariate

How do we do this? We can generalize the prediction equation:

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_{1i} + \hat{\beta}_2 x_{2i}$$

This implies that we want to predict y using the information we have about x_1 and x_2 , and we are assuming a linear functional form.

Notice that now we write X_{ji} where:

- $j = 1, \dots, k$ is the index for the explanatory variables
- $i = 1, \dots, n$ is the index for the observation
- we often omit i to avoid clutter

Adding a Covariate

How do we do this? We can generalize the prediction equation:

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_{1i} + \hat{\beta}_2 x_{2i}$$

This implies that we want to predict y using the information we have about x_1 and x_2 , and we are assuming a linear functional form.

Notice that now we write X_{ji} where:

- $j = 1, \dots, k$ is the index for the explanatory variables
- $i = 1, \dots, n$ is the index for the observation
- we often omit i to avoid clutter

In words:

$$\widehat{Democracy} = \hat{\beta}_0 + \hat{\beta}_1 \text{Log}(GDP) + \hat{\beta}_2 \text{Colony}$$

Interpreting a Binary Covariate

Interpreting a Binary Covariate

Assume X_{2i} indicates whether country i used to be a British colony.

When $X_2 = 0$, the model becomes:

Interpreting a Binary Covariate

Assume X_{2i} indicates whether country i used to be a British colony.

When $X_2 = 0$, the model becomes:

$$\begin{aligned}\hat{y} &= \hat{\beta}_0 + \hat{\beta}_1 x_1 + \hat{\beta}_2 0 \\ &= \hat{\beta}_0 + \hat{\beta}_1 x_1\end{aligned}$$

Interpreting a Binary Covariate

Assume X_{2i} indicates whether country i used to be a British colony.

When $X_2 = 0$, the model becomes:

$$\begin{aligned}\hat{y} &= \hat{\beta}_0 + \hat{\beta}_1 x_1 + \hat{\beta}_2 0 \\ &= \hat{\beta}_0 + \hat{\beta}_1 x_1\end{aligned}$$

When $X_2 = 1$, the model becomes:

Interpreting a Binary Covariate

Assume X_{2i} indicates whether country i used to be a British colony.

When $X_2 = 0$, the model becomes:

$$\begin{aligned}\hat{y} &= \hat{\beta}_0 + \hat{\beta}_1 x_1 + \hat{\beta}_2 0 \\ &= \hat{\beta}_0 + \hat{\beta}_1 x_1\end{aligned}$$

When $X_2 = 1$, the model becomes:

$$\begin{aligned}\hat{y} &= \hat{\beta}_0 + \hat{\beta}_1 x_1 + \hat{\beta}_2 1 \\ &= (\hat{\beta}_0 + \hat{\beta}_2) + \hat{\beta}_1 x_1\end{aligned}$$

What does this mean?

Interpreting a Binary Covariate

Assume X_{2i} indicates whether country i used to be a British colony.

When $X_2 = 0$, the model becomes:

$$\begin{aligned}\hat{y} &= \hat{\beta}_0 + \hat{\beta}_1 x_1 + \hat{\beta}_2 0 \\ &= \hat{\beta}_0 + \hat{\beta}_1 x_1\end{aligned}$$

When $X_2 = 1$, the model becomes:

$$\begin{aligned}\hat{y} &= \hat{\beta}_0 + \hat{\beta}_1 x_1 + \hat{\beta}_2 1 \\ &= (\hat{\beta}_0 + \hat{\beta}_2) + \hat{\beta}_1 x_1\end{aligned}$$

What does this mean? We are fitting two lines with the **same slope** but **different intercepts**.

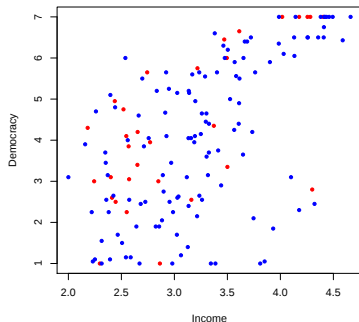
Regression of Democracy on Income

From R, we obtain estimates

$\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$:

Coefficients:

	Estimate
(Intercept)	-1.5060
GDP90LGN	1.7059
BRITCOL	0.5881



Regression of Democracy on Income

From R, we obtain estimates

$\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$:

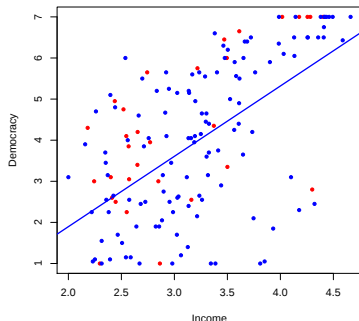
Coefficients:

	Estimate
(Intercept)	-1.5060
GDP90LGN	1.7059
BRITCOL	0.5881

- Non-British colonies:

$$\hat{y} = \hat{\beta}_0 + \hat{\beta}_1 x_1$$

$$\hat{y} = -1.5 + 1.7 x_1$$



Regression of Democracy on Income

From R, we obtain estimates

$\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$:

Coefficients:

	Estimate
(Intercept)	-1.5060
GDP90LGN	1.7059
BRITCOL	0.5881

- Non-British colonies:

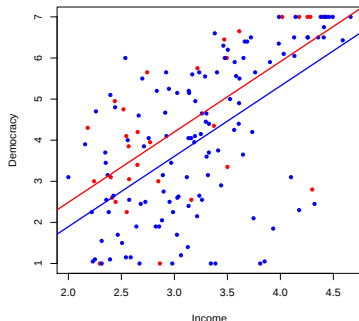
$$\hat{y} = \hat{\beta}_0 + \hat{\beta}_1 x_1$$

$$\hat{y} = -1.5 + 1.7 x_1$$

- Former British colonies:

$$\hat{y} = (\hat{\beta}_0 + \hat{\beta}_2) + \hat{\beta}_1 x_1$$

$$\hat{y} = -.92 + 1.7 x_1$$

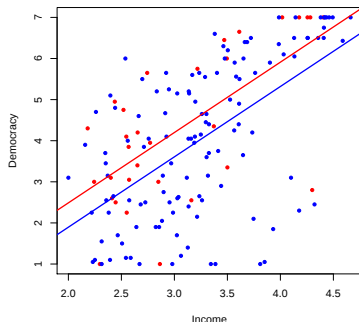


Regression of Democracy on Income

Our prediction equation is:

$$\hat{y} = -1.5 + 1.7x_1 + .58x_2$$

Where do these quantities appear on the graph?



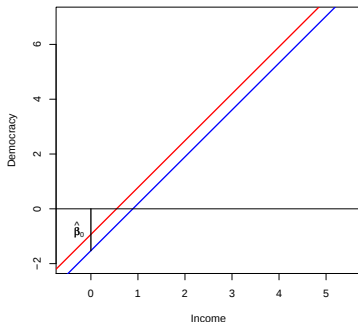
Regression of Democracy on Income

Our prediction equation is:

$$\hat{y} = -1.5 + 1.7x_1 + .58x_2$$

Where do these quantities appear on the graph?

- $\hat{\beta}_0 = -1.5$ is the intercept for the prediction line for non-British colonies.



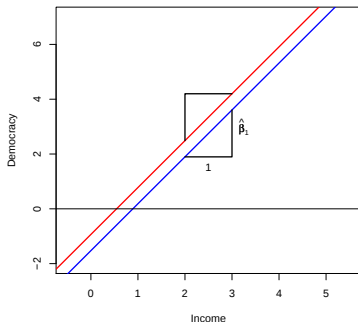
Regression of Democracy on Income

Our prediction equation is:

$$\hat{y} = -1.5 + 1.7 x_1 + .58 x_2$$

Where do these quantities appear on the graph?

- $\hat{\beta}_0 = -1.5$ is the intercept for the prediction line for non-British colonies.
- $\hat{\beta}_1 = 1.7$ is the slope for both lines.



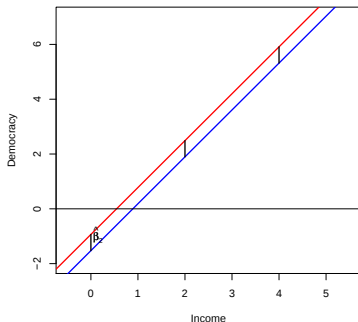
Regression of Democracy on Income

Our prediction equation is:

$$\hat{y} = -1.5 + 1.7x_1 + .58x_2$$

Where do these quantities appear on the graph?

- $\hat{\beta}_0 = -1.5$ is the intercept for the prediction line for non-British colonies.
- $\hat{\beta}_1 = 1.7$ is the slope for both lines.
- $\hat{\beta}_2 = .58$ is the vertical distance between the two lines for Ex-British colonies and non-colonies respectively



Example Interpretation of the Coefficients

- Let's review what we've seen so far:

	Intercept for X_1	Slope for X_1
Non-Colony ($X_2 = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Former Colony ($X_2 = 1$)	$\hat{\beta}_0 + \hat{\beta}_2$	$\hat{\beta}_1$

Example Interpretation of the Coefficients

- Let's review what we've seen so far:

	Intercept for X_1	Slope for X_1
Non-Colony ($X_2 = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Former Colony ($X_2 = 1$)	$\hat{\beta}_0 + \hat{\beta}_2$	$\hat{\beta}_1$

- In this example, we have:

$$\hat{Y}_i = -1.5060 + 1.7059 \cdot X_1 + 0.5881 \cdot X_2$$

Example Interpretation of the Coefficients

- Let's review what we've seen so far:

	Intercept for X_1	Slope for X_1
Non-Colony ($X_2 = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Former Colony ($X_2 = 1$)	$\hat{\beta}_0 + \hat{\beta}_2$	$\hat{\beta}_1$

- In this example, we have:

$$\hat{Y}_i = -1.5060 + 1.7059 \cdot X_1 + 0.5881 \cdot X_2$$

- We can read these as:

Example Interpretation of the Coefficients

- Let's review what we've seen so far:

	Intercept for X_1	Slope for X_1
Non-Colony ($X_2 = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Former Colony ($X_2 = 1$)	$\hat{\beta}_0 + \hat{\beta}_2$	$\hat{\beta}_1$

- In this example, we have:

$$\hat{Y}_i = -1.5060 + 1.7059 \cdot X_1 + 0.5881 \cdot X_2$$

- We can read these as:
 - $\hat{\beta}_0$: average democracy for non-British colony with log income of 0 is **-1.5060** (note this is an extrapolation for this data!).

Example Interpretation of the Coefficients

- Let's review what we've seen so far:

	Intercept for X_1	Slope for X_1
Non-Colony ($X_2 = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Former Colony ($X_2 = 1$)	$\hat{\beta}_0 + \hat{\beta}_2$	$\hat{\beta}_1$

- In this example, we have:

$$\hat{Y}_i = -1.5060 + 1.7059 \cdot X_1 + 0.5881 \cdot X_2$$

- We can read these as:
 - ▶ $\hat{\beta}_0$: average democracy for non-British colony with log income of 0 is -1.5060 (note this is an extrapolation for this data!).
 - ▶ $\hat{\beta}_1$: countries with a one unit higher log income have on average a 1.7059 higher democracy score.

Example Interpretation of the Coefficients

- Let's review what we've seen so far:

	Intercept for X_1	Slope for X_1
Non-Colony ($X_2 = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Former Colony ($X_2 = 1$)	$\hat{\beta}_0 + \hat{\beta}_2$	$\hat{\beta}_1$

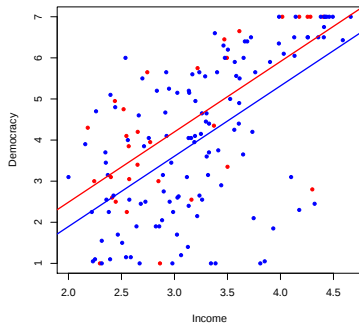
- In this example, we have:

$$\hat{Y}_i = -1.5060 + 1.7059 \cdot X_1 + 0.5881 \cdot X_2$$

- We can read these as:
 - $\hat{\beta}_0$: average democracy for non-British colony with log income of 0 is -1.5060 (note this is an extrapolation for this data!).
 - $\hat{\beta}_1$: countries with a one unit higher log income have on average a 1.7059 higher democracy score.
 - $\hat{\beta}_2$: former british colonies are predicted to have a 0.5881 higher average democracy score than non-british colonies with the same level of income.

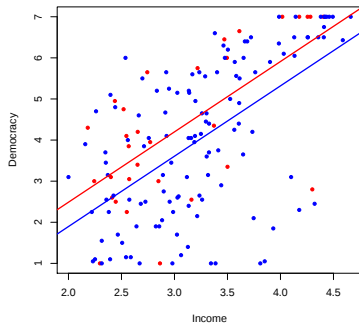
Fitting a regression plane

- We have considered an example of multiple regression with one **continuous** explanatory variable and one **binary** explanatory variable.



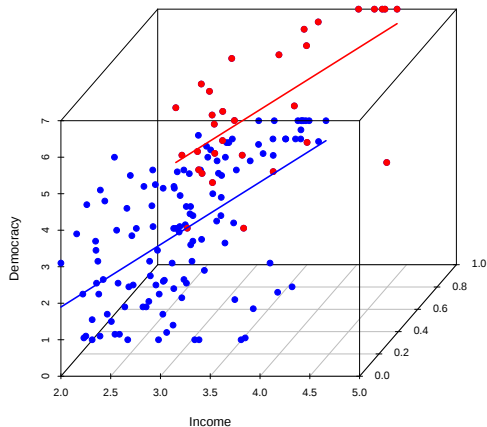
Fitting a regression plane

- We have considered an example of multiple regression with one **continuous** explanatory variable and one **binary** explanatory variable.
- This is easy to represent graphically in **two dimensions** because we can use colors to distinguish the two groups in the data.



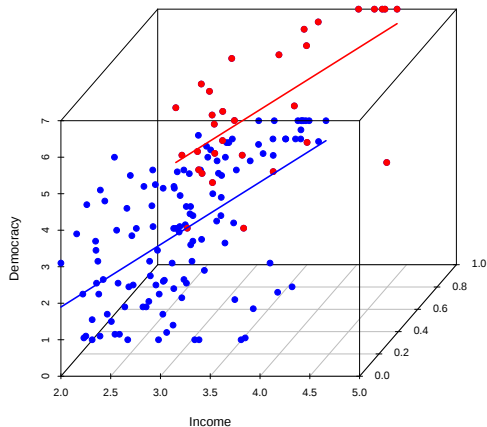
Regression of Democracy on Income

- These observations are actually located in a **three-dimensional** space.



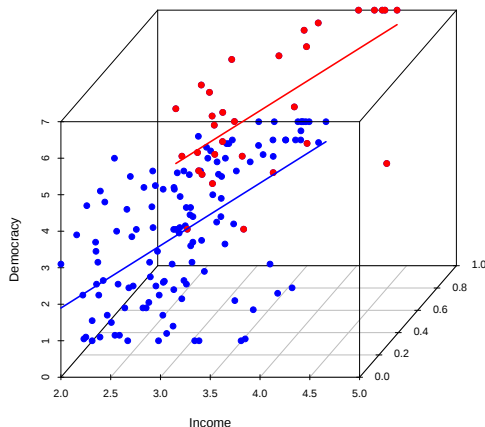
Regression of Democracy on Income

- These observations are actually located in a **three-dimensional** space.
- We can try to represent this using a **3D scatterplot**.



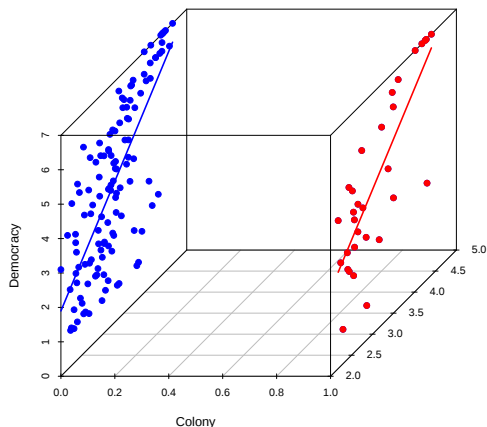
Regression of Democracy on Income

- These observations are actually located in a **three-dimensional** space.
- We can try to represent this using a **3D scatterplot**.
- In this view, we are looking at the data from the **Income side**; the two regression lines are drawn in the appropriate locations.



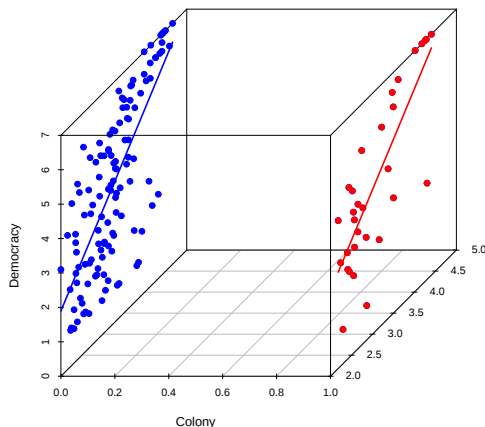
Regression of Democracy on Income

- We can also look at the 3D scatterplot from the **British colony side**.



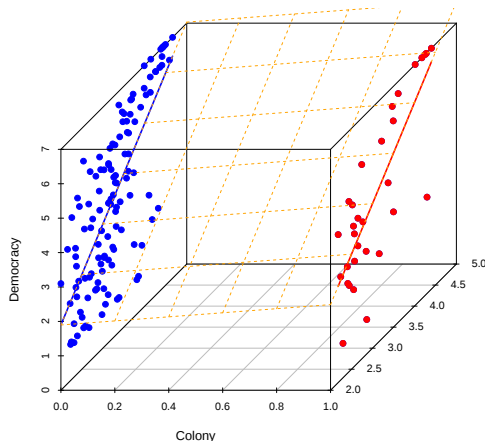
Regression of Democracy on Income

- We can also look at the 3D scatterplot from the **British colony side**.
- While the British colonial status variable is either 0 or 1, there is nothing in the prediction equation that requires this to be the case.



Regression of Democracy on Income

- We can also look at the 3D scatterplot from the **British colony side**.
- While the British colonial status variable is either 0 or 1, there is nothing in the prediction equation that requires this to be the case.
- In fact, the prediction equation defines a **regression plane** that connects the lines when $x_2 = 0$ and $x_2 = 1$.



Regression with two continuous variables

- Since we fit a regression plane to the data whenever we have two explanatory variables, it is easy to move to a case with **two continuous** explanatory variables.

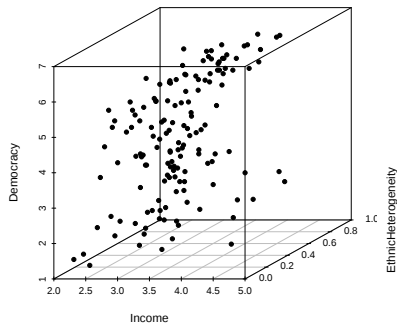
Regression with two continuous variables

- Since we fit a regression plane to the data whenever we have two explanatory variables, it is easy to move to a case with **two continuous** explanatory variables.
- For example, we might want to use:
 - ▶ X_1 Income and X_2 Ethnic Heterogeneity
 - ▶ Y Democracy

$$\widehat{\text{Democracy}} = \hat{\beta}_0 + \hat{\beta}_1 \text{Income} + \hat{\beta}_2 \text{Ethnic Heterogeneity}$$

Regression of Democracy on Income

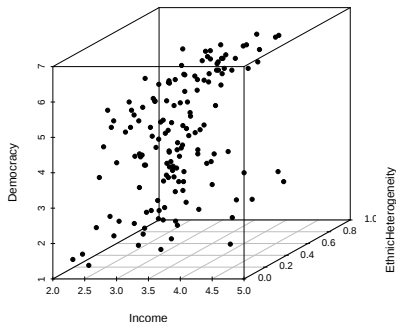
- We can plot the points in a 3D scatterplot.



Regression of Democracy on Income

- We can plot the points in a 3D scatterplot.
- R returns:
 - ▶ $\hat{\beta}_0 = -.71$
 - ▶ $\hat{\beta}_1 = 1.6$ for Income
 - ▶ $\hat{\beta}_2 = -.6$ for Ethnic Heterogeneity

How does this look graphically?

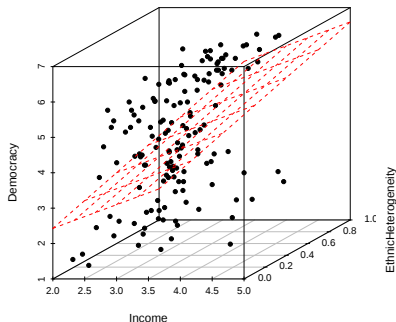


Regression of Democracy on Income

- We can plot the points in a 3D scatterplot.
- R returns:
 - ▶ $\hat{\beta}_0 = -.71$
 - ▶ $\hat{\beta}_1 = 1.6$ for Income
 - ▶ $\hat{\beta}_2 = -.6$ for Ethnic Heterogeneity

How does this look graphically?

- These estimates define a **regression plane** through the data.

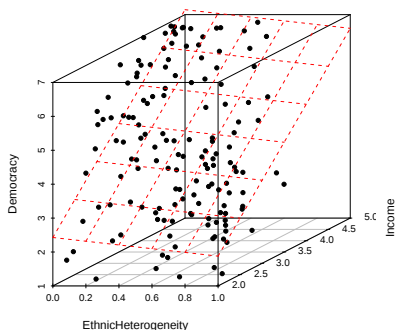


Regression of Democracy on Income

- We can plot the points in a 3D scatterplot.
- R returns:
 - ▶ $\hat{\beta}_0 = -.71$
 - ▶ $\hat{\beta}_1 = 1.6$ for Income
 - ▶ $\hat{\beta}_2 = -.6$ for Ethnic Heterogeneity

How does this look graphically?

- These estimates define a **regression plane** through the data.



Interpreting a Continuous Covariate

- The coefficient estimates have a similar interpretation in this case as they did in the Income-British Colony example.

Interpreting a Continuous Covariate

- The coefficient estimates have a similar interpretation in this case as they did in the Income-British Colony example.
- For example, $\hat{\beta}_1 = 1.6$ represents our prediction of the difference in Democracy between two observations that differ by one unit of Income **but have the same value of Ethnic Heterogeneity**.

Interpreting a Continuous Covariate

- The coefficient estimates have a similar interpretation in this case as they did in the Income-British Colony example.
- For example, $\hat{\beta}_1 = 1.6$ represents our prediction of the difference in Democracy between two observations that differ by one unit of Income **but have the same value of Ethnic Heterogeneity**.
- The slope estimates have an interpretation in terms of the partial derivative:

$$\frac{\partial(y = \beta_0 + \beta_1 X_1 + \beta_2 X_2)}{\partial X_1} =$$

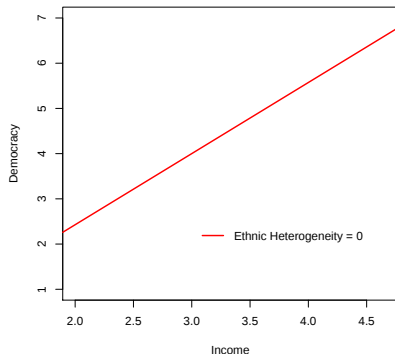
Interpreting a Continuous Covariate

- The coefficient estimates have a similar interpretation in this case as they did in the Income-British Colony example.
- For example, $\hat{\beta}_1 = 1.6$ represents our prediction of the difference in Democracy between two observations that differ by one unit of Income **but have the same value of Ethnic Heterogeneity**.
- The slope estimates have an interpretation in terms of the partial derivative:

$$\frac{\partial(y = \beta_0 + \beta_1 X_1 + \beta_2 X_2)}{\partial X_1} = \beta_1$$

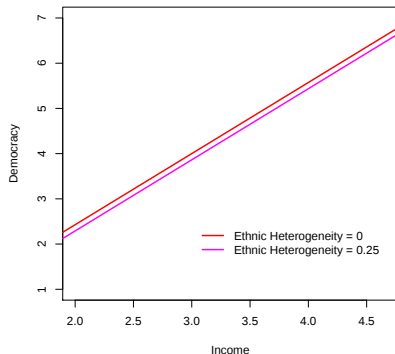
Interpreting a Continuous Covariate

- Again, we can think of this as defining a regression line for the relationship between Democracy and Income at every level of Ethnic Heterogeneity.



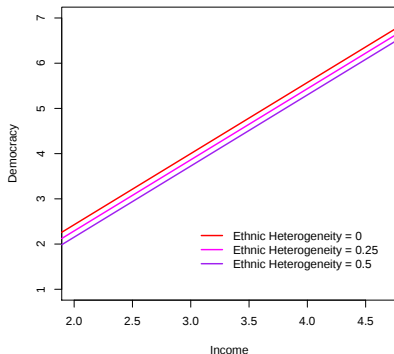
Interpreting a Continuous Covariate

- Again, we can think of this as defining a regression line for the relationship between Democracy and Income at every level of Ethnic Heterogeneity.
- All of these lines are parallel since they have the slope $\hat{\beta}_1 = 1.6$



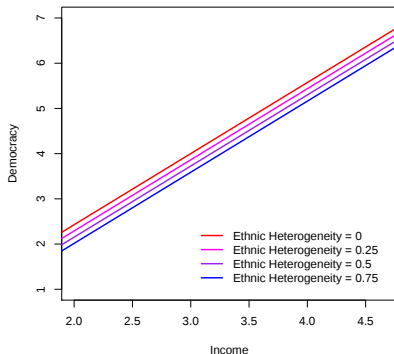
Interpreting a Continuous Covariate

- Again, we can think of this as defining a regression line for the relationship between Democracy and Income at every level of Ethnic Heterogeneity.
- All of these lines are parallel since they have the slope $\hat{\beta}_1 = 1.6$
- The lines shift up or down based on the value of Ethnic Heterogeneity.



Interpreting a Continuous Covariate

- Again, we can think of this as defining a regression line for the relationship between Democracy and Income at every level of Ethnic Heterogeneity.
- All of these lines are parallel since they have the slope $\hat{\beta}_1 = 1.6$
- The lines shift up or down based on the value of Ethnic Heterogeneity.



More Complex Predictions

- We can also use the coefficient estimates for more complex predictions that involve changing multiple variables simultaneously.

More Complex Predictions

- We can also use the coefficient estimates for more complex predictions that involve changing multiple variables simultaneously.
- Consider our results for the regression of democracy on X_1 income and X_2 ethnic heterogeneity:
 - ▶ $\hat{\beta}_0 = -.71$
 - ▶ $\hat{\beta}_1 = 1.6$
 - ▶ $\hat{\beta}_2 = -.6$

More Complex Predictions

- We can also use the coefficient estimates for more complex predictions that involve changing multiple variables simultaneously.
- Consider our results for the regression of democracy on X_1 income and X_2 ethnic heterogeneity:
 - ▶ $\hat{\beta}_0 = -.71$
 - ▶ $\hat{\beta}_1 = 1.6$
 - ▶ $\hat{\beta}_2 = -.6$
- What is the predicted difference in democracy between
 - ▶ **Chile** with $X_1 = 3.5$ and $X_2 = .06$?
 - ▶ **China** with $X_1 = 2.5$ and $X_2 = .5$?

More Complex Predictions

- We can also use the coefficient estimates for more complex predictions that involve changing multiple variables simultaneously.
- Consider our results for the regression of democracy on X_1 income and X_2 ethnic heterogeneity:
 - ▶ $\hat{\beta}_0 = -.71$
 - ▶ $\hat{\beta}_1 = 1.6$
 - ▶ $\hat{\beta}_2 = -.6$
- What is the predicted difference in democracy between
 - ▶ **Chile** with $X_1 = 3.5$ and $X_2 = .06$?
 - ▶ **China** with $X_1 = 2.5$ and $X_2 = .5$?
- Predicted democracy is
 - ▶ $-.71 + 1.6 \cdot 3.5 - .6 \cdot .06 = 4.8$ for **Chile**

More Complex Predictions

- We can also use the coefficient estimates for more complex predictions that involve changing multiple variables simultaneously.
- Consider our results for the regression of democracy on X_1 income and X_2 ethnic heterogeneity:
 - ▶ $\hat{\beta}_0 = -.71$
 - ▶ $\hat{\beta}_1 = 1.6$
 - ▶ $\hat{\beta}_2 = -.6$
- What is the predicted difference in democracy between
 - ▶ **Chile** with $X_1 = 3.5$ and $X_2 = .06$?
 - ▶ **China** with $X_1 = 2.5$ and $X_2 = .5$?
- Predicted democracy is
 - ▶ $-.71 + 1.6 \cdot 3.5 - .6 \cdot .06 = 4.8$ for **Chile**
 - ▶ $-.71 + 1.6 \cdot 2.5 - .6 \cdot 0.5 = 3$ for **China**.

More Complex Predictions

- We can also use the coefficient estimates for more complex predictions that involve changing multiple variables simultaneously.
- Consider our results for the regression of democracy on X_1 income and X_2 ethnic heterogeneity:
 - ▶ $\hat{\beta}_0 = -.71$
 - ▶ $\hat{\beta}_1 = 1.6$
 - ▶ $\hat{\beta}_2 = -.6$
- What is the predicted difference in democracy between
 - ▶ **Chile** with $X_1 = 3.5$ and $X_2 = .06$?
 - ▶ **China** with $X_1 = 2.5$ and $X_2 = .5$?
- Predicted democracy is
 - ▶ $-.71 + 1.6 \cdot 3.5 - .6 \cdot .06 = 4.8$ for **Chile**
 - ▶ $-.71 + 1.6 \cdot 2.5 - .6 \cdot 0.5 = 3$ for **China**.

Predicted difference is thus: 1.8 or $(3.5 - 2.5)\hat{\beta}_1 + (.06 - .5)\hat{\beta}_2$

Dummy Variables

Dummy Variables

- A **dummy variable** (a.k.a. indicator variable, binary variable, etc.) is a variable that is coded 1 or 0 only.

Dummy Variables

- A **dummy variable** (a.k.a. indicator variable, binary variable, etc.) is a variable that is coded 1 or 0 only.
- We use dummy variables in regression to represent qualitative information through **categorical variables** such as different subgroups of the sample (e.g. regions, old and young respondents, etc.)

Dummy Variables

- A **dummy variable** (a.k.a. indicator variable, binary variable, etc.) is a variable that is coded 1 or 0 only.
- We use dummy variables in regression to represent qualitative information through **categorical variables** such as different subgroups of the sample (e.g. regions, old and young respondents, etc.)
- By including dummy variables into our regression function, we can easily obtain the **conditional mean of the outcome variable for each category**.

How Can I Use a Dummy Variable?

- Consider the easiest case with two categories. The type of electoral system of country i is given by:

$$X_i \in \{Proportional, Majoritarian\}$$

How Can I Use a Dummy Variable?

- Consider the easiest case with two categories. The type of electoral system of country i is given by:

$$X_i \in \{Proportional, Majoritarian\}$$

- For this we use a single dummy variable which is coded like:

$$D_i = \begin{cases} 1 & \text{if country } i \text{ has a Majoritarian Electoral System} \\ 0 & \text{if country } i \text{ has a Proportional Electoral System} \end{cases}$$

Dummy Variables for Multiple Categories

- More generally, let's say X measures which of m categories each unit i belongs to. E.g. the type of electoral system or region of country i is given by:
 - ▶ $X_i \in \{Proportional, Majoritarian\}$ so $m = 2$
 - ▶ $X_i \in \{Asia, Africa, LatinAmerica, OECD, Transition\}$ so $m = 5$

Dummy Variables for Multiple Categories

- More generally, let's say X measures which of m categories each unit i belongs to. E.g. the type of electoral system or region of country i is given by:
 - ▶ $X_i \in \{Proportional, Majoritarian\}$ so $m = 2$
 - ▶ $X_i \in \{Asia, Africa, LatinAmerica, OECD, Transition\}$ so $m = 5$
- To incorporate this information into our regression function we usually create $m - 1$ dummy variables, one for each of the $m - 1$ categories.

Dummy Variables for Multiple Categories

- More generally, let's say X measures which of m categories each unit i belongs to. E.g. the type of electoral system or region of country i is given by:
 - ▶ $X_i \in \{Proportional, Majoritarian\}$ so $m = 2$
 - ▶ $X_i \in \{Asia, Africa, LatinAmerica, OECD, Transition\}$ so $m = 5$
- To incorporate this information into our regression function we usually create $m - 1$ dummy variables, one for each of the $m - 1$ categories.
- Why not all m ?

Dummy Variables for Multiple Categories

- More generally, let's say X measures which of m categories each unit i belongs to. E.g. the type of electoral system or region of country i is given by:
 - ▶ $X_i \in \{Proportional, Majoritarian\}$ so $m = 2$
 - ▶ $X_i \in \{Asia, Africa, LatinAmerica, OECD, Transition\}$ so $m = 5$
- To incorporate this information into our regression function we usually create $m - 1$ dummy variables, one for each of the $m - 1$ categories.
- Why not all m ? Including all m category indicators as dummies would be indistinguishable from the intercept:

$$D_m = 1 - (D_1 + \dots + D_{m-1})$$

Dummy Variables for Multiple Categories

- More generally, let's say X measures which of m categories each unit i belongs to. E.g. the type of electoral system or region of country i is given by:
 - ▶ $X_i \in \{Proportional, Majoritarian\}$ so $m = 2$
 - ▶ $X_i \in \{Asia, Africa, LatinAmerica, OECD, Transition\}$ so $m = 5$
- To incorporate this information into our regression function we usually create $m - 1$ dummy variables, one for each of the $m - 1$ categories.
- Why not all m ? Including all m category indicators as dummies would be indistinguishable from the intercept:

$$D_m = 1 - (D_1 + \dots + D_{m-1})$$

- The omitted category is our **baseline case** (also called a **reference category**) against which we compare the conditional means of Y for the other $m - 1$ categories.

Example: Regions of the World

- Consider the case of our “polytomous” variable world region with $m = 5$:

$$X_i \in \{Asia, Africa, LatinAmerica, OECD, Transition\}$$

Example: Regions of the World

- Consider the case of our “polytomous” variable world region with $m = 5$:
 $X_i \in \{Asia, Africa, LatinAmerica, OECD, Transition\}$
- This five-category classification can be represented in the regression equation by introducing $m - 1 = 4$ dummy regressors:

Category	D_1	D_2	D_3	D_4
Asia	1	0	0	0
Africa	0	1	0	0
LatinAmerica	0	0	1	0
OECD	0	0	0	1
Transition	0	0	0	0

Example: Regions of the World

- Consider the case of our “polytomous” variable world region with $m = 5$:
 $X_i \in \{Asia, Africa, LatinAmerica, OECD, Transition\}$
- This five-category classification can be represented in the regression equation by introducing $m - 1 = 4$ dummy regressors:

Category	D_1	D_2	D_3	D_4
Asia	1	0	0	0
Africa	0	1	0	0
LatinAmerica	0	0	1	0
OECD	0	0	0	1
Transition	0	0	0	0

Our regression equation is:

$$Y = \beta_0 + \beta_1 D_1 + \beta_2 D_2 + \beta_3 D_3 + \beta_4 D_4 + u$$

Example: GDP per capita on Regions

```
----- R Code -----
> summary(lm(REALGDPCAP ~ Asia + Africa + LatAmerica + Oecd, data = D))

~~~~~

Coefficients:
              Estimate Std. Error t value Pr(>|t|)
(Intercept)   4452.7      783.4    5.684 2.07e-07 ***
Asia           148.9     1149.8    0.129  0.8973
Africa        -2552.8     1204.5   -2.119  0.0372 *
LatAmerica    -271.3     1007.0   -0.269  0.7883
Oecd           9671.3     1007.0    9.604 5.74e-15 ***
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 3034 on 80 degrees of freedom
Multiple R-squared:  0.7096,    Adjusted R-squared:  0.6951
F-statistic: 48.88 on 4 and 80 DF,  p-value: < 2.2e-16
```

What does β_0 mean?

Example: GDP per capita on Regions

```
----- R Code -----  
> summary(lm(REALGDPCAP ~ Asia + Africa + LatAmerica + Oecd, data = D))  
  
~~~~~  
  
Coefficients:  
              Estimate Std. Error t value Pr(>|t|)  
(Intercept)  4452.7      783.4   5.684 2.07e-07 ***  
Asia         148.9      1149.8   0.129  0.8973  
Africa      -2552.8     1204.5  -2.119  0.0372 *  
LatAmerica  -271.3     1007.0  -0.269  0.7883  
Oecd        9671.3     1007.0   9.604 5.74e-15 ***  
---  
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1  
  
Residual standard error: 3034 on 80 degrees of freedom  
Multiple R-squared:  0.7096,    Adjusted R-squared:  0.6951  
F-statistic: 48.88 on 4 and 80 DF,  p-value: < 2.2e-16
```

What does β_0 mean?

$$\beta_0 = E[GDP | D_j = 0 \text{ for all } j] = E[GDP | \text{Transition}]$$

So the mean for the baseline category shows up as the intercept.

Example: GDP per capita on Regions

```
----- R Code -----  
> summary(lm(REALGDPCAP ~ Asia + Africa + LatAmerica + Oecd, data = D))  
  
~~~~~  
  
Coefficients:  
              Estimate Std. Error t value Pr(>|t|)  
(Intercept)   4452.7      783.4    5.684 2.07e-07 ***  
Asia           148.9     1149.8    0.129  0.8973  
Africa       -2552.8     1204.5   -2.119  0.0372 *  
LatAmerica    -271.3     1007.0   -0.269  0.7883  
Oecd          9671.3     1007.0    9.604 5.74e-15 ***  
---  
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1  
  
Residual standard error: 3034 on 80 degrees of freedom  
Multiple R-squared:  0.7096,    Adjusted R-squared:  0.6951  
F-statistic: 48.88 on 4 and 80 DF,  p-value: < 2.2e-16
```

What does β_{Africa} mean?

Example: GDP per capita on Regions

```
_____ R Code _____  
> summary(lm(REALGDPCAP ~ Asia + Africa + LatAmerica + Oecd, data = D))  
  
~~~~~  
  
Coefficients:  
              Estimate Std. Error t value Pr(>|t|)  
(Intercept)   4452.7      783.4    5.684 2.07e-07 ***  
Asia           148.9     1149.8    0.129  0.8973  
Africa        -2552.8     1204.5   -2.119  0.0372 *  
LatAmerica    -271.3     1007.0   -0.269  0.7883  
Oecd          9671.3     1007.0    9.604 5.74e-15 ***  
---  
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1  
  
Residual standard error: 3034 on 80 degrees of freedom  
Multiple R-squared:  0.7096,    Adjusted R-squared:  0.6951  
F-statistic: 48.88 on 4 and 80 DF,  p-value: < 2.2e-16
```

What does β_{Africa} mean?

$$\beta_{Africa} = E[GDP|Africa] - E[GDP|Transition]$$

The difference in means between the baseline and that category.

Example: GDP per capita on Regions

```
_____ R Code _____  
> summary(lm(REALGDPCAP ~ Asia + Africa + LatAmerica + Oecd, data = D))  
  
~~~~~  
  
Coefficients:  
              Estimate Std. Error t value Pr(>|t|)  
(Intercept)   4452.7      783.4   5.684 2.07e-07 ***  
Asia           148.9      1149.8   0.129  0.8973  
Africa        -2552.8     1204.5  -2.119  0.0372 *  
LatAmerica    -271.3      1007.0  -0.269  0.7883  
Oecd          9671.3      1007.0   9.604 5.74e-15 ***  
---  
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1  
  
Residual standard error: 3034 on 80 degrees of freedom  
Multiple R-squared:  0.7096,    Adjusted R-squared:  0.6951  
F-statistic: 48.88 on 4 and 80 DF,  p-value: < 2.2e-16
```

Do Latin America economies have higher or lower average GDP than Asian economies?

Example: GDP per capita on Regions

```
_____ R Code _____
> summary(lm(REALGDPCAP ~ Asia + Africa + LatAmerica + Oecd, data = D))

~~~~~

Coefficients:
              Estimate Std. Error t value Pr(>|t|)
(Intercept)   4452.7      783.4    5.684 2.07e-07 ***
Asia          148.9       1149.8    0.129  0.8973
Africa       -2552.8      1204.5   -2.119  0.0372 *
LatAmerica   -271.3       1007.0   -0.269  0.7883
Oecd         9671.3       1007.0    9.604 5.74e-15 ***
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 3034 on 80 degrees of freedom
Multiple R-squared:  0.7096,    Adjusted R-squared:  0.6951
F-statistic: 48.88 on 4 and 80 DF,  p-value: < 2.2e-16
```

Do Latin America economies have higher or lower average GDP than Asian economies?

$\beta_{LatAmerica} = E[GDP|LatAmerica] - E[GDP|Transition]$, and

$\beta_{Asia} = E[GDP|Asia] - E[GDP|Transition]$, so

Example: GDP per capita on Regions

```
_____ R Code _____
> summary(lm(REALGDPCAP ~ Asia + Africa + LatAmerica + Oecd, data = D))

~~~~~

Coefficients:
              Estimate Std. Error t value Pr(>|t|)
(Intercept)   4452.7      783.4    5.684 2.07e-07 ***
Asia          148.9       1149.8    0.129  0.8973
Africa       -2552.8     1204.5   -2.119  0.0372 *
LatAmerica   -271.3      1007.0   -0.269  0.7883
Oecd         9671.3      1007.0    9.604 5.74e-15 ***
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 3034 on 80 degrees of freedom
Multiple R-squared:  0.7096,    Adjusted R-squared:  0.6951
F-statistic: 48.88 on 4 and 80 DF,  p-value: < 2.2e-16
```

Do Latin America economies have higher or lower average GDP than Asian economies?

$$\beta_{\text{LatAmerica}} = E[\text{GDP}|\text{LatAmerica}] - E[\text{GDP}|\text{Transition}], \quad \text{and}$$

$$\beta_{\text{Asia}} = E[\text{GDP}|\text{Asia}] - E[\text{GDP}|\text{Transition}], \quad \text{so}$$

$$\beta_{\text{LatAmerica}} - \beta_{\text{Asia}} = E[\text{GDP}|\text{LatAmerica}] - E[\text{GDP}|\text{Asia}] = -420$$

Dealing with a Categorical Variable in R

R automatically expands an m -category variable into an $m - 1$ dummy variables:

```
R Code
> head(D$Region)
[1] LatAmerica Oecd      Oecd      LatAmerica Asia      LatAmerica
Levels: Africa Asia LatAmerica Oecd Transition

> summary(lm(REALGDPCAP ~ Region, data = D))

~~~~~
Coefficients:
              Estimate Std. Error t value Pr(>|t|)
(Intercept)    1899.9      914.9    2.077  0.0410 *
RegionAsia      2701.7     1243.0    2.173  0.0327 *
RegionLatAmerica 2281.5     1112.3    2.051  0.0435 *
RegionOecd      12224.2     1112.3   10.990 <2e-16 ***
RegionTransition 2552.8     1204.5    2.119  0.0372 *
---
Signif. codes:  0 *** 0.001 ** 0.01 * 0.05 . 0.1 1

Residual standard error: 3034 on 80 degrees of freedom
Multiple R-squared:  0.7096,    Adjusted R-squared:  0.6951
F-statistic: 48.88 on 4 and 80 DF,  p-value: < 2.2e-16
```

Dealing with a Categorical Variable in R

You can change the baseline category by the `relevel()` function:

```
_____ R Code _____
> D$Region <- relevel(D$Region, ref="Transition")
> summary(lm(REALGDPCAP ~ Region, data = D))

~~~~~
Coefficients:
                Estimate Std. Error t value Pr(>|t|)
(Intercept)      4452.7      783.4    5.684 2.07e-07 ***
RegionAfrica     -2552.8     1204.5   -2.119  0.0372 *
RegionAsia        148.9      1149.8    0.129  0.8973
RegionLatAmerica -271.3      1007.0   -0.269  0.7883
RegionOecd        9671.3      1007.0    9.604 5.74e-15 ***
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

Residual standard error: 3034 on 80 degrees of freedom
Multiple R-squared:  0.7096,    Adjusted R-squared:  0.6951
F-statistic: 48.88 on 4 and 80 DF,  p-value: < 2.2e-16
```

Saturated Models

Saturated Models

- A model is **saturated** if there are as many parameters as there are possible combination of the X_i variables.

Saturated Models

- A model is **saturated** if there are as many parameters as there are possible combination of the X_i variables.
- This happens when we have a dummy variable for every possible configuration of X variables in the data.

Saturated Models

- A model is **saturated** if there are as many parameters as there are possible combination of the X_i variables.
- This happens when we have a dummy variable for every possible configuration of X variables in the data.
- In this setting, linearity holds **by construction** because we are estimating a single mean for every combination of X_i variables.

Saturated Model Example

- Two binary variables, X_{1i} for marriage status and X_{2i} for having children.

Saturated Model Example

- Two binary variables, X_{1i} for marriage status and X_{2i} for having children.

Saturated Model Example

- Two binary variables, X_{1i} for marriage status and X_{2i} for having children.
- Four possible values of X_i , four possible values of $\mu(X_i)$:

$$E[Y_i | X_{1i} = 0, X_{2i} = 0]$$

$$E[Y_i | X_{1i} = 1, X_{2i} = 0]$$

$$E[Y_i | X_{1i} = 0, X_{2i} = 1]$$

$$E[Y_i | X_{1i} = 1, X_{2i} = 1]$$

Saturated Model Example

- Two binary variables, X_{1i} for marriage status and X_{2i} for having children.
- Four possible values of X_i , four possible values of $\mu(X_i)$:

$$E[Y_i|X_{1i} = 0, X_{2i} = 0] = \alpha$$

$$E[Y_i|X_{1i} = 1, X_{2i} = 0]$$

$$E[Y_i|X_{1i} = 0, X_{2i} = 1]$$

$$E[Y_i|X_{1i} = 1, X_{2i} = 1]$$

- We can write the CEF as follows:

$$E[Y_i|X_{1i}, X_{2i}] = \alpha + \beta X_{1i} + \gamma X_{2i} + \delta(X_{1i}X_{2i})$$

Saturated Model Example

- Two binary variables, X_{1i} for marriage status and X_{2i} for having children.
- Four possible values of X_i , four possible values of $\mu(X_i)$:

$$E[Y_i|X_{1i} = 0, X_{2i} = 0] = \alpha$$

$$E[Y_i|X_{1i} = 1, X_{2i} = 0] = \alpha + \beta$$

$$E[Y_i|X_{1i} = 0, X_{2i} = 1]$$

$$E[Y_i|X_{1i} = 1, X_{2i} = 1]$$

- We can write the CEF as follows:

$$E[Y_i|X_{1i}, X_{2i}] = \alpha + \beta X_{1i} + \gamma X_{2i} + \delta(X_{1i}X_{2i})$$

Saturated Model Example

- Two binary variables, X_{1i} for marriage status and X_{2i} for having children.
- Four possible values of X_i , four possible values of $\mu(X_i)$:

$$E[Y_i|X_{1i} = 0, X_{2i} = 0] = \alpha$$

$$E[Y_i|X_{1i} = 1, X_{2i} = 0] = \alpha + \beta$$

$$E[Y_i|X_{1i} = 0, X_{2i} = 1] = \alpha + \gamma$$

$$E[Y_i|X_{1i} = 1, X_{2i} = 1]$$

- We can write the CEF as follows:

$$E[Y_i|X_{1i}, X_{2i}] = \alpha + \beta X_{1i} + \gamma X_{2i} + \delta(X_{1i}X_{2i})$$

Saturated Model Example

- Two binary variables, X_{1i} for marriage status and X_{2i} for having children.
- Four possible values of X_i , four possible values of $\mu(X_i)$:

$$E[Y_i|X_{1i} = 0, X_{2i} = 0] = \alpha$$

$$E[Y_i|X_{1i} = 1, X_{2i} = 0] = \alpha + \beta$$

$$E[Y_i|X_{1i} = 0, X_{2i} = 1] = \alpha + \gamma$$

$$E[Y_i|X_{1i} = 1, X_{2i} = 1] = \alpha + \beta + \gamma + \delta$$

- We can write the CEF as follows:

$$E[Y_i|X_{1i}, X_{2i}] = \alpha + \beta X_{1i} + \gamma X_{2i} + \delta(X_{1i}X_{2i})$$

Saturated model example

$$E[Y_i|X_{1i}, X_{2i}] = \alpha + \beta X_{1i} + \gamma X_{2i} + \delta(X_{1i}X_{2i})$$

- Basically, each value of the CEF is being estimated separately.

Saturated model example

$$E[Y_i|X_{1i}, X_{2i}] = \alpha + \beta X_{1i} + \gamma X_{2i} + \delta(X_{1i}X_{2i})$$

- Basically, each value of the CEF is being estimated separately.
 - ▶ $\rightsquigarrow$ within-strata estimation.

Saturated model example

$$E[Y_i|X_{1i}, X_{2i}] = \alpha + \beta X_{1i} + \gamma X_{2i} + \delta(X_{1i}X_{2i})$$

- Basically, each value of the CEF is being estimated separately.
 - ▶ $\rightsquigarrow$ within-strata estimation.
 - ▶ No borrowing of information from across values of X_i .

Saturated model example

$$E[Y_i|X_{1i}, X_{2i}] = \alpha + \beta X_{1i} + \gamma X_{2i} + \delta(X_{1i}X_{2i})$$

- Basically, each value of the CEF is being estimated separately.
 - ▶ $\rightsquigarrow$ within-strata estimation.
 - ▶ No borrowing of information from across values of X_i .
- Requires a set of dummies for each categorical variable plus **all interactions**.

Saturated model example

$$E[Y_i|X_{1i}, X_{2i}] = \alpha + \beta X_{1i} + \gamma X_{2i} + \delta(X_{1i}X_{2i})$$

- Basically, each value of the CEF is being estimated separately.
 - ▶ $\rightsquigarrow$ within-strata estimation.
 - ▶ No borrowing of information from across values of X_i .
- Requires a set of dummies for each categorical variable plus **all interactions**.
- i.e. a series of dummies for each unique combination of X_i .

Saturated model example

- Ebonya Washington (AER) data from *AER* paper “Female socialization: how daughters affect their legislator fathers”

Saturated model example

- Ebonya Washington (AER) data from *AER* paper “Female socialization: how daughters affect their legislator fathers”
- We'll look at the relationship between voting and number of kids.

Saturated model example

- Ebonya Washington (AER) data from *AER* paper “Female socialization: how daughters affect their legislator fathers”
- We'll look at the relationship between voting and number of kids.

```
girls <- foreign::read.dta("girls.dta")
head(girls[, c("name", "totchi", "aauw")])
```

##		name	totchi	aauw
## 1		ABERCROMBIE, NEIL	0	100
## 2		ACKERMAN, GARY L.	3	88
## 3		ADERHOLT, ROBERT B.	0	0
## 4		ALLEN, THOMAS H.	2	100
## 5		ANDREWS, ROBERT E.	2	100
## 6		ARCHER, W.R.	7	0

Linear model

```
summary(lm(aauw ~ totchi, data = girls))
```

```
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)    61.31      1.81    33.81  <2e-16 ***
## totchi         -5.33      0.62   -8.59  <2e-16 ***
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 42 on 1733 degrees of freedom
## (5 observations deleted due to missingness)
## Multiple R-squared:  0.0408, Adjusted R-squared:  0.0403
## F-statistic: 73.8 on 1 and 1733 DF,  p-value: <2e-16
```

Saturated model

```
summary(lm(aauw ~ as.factor(totchi), data = girls))
```

```
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)      56.41      2.76   20.42 < 2e-16 ***
## as.factor(totchi)1    5.45      4.11    1.33  0.1851
## as.factor(totchi)2   -3.80      3.27   -1.16  0.2454
## as.factor(totchi)3  -13.65      3.45   -3.95 8.1e-05 ***
## as.factor(totchi)4  -19.31      4.01   -4.82 1.6e-06 ***
## as.factor(totchi)5  -15.46      4.85   -3.19  0.0015 **
## as.factor(totchi)6  -33.59     10.42   -3.22  0.0013 **
## as.factor(totchi)7  -17.13     11.41   -1.50  0.1336
## as.factor(totchi)8  -55.33     12.28   -4.51 7.0e-06 ***
## as.factor(totchi)9  -50.41     24.08   -2.09  0.0364 *
## as.factor(totchi)10 -53.41     20.90   -2.56  0.0107 *
## as.factor(totchi)12 -56.41     41.53   -1.36  0.1745
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 41 on 1723 degrees of freedom
## (5 observations deleted due to missingness)
## Multiple R-squared:  0.0506, Adjusted R-squared:  0.0446
## F-statistic: 8.36 on 11 and 1723 DF,  p-value: 1.84e-14
```

Saturated model minus the constant

```
summary(lm(aauw ~ as.factor(totchi) - 1, data = girls))
```

```
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## as.factor(totchi)0    56.41      2.76  20.42 <2e-16 ***
## as.factor(totchi)1    61.86      3.05  20.31 <2e-16 ***
## as.factor(totchi)2    52.62      1.75  30.13 <2e-16 ***
## as.factor(totchi)3    42.76      2.07  20.62 <2e-16 ***
## as.factor(totchi)4    37.11      2.90  12.79 <2e-16 ***
## as.factor(totchi)5    40.95      3.99  10.27 <2e-16 ***
## as.factor(totchi)6    22.82     10.05   2.27  0.0233 *
## as.factor(totchi)7    39.29     11.07   3.55  0.0004 ***
## as.factor(totchi)8     1.08     11.96   0.09  0.9278
## as.factor(totchi)9     6.00     23.92   0.25  0.8020
## as.factor(totchi)10    3.00     20.72   0.14  0.8849
## as.factor(totchi)12    0.00     41.43   0.00  1.0000
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 41 on 1723 degrees of freedom
## (5 observations deleted due to missingness)
## Multiple R-squared:  0.587, Adjusted R-squared:  0.584
## F-statistic: 204 on 12 and 1723 DF, p-value: <2e-16
```

Compare to within-strata means

- The saturated model makes no assumptions about the between-strata relationships.

Compare to within-strata means

- The saturated model makes no assumptions about the between-strata relationships.
- Just calculates within-strata means:

Compare to within-strata means

- The saturated model makes no assumptions about the between-strata relationships.
- Just calculates within-strata means:

```
c1 <- coef(lm(aauw ~ as.factor(totchi) - 1, data = girls))
c2 <- with(girls, tapply(aauw, totchi, mean, na.rm = TRUE))
rbind(c1, c2)
```

```
##      0  1  2  3  4  5  6  7  8  9 10 12
## c1 56 62 53 43 37 41 23 39 1.1 6  3  0
## c2 56 62 53 43 37 41 23 39 1.1 6  3  0
```

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

Fitted Values and Residuals

- Where do we get our hats?

Fitted Values and Residuals

- Where do we get our hats?

Fitted Values and Residuals

- Where do we get our hats? $\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$

Fitted Values and Residuals

- Where do we get our hats? $\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$
- To answer this, we first need to redefine and generalize some terms from linear regression with one predictor.

Fitted Values and Residuals

- Where do we get our hats? $\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$
- To answer this, we first need to redefine and generalize some terms from linear regression with one predictor.
- Let's change notation to call our second variable Z_i so its a bit clearer.

Fitted Values and Residuals

- Where do we get our hats? $\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$
- To answer this, we first need to redefine and generalize some terms from linear regression with one predictor.
- Let's change notation to call our second variable Z_i so its a bit clearer.
- Fitted values for $i = 1, \dots, n$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i$$

Fitted Values and Residuals

- Where do we get our hats? $\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$
- To answer this, we first need to redefine and generalize some terms from linear regression with one predictor.
- Let's change notation to call our second variable Z_i so its a bit clearer.
- Fitted values for $i = 1, \dots, n$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i$$

- Residuals for $i = 1, \dots, n$:

$$\hat{u}_i = Y_i - \hat{Y}_i$$

Least Squares is Still Least Squares

- How do we estimate $\hat{\beta}_0$, $\hat{\beta}_1$, and $\hat{\beta}_2$?

Least Squares is Still Least Squares

- How do we estimate $\hat{\beta}_0$, $\hat{\beta}_1$, and $\hat{\beta}_2$?
- Minimize the sum of the squared residuals,

$$(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2) = \arg \min_{b_0, b_1, b_2} \sum_{i=1}^n (Y_i - b_0 - b_1 X_i - b_2 Z_i)^2$$

Least Squares is Still Least Squares

- How do we estimate $\hat{\beta}_0$, $\hat{\beta}_1$, and $\hat{\beta}_2$?
- Minimize the sum of the squared residuals,

$$(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2) = \arg \min_{b_0, b_1, b_2} \sum_{i=1}^n (Y_i - b_0 - b_1 X_i - b_2 Z_i)^2$$

Plan is conceptually the same as before

Least Squares is Still Least Squares

- How do we estimate $\hat{\beta}_0$, $\hat{\beta}_1$, and $\hat{\beta}_2$?
- Minimize the sum of the squared residuals,

$$(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2) = \arg \min_{b_0, b_1, b_2} \sum_{i=1}^n (Y_i - b_0 - b_1 X_i - b_2 Z_i)^2$$

Plan is conceptually the same as before

- 1 Take the partial derivatives of S with respect to b_0, b_1, b_2 .

Least Squares is Still Least Squares

- How do we estimate $\hat{\beta}_0$, $\hat{\beta}_1$, and $\hat{\beta}_2$?
- Minimize the sum of the squared residuals,

$$(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2) = \arg \min_{b_0, b_1, b_2} \sum_{i=1}^n (Y_i - b_0 - b_1 X_i - b_2 Z_i)^2$$

Plan is conceptually the same as before

- 1 Take the partial derivatives of S with respect to b_0, b_1, b_2 .
- 2 Set each of the partial derivatives to 0 to obtain the **first order conditions**.

Least Squares is Still Least Squares

- How do we estimate $\hat{\beta}_0$, $\hat{\beta}_1$, and $\hat{\beta}_2$?
- Minimize the sum of the squared residuals,

$$(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2) = \arg \min_{b_0, b_1, b_2} \sum_{i=1}^n (Y_i - b_0 - b_1 X_i - b_2 Z_i)^2$$

Plan is conceptually the same as before

- 1 Take the partial derivatives of S with respect to b_0, b_1, b_2 .
- 2 Set each of the partial derivatives to 0 to obtain the **first order conditions**.
- 3 Substitute $\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$ for b_0, b_1, b_2 and solve for $\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2$ to obtain the OLS estimator.

Take Partial Derivatives

$$(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2) = \arg \min_{b_0, b_1, b_2} \sum_{i=1}^n (Y_i - b_0 - b_1 X_i - b_2 Z_i)^2$$

After some calculus and algebra we can show that:

$$\frac{\partial S}{\partial b_0} = \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i)$$

$$\frac{\partial S}{\partial b_1} = \sum_{i=1}^n x_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i)$$

$$\frac{\partial S}{\partial b_2} = \sum_{i=1}^n z_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i)$$

First Order Conditions

Setting the partial derivatives equal to zero leads to a system of 3 linear equations in 3 unknowns: $\hat{\beta}_0$, $\hat{\beta}_1$ and $\hat{\beta}_2$

$$\frac{\partial S}{\partial b_0} = \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

$$\frac{\partial S}{\partial b_1} = \sum_{i=1}^n x_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

$$\frac{\partial S}{\partial b_2} = \sum_{i=1}^n z_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

When will this linear system have a unique solution?

First Order Conditions

Setting the partial derivatives equal to zero leads to a system of 3 linear equations in 3 unknowns: $\hat{\beta}_0$, $\hat{\beta}_1$ and $\hat{\beta}_2$

$$\frac{\partial S}{\partial b_0} = \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

$$\frac{\partial S}{\partial b_1} = \sum_{i=1}^n x_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

$$\frac{\partial S}{\partial b_2} = \sum_{i=1}^n z_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

When will this linear system have a unique solution?

- More observations than predictors (i.e. $n > 2$)

First Order Conditions

Setting the partial derivatives equal to zero leads to a system of 3 linear equations in 3 unknowns: $\hat{\beta}_0$, $\hat{\beta}_1$ and $\hat{\beta}_2$

$$\frac{\partial S}{\partial b_0} = \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

$$\frac{\partial S}{\partial b_1} = \sum_{i=1}^n x_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

$$\frac{\partial S}{\partial b_2} = \sum_{i=1}^n z_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

When will this linear system have a unique solution?

- More observations than predictors (i.e. $n > 2$)
- x and z are **linearly independent**, i.e.,
 - ▶ neither x nor z is a constant
 - ▶ x is not a linear function of z (or vice versa)

First Order Conditions

Setting the partial derivatives equal to zero leads to a system of 3 linear equations in 3 unknowns: $\hat{\beta}_0$, $\hat{\beta}_1$ and $\hat{\beta}_2$

$$\frac{\partial S}{\partial b_0} = \sum_{i=1}^n (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

$$\frac{\partial S}{\partial b_1} = \sum_{i=1}^n x_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

$$\frac{\partial S}{\partial b_2} = \sum_{i=1}^n z_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i - \hat{\beta}_2 z_i) = 0$$

When will this linear system have a unique solution?

- More observations than predictors (i.e. $n > 2$)
- x and z are **linearly independent**, i.e.,
 - ▶ neither x nor z is a constant
 - ▶ x is not a linear function of z (or vice versa)
- Typically called **no perfect collinearity**

The OLS Estimator

After lots of algebra, the OLS estimator for $(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2)$ can be written as

$$\begin{aligned}\hat{\beta}_0 &= \bar{y} - \hat{\beta}_1 \bar{x} - \hat{\beta}_2 \bar{z} \\ \hat{\beta}_1 &= \frac{\text{Cov}(x, y) \text{Var}(z) - \text{Cov}(z, y) \text{Cov}(x, z)}{\text{Var}(x) \text{Var}(z) - \text{Cov}(x, z)^2} \\ \hat{\beta}_2 &= \frac{\text{Cov}(z, y) \text{Var}(x) - \text{Cov}(x, y) \text{Cov}(z, x)}{\text{Var}(x) \text{Var}(z) - \text{Cov}(x, z)^2}\end{aligned}$$

The OLS Estimator

After lots of algebra, the OLS estimator for $(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2)$ can be written as

$$\begin{aligned}\hat{\beta}_0 &= \bar{y} - \hat{\beta}_1 \bar{x} - \hat{\beta}_2 \bar{z} \\ \hat{\beta}_1 &= \frac{\text{Cov}(x, y) \text{Var}(z) - \text{Cov}(z, y) \text{Cov}(x, z)}{\text{Var}(x) \text{Var}(z) - \text{Cov}(x, z)^2} \\ \hat{\beta}_2 &= \frac{\text{Cov}(z, y) \text{Var}(x) - \text{Cov}(x, y) \text{Cov}(z, x)}{\text{Var}(x) \text{Var}(z) - \text{Cov}(x, z)^2}\end{aligned}$$

For $(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2)$ to be well-defined we need:

$$\text{Var}(x) \text{Var}(z) \neq \text{Cov}(x, z)^2$$

This requirement fails if:

The OLS Estimator

After lots of algebra, the OLS estimator for $(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2)$ can be written as

$$\begin{aligned}\hat{\beta}_0 &= \bar{y} - \hat{\beta}_1 \bar{x} - \hat{\beta}_2 \bar{z} \\ \hat{\beta}_1 &= \frac{\text{Cov}(x, y) \text{Var}(z) - \text{Cov}(z, y) \text{Cov}(x, z)}{\text{Var}(x) \text{Var}(z) - \text{Cov}(x, z)^2} \\ \hat{\beta}_2 &= \frac{\text{Cov}(z, y) \text{Var}(x) - \text{Cov}(x, y) \text{Cov}(z, x)}{\text{Var}(x) \text{Var}(z) - \text{Cov}(x, z)^2}\end{aligned}$$

For $(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2)$ to be well-defined we need:

$$\text{Var}(x) \text{Var}(z) \neq \text{Cov}(x, z)^2$$

This requirement fails if:

- 1 If x or z is a constant ($\Rightarrow \text{Var}(x) \text{Var}(z) = \text{Cov}(x, z) = 0$)

The OLS Estimator

After lots of algebra, the OLS estimator for $(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2)$ can be written as

$$\begin{aligned}\hat{\beta}_0 &= \bar{y} - \hat{\beta}_1 \bar{x} - \hat{\beta}_2 \bar{z} \\ \hat{\beta}_1 &= \frac{\text{Cov}(x, y) \text{Var}(z) - \text{Cov}(z, y) \text{Cov}(x, z)}{\text{Var}(x) \text{Var}(z) - \text{Cov}(x, z)^2} \\ \hat{\beta}_2 &= \frac{\text{Cov}(z, y) \text{Var}(x) - \text{Cov}(x, y) \text{Cov}(z, x)}{\text{Var}(x) \text{Var}(z) - \text{Cov}(x, z)^2}\end{aligned}$$

For $(\hat{\beta}_0, \hat{\beta}_1, \hat{\beta}_2)$ to be well-defined we need:

$$\text{Var}(x) \text{Var}(z) \neq \text{Cov}(x, z)^2$$

This requirement fails if:

- 1 If x or z is a constant ($\Rightarrow \text{Var}(x) \text{Var}(z) = \text{Cov}(x, z) = 0$)
- 2 One explanatory variable is an exact linear function of another ($\Rightarrow \text{Cor}(x, z) = 1 \Rightarrow \text{Var}(x) \text{Var}(z) = \text{Cov}(x, z)^2$)

OLS Assumptions for Unbiasedness

OLS Assumptions for Unbiasedness

- When we have more than one independent variable, we need the following assumptions in order for OLS to be unbiased:

OLS Assumptions for Unbiasedness

- When we have more than one independent variable, we need the following assumptions in order for OLS to be unbiased:

- 1 Linearity

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

OLS Assumptions for Unbiasedness

- When we have more than one independent variable, we need the following assumptions in order for OLS to be unbiased:

- 1 Linearity

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

- 2 Random/iid sample

OLS Assumptions for Unbiasedness

- When we have more than one independent variable, we need the following assumptions in order for OLS to be unbiased:

- 1 Linearity

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

- 2 Random/iid sample

- 3 No perfect collinearity

OLS Assumptions for Unbiasedness

- When we have more than one independent variable, we need the following assumptions in order for OLS to be unbiased:

- ① Linearity

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

- ② Random/iid sample

- ③ No perfect collinearity

- ④ Zero conditional mean error

$$E[u_i | X_i, Z_i] = 0$$

New Assumption

Assumption 3: No perfect collinearity

(1) No explanatory variable is constant in the sample and (2) there are no exactly linear relationships among the explanatory variables.

- Two components

New Assumption

Assumption 3: No perfect collinearity

(1) No explanatory variable is constant in the sample and (2) there are no exactly linear relationships among the explanatory variables.

- Two components
 - ① Both X_i and Z_i have to vary.

New Assumption

Assumption 3: No perfect collinearity

(1) No explanatory variable is constant in the sample and (2) there are no exactly linear relationships among the explanatory variables.

- Two components
 - ① Both X_i and Z_i have to vary.
 - ② Z_i cannot be a deterministic, linear function of X_i .

New Assumption

Assumption 3: No perfect collinearity

(1) No explanatory variable is constant in the sample and (2) there are no exactly linear relationships among the explanatory variables.

- Two components
 - ① Both X_i and Z_i have to vary.
 - ② Z_i cannot be a deterministic, linear function of X_i .
- Part 2 rules out anything of the form:

$$Z_i = a + bX_i$$

New Assumption

Assumption 3: No perfect collinearity

(1) No explanatory variable is constant in the sample and (2) there are no exactly linear relationships among the explanatory variables.

- Two components
 - ① Both X_i and Z_i have to vary.
 - ② Z_i cannot be a deterministic, linear function of X_i .
- Part 2 rules out anything of the form:

$$Z_i = a + bX_i$$

- Notice how this is linear (equation of a line) and there is no error, so it is deterministic.

New Assumption

Assumption 3: No perfect collinearity

(1) No explanatory variable is constant in the sample and (2) there are no exactly linear relationships among the explanatory variables.

- Two components
 - ① Both X_i and Z_i have to vary.
 - ② Z_i cannot be a deterministic, linear function of X_i .
- Part 2 rules out anything of the form:

$$Z_i = a + bX_i$$

- Notice how this is linear (equation of a line) and there is no error, so it is deterministic.
- What's the correlation between Z_i and X_i ? 1!

Perfect Collinearity Example (I)

- Simple example:

Perfect Collinearity Example (I)

- Simple example:
 - ▶ $X_i = 1$ if a country is **not** in Africa and 0 otherwise.

Perfect Collinearity Example (I)

- Simple example:
 - ▶ $X_i = 1$ if a country is **not** in Africa and 0 otherwise.
 - ▶ $Z_i = 1$ if a country **is** in Africa and 0 otherwise.

Perfect Collinearity Example (I)

- Simple example:
 - ▶ $X_i = 1$ if a country is **not** in Africa and 0 otherwise.
 - ▶ $Z_i = 1$ if a country **is** in Africa and 0 otherwise.
- But, clearly we have the following:

$$Z_i = 1 - X_i$$

Perfect Collinearity Example (I)

- Simple example:
 - ▶ $X_i = 1$ if a country is **not** in Africa and 0 otherwise.
 - ▶ $Z_i = 1$ if a country **is** in Africa and 0 otherwise.
- But, clearly we have the following:

$$Z_i = 1 - X_i$$

- These two variables are perfectly collinear.

Perfect Collinearity Example (I)

- Simple example:
 - ▶ $X_i = 1$ if a country is **not** in Africa and 0 otherwise.
 - ▶ $Z_i = 1$ if a country **is** in Africa and 0 otherwise.
- But, clearly we have the following:

$$Z_i = 1 - X_i$$

- These two variables are perfectly collinear.
- What about the following:

Perfect Collinearity Example (I)

- Simple example:
 - ▶ $X_i = 1$ if a country is **not** in Africa and 0 otherwise.
 - ▶ $Z_i = 1$ if a country **is** in Africa and 0 otherwise.
- But, clearly we have the following:

$$Z_i = 1 - X_i$$

- These two variables are perfectly collinear.
- What about the following:
 - ▶ $X_i = \text{income}$

Perfect Collinearity Example (I)

- Simple example:
 - ▶ $X_i = 1$ if a country is **not** in Africa and 0 otherwise.
 - ▶ $Z_i = 1$ if a country **is** in Africa and 0 otherwise.
- But, clearly we have the following:

$$Z_i = 1 - X_i$$

- These two variables are perfectly collinear.
- What about the following:
 - ▶ $X_i = \text{income}$
 - ▶ $Z_i = X_i^2$

Perfect Collinearity Example (I)

- Simple example:
 - ▶ $X_i = 1$ if a country is **not** in Africa and 0 otherwise.
 - ▶ $Z_i = 1$ if a country **is** in Africa and 0 otherwise.
- But, clearly we have the following:

$$Z_i = 1 - X_i$$

- These two variables are perfectly collinear.
- What about the following:
 - ▶ $X_i = \text{income}$
 - ▶ $Z_i = X_i^2$
- Do we have to worry about collinearity here?

Perfect Collinearity Example (I)

- Simple example:
 - ▶ $X_i = 1$ if a country is **not** in Africa and 0 otherwise.
 - ▶ $Z_i = 1$ if a country **is** in Africa and 0 otherwise.
- But, clearly we have the following:

$$Z_i = 1 - X_i$$

- These two variables are perfectly collinear.
- What about the following:
 - ▶ $X_i = \text{income}$
 - ▶ $Z_i = X_i^2$
- Do we have to worry about collinearity here?
- No! Because while Z_i is a deterministic function of X_i , it is not a linear function of X_i .

R and Perfect Collinearity

- R, and all other packages, will drop one of the variables if there is perfect collinearity:

R and Perfect Collinearity

- R, and all other packages, will drop one of the variables if there is perfect collinearity:

R and Perfect Collinearity

- R, and all other packages, will drop one of the variables if there is perfect collinearity:

```
##  
## Coefficients: (1 not defined because of singularities)  
##           Estimate Std. Error t value Pr(>|t|)  
## (Intercept)  8.71638    0.08991  96.941 < 2e-16 ***  
## africa      -1.36119    0.16306  -8.348 4.87e-14 ***  
## nonafrica           NA          NA      NA      NA  
## ---  
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1  
##  
## Residual standard error: 0.9125 on 146 degrees of freedom  
## (15 observations deleted due to missingness)  
## Multiple R-squared:  0.3231, Adjusted R-squared:  0.3184  
## F-statistic: 69.68 on 1 and 146 DF,  p-value: 4.87e-14
```

Perfect Collinearity Example (II)

- Another example:

Perfect Collinearity Example (II)

- Another example:
 - ▶ X_i = mean temperature in Celsius

Perfect Collinearity Example (II)

- Another example:
 - ▶ X_i = mean temperature in Celsius
 - ▶ $Z_i = 1.8X_i + 32$ (mean temperature in Fahrenheit)

Perfect Collinearity Example (II)

- Another example:
 - ▶ X_i = mean temperature in Celsius
 - ▶ $Z_i = 1.8X_i + 32$ (mean temperature in Fahrenheit)

Perfect Collinearity Example (II)

- Another example:

- ▶ X_i = mean temperature in Celsius
- ▶ $Z_i = 1.8X_i + 32$ (mean temperature in Fahrenheit)

```
## (Intercept)      meantemp  meantemp.f
##  10.8454999   -0.1206948              NA
```

OLS Assumptions for Large-sample Inference

OLS Assumptions for Large-sample Inference

For large-sample inference and calculating SEs, we need the two-variable version of the Gauss-Markov assumptions:

OLS Assumptions for Large-sample Inference

For large-sample inference and calculating SEs, we need the two-variable version of the Gauss-Markov assumptions:

① Linearity

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

② Random/iid sample

③ No perfect collinearity

④ Zero conditional mean error

$$E[u_i | X_i, Z_i] = 0$$

OLS Assumptions for Large-sample Inference

For large-sample inference and calculating SEs, we need the two-variable version of the Gauss-Markov assumptions:

① Linearity

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

② Random/iid sample

③ No perfect collinearity

④ Zero conditional mean error

$$E[u_i | X_i, Z_i] = 0$$

⑤ Homoskedasticity

$$\text{var}[u_i | X_i, Z_i] = \sigma_u^2$$

OLS Assumptions for Large-sample Inference

- We have our OLS estimate $\hat{\beta}_1$

OLS Assumptions for Large-sample Inference

- We have our OLS estimate $\hat{\beta}_1$
- We have an estimate of the standard error for that coefficient, $\widehat{SE}[\hat{\beta}_1]$.

OLS Assumptions for Large-sample Inference

- We have our OLS estimate $\hat{\beta}_1$
- We have an estimate of the standard error for that coefficient, $\widehat{SE}[\hat{\beta}_1]$.
- Under assumption 1-5, in large samples, we'll have the following:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim N(0, 1)$$

OLS Assumptions for Large-sample Inference

- We have our OLS estimate $\hat{\beta}_1$
- We have an estimate of the standard error for that coefficient, $\widehat{SE}[\hat{\beta}_1]$.
- Under assumption 1-5, in large samples, we'll have the following:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim N(0, 1)$$

- The same holds for the other coefficient:

$$\frac{\hat{\beta}_2 - \beta_2}{\widehat{SE}[\hat{\beta}_2]} \sim N(0, 1)$$

OLS Assumptions for Large-sample Inference

- We have our OLS estimate $\hat{\beta}_1$
- We have an estimate of the standard error for that coefficient, $\widehat{SE}[\hat{\beta}_1]$.
- Under assumption 1-5, in large samples, we'll have the following:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim N(0, 1)$$

- The same holds for the other coefficient:

$$\frac{\hat{\beta}_2 - \beta_2}{\widehat{SE}[\hat{\beta}_2]} \sim N(0, 1)$$

- Inference is exactly the same in large samples!

OLS Assumptions for Large-sample Inference

- We have our OLS estimate $\hat{\beta}_1$
- We have an estimate of the standard error for that coefficient, $\widehat{SE}[\hat{\beta}_1]$.
- Under assumption 1-5, in large samples, we'll have the following:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim N(0, 1)$$

- The same holds for the other coefficient:

$$\frac{\hat{\beta}_2 - \beta_2}{\widehat{SE}[\hat{\beta}_2]} \sim N(0, 1)$$

- Inference is exactly the same in large samples!
- Hypothesis tests and CIs are good to go

OLS Assumptions for Large-sample Inference

- We have our OLS estimate $\hat{\beta}_1$
- We have an estimate of the standard error for that coefficient, $\widehat{SE}[\hat{\beta}_1]$.
- Under assumption 1-5, in large samples, we'll have the following:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim N(0, 1)$$

- The same holds for the other coefficient:

$$\frac{\hat{\beta}_2 - \beta_2}{\widehat{SE}[\hat{\beta}_2]} \sim N(0, 1)$$

- Inference is exactly the same in large samples!
- Hypothesis tests and CIs are good to go
- The SE's will change, though

Inference with two independent variables in large samples

Inference with two independent variables in large samples

For small-sample inference, we need the Gauss-Markov plus Normal errors:

Inference with two independent variables in large samples

For small-sample inference, we need the Gauss-Markov plus Normal errors:

- 1 Linearity

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

- 2 Random/iid sample

- 3 No perfect collinearity

- 4 Zero conditional mean error

$$E[u_i | X_i, Z_i] = 0$$

- 5 Homoskedasticity

$$\text{var}[u_i | X_i, Z_i] = \sigma_u^2$$

Inference with two independent variables in large samples

For small-sample inference, we need the Gauss-Markov plus Normal errors:

- 1 Linearity

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + u_i$$

- 2 Random/iid sample
- 3 No perfect collinearity
- 4 Zero conditional mean error

$$E[u_i | X_i, Z_i] = 0$$

- 5 Homoskedasticity

$$\text{var}[u_i | X_i, Z_i] = \sigma_u^2$$

- 6 Normal conditional errors

$$u_i \sim N(0, \sigma_u^2)$$

Inference with two independent variables in small samples

- Under assumptions 1-6, we have the following small change to our small- n sampling distribution:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim t_{n-3}$$

Inference with two independent variables in small samples

- Under assumptions 1-6, we have the following small change to our small- n sampling distribution:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim t_{n-3}$$

- The same is true for the other coefficient:

$$\frac{\hat{\beta}_2 - \beta_2}{\widehat{SE}[\hat{\beta}_2]} \sim t_{n-3}$$

Inference with two independent variables in small samples

- Under assumptions 1-6, we have the following small change to our small- n sampling distribution:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim t_{n-3}$$

- The same is true for the other coefficient:

$$\frac{\hat{\beta}_2 - \beta_2}{\widehat{SE}[\hat{\beta}_2]} \sim t_{n-3}$$

- Why $n - 3$?

Inference with two independent variables in small samples

- Under assumptions 1-6, we have the following small change to our small- n sampling distribution:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim t_{n-3}$$

- The same is true for the other coefficient:

$$\frac{\hat{\beta}_2 - \beta_2}{\widehat{SE}[\hat{\beta}_2]} \sim t_{n-3}$$

- Why $n - 3$?
 - ▶ We've estimated another parameter, so we need to take off another degree of freedom.

Inference with two independent variables in small samples

- Under assumptions 1-6, we have the following small change to our small- n sampling distribution:

$$\frac{\hat{\beta}_1 - \beta_1}{\widehat{SE}[\hat{\beta}_1]} \sim t_{n-3}$$

- The same is true for the other coefficient:

$$\frac{\hat{\beta}_2 - \beta_2}{\widehat{SE}[\hat{\beta}_2]} \sim t_{n-3}$$

- Why $n - 3$?
 - ▶ We've estimated another parameter, so we need to take off another degree of freedom.
- $\rightsquigarrow$ small adjustments to the critical values and the t-values for our hypothesis tests and confidence intervals.

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

Another view of OLS with two predictors

- “Partialling out” OLS recipe:

Another view of OLS with two predictors

- “Partialling out” OLS recipe:
 - 1 Run regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

Another view of OLS with two predictors

- “Partialling out” OLS recipe:

- 1 Run regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- 2 Calculate residuals from this regression:

$$\hat{r}_{xz,i} = X_i - \hat{X}_i$$

Another view of OLS with two predictors

- “Partialling out” OLS recipe:

- 1 Run regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- 2 Calculate residuals from this regression:

$$\hat{r}_{xz,i} = X_i - \hat{X}_i$$

- 3 Run a simple regression of Y_i on residuals, $\hat{r}_{xz,i}$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 \hat{r}_{xz,i}$$

Another view of OLS with two predictors

- “Partialling out” OLS recipe:

- 1 Run regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- 2 Calculate residuals from this regression:

$$\hat{r}_{xz,i} = X_i - \hat{X}_i$$

- 3 Run a simple regression of Y_i on residuals, $\hat{r}_{xz,i}$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 \hat{r}_{xz,i}$$

- Estimate of $\hat{\beta}_1$ will be the same as running:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i$$

Another view of OLS with two predictors

- “Partialling out” OLS recipe:

- 1 Run regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- 2 Calculate residuals from this regression:

$$\hat{r}_{xz,i} = X_i - \hat{X}_i$$

- 3 Run a simple regression of Y_i on residuals, $\hat{r}_{xz,i}$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 \hat{r}_{xz,i}$$

- Estimate of $\hat{\beta}_1$ will be the same as running:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i$$

Another view of OLS with two predictors

- “Partialling out” OLS recipe:

- 1 Run regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- 2 Calculate residuals from this regression:

$$\hat{r}_{xz,i} = X_i - \hat{X}_i$$

- 3 Run a simple regression of Y_i on residuals, $\hat{r}_{xz,i}$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 \hat{r}_{xz,i}$$

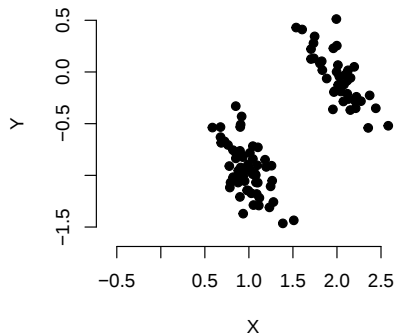
- Estimate of $\hat{\beta}_1$ will be the same as running:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i$$

This is based on the **Frisch-Waugh-Lovell (FWL) theorem**.

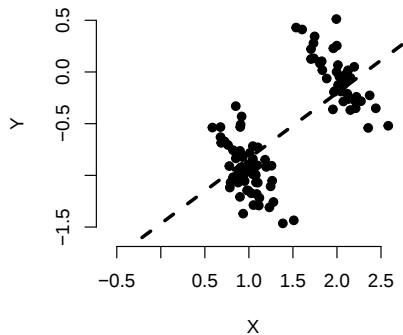
A Visual of Partialling Out

Original



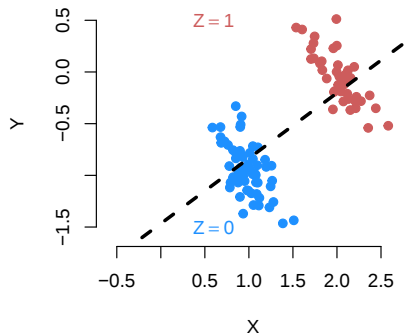
A Visual of Partialling Out

Original



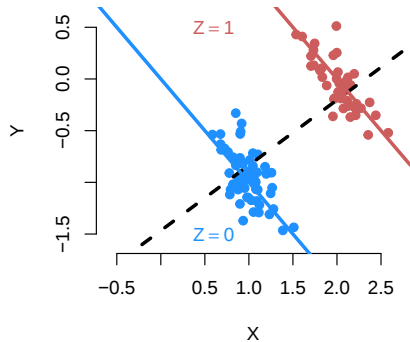
A Visual of Partialling Out

Original



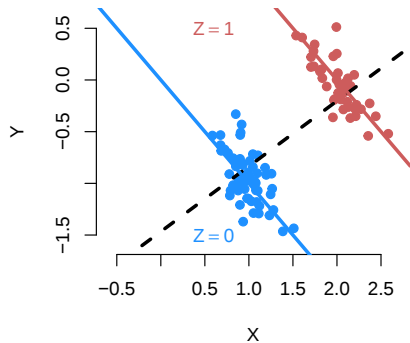
A Visual of Partialling Out

Original

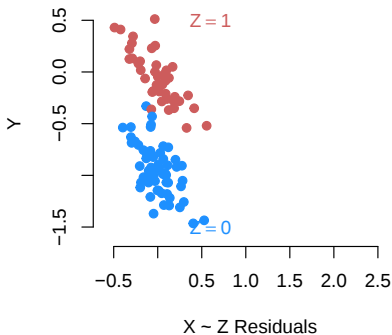


A Visual of Partialling Out

Original

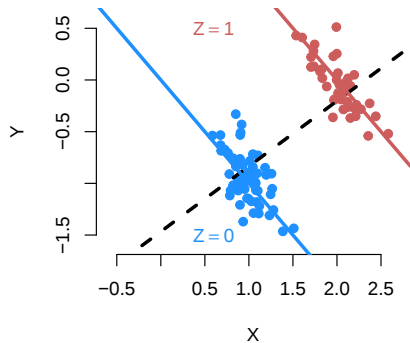


Residualizing X

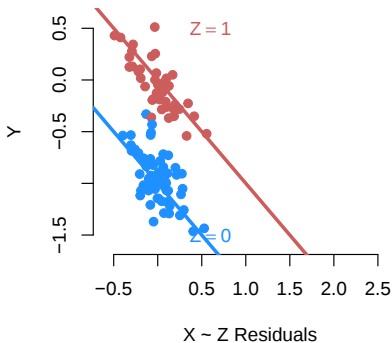


A Visual of Partialling Out

Original

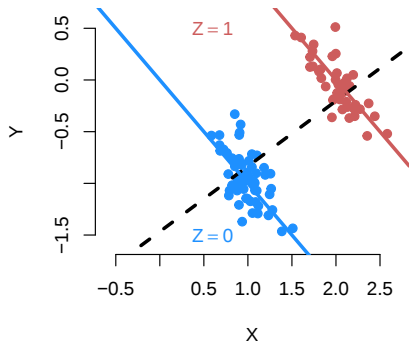


Residualizing X

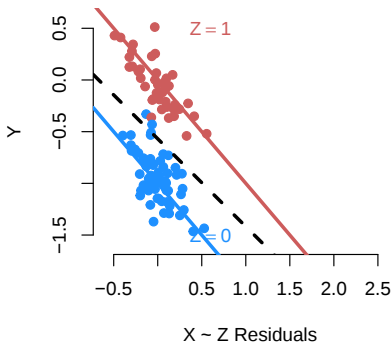


A Visual of Partialling Out

Original

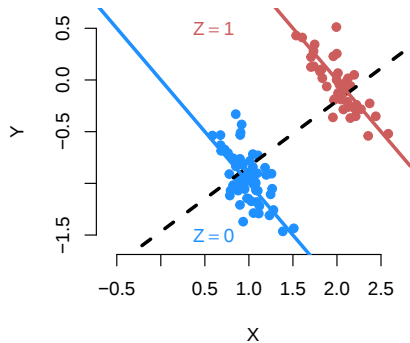


Residualizing X

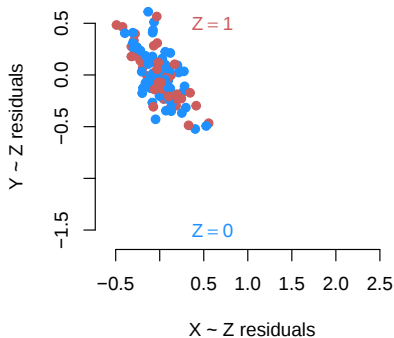


A Visual of Partialling Out

Original

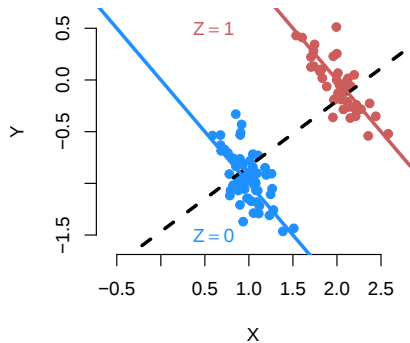


Residualizing X and Y

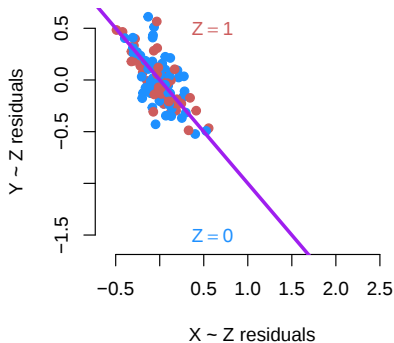


A Visual of Partialling Out

Original

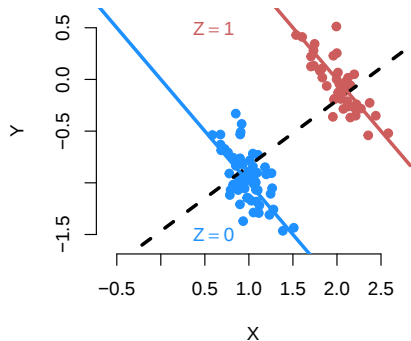


Residualizing X and Y

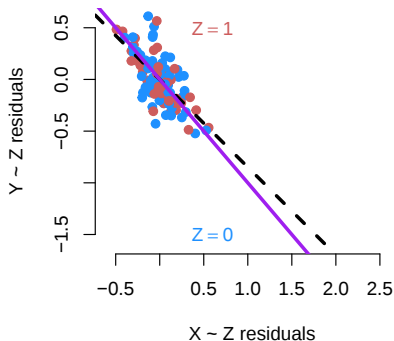


A Visual of Partialling Out

Original



Residualizing X and Y



Origin of the Partial Out Recipe

Assume $Y = \beta_0 + \beta_1 X + \beta_2 Z + u$. Another way to write the OLS estimator is:

$$\hat{\beta}_1 = \frac{\sum_i^n \hat{r}_{xz,i} y_i}{\sum_i^n \hat{r}_{xz,i}^2}$$

where $\hat{r}_{xz,i}$ are the residuals from the regression of X on Z :

$$X = \lambda + \delta Z + r_{xz}$$

In other words, both of these regressions yield identical estimates $\hat{\beta}_1$:

$$y = \hat{\gamma}_0 + \hat{\beta}_1 \hat{r}_{xz} \quad \text{and} \quad y = \hat{\beta}_0 + \hat{\beta}_1 x + \hat{\beta}_2 z$$

Origin of the Partial Out Recipe

Assume $Y = \beta_0 + \beta_1 X + \beta_2 Z + u$. Another way to write the OLS estimator is:

$$\hat{\beta}_1 = \frac{\sum_i^n \hat{r}_{xz,i} y_i}{\sum_i^n \hat{r}_{xz,i}^2}$$

where $\hat{r}_{xz,i}$ are the residuals from the regression of X on Z :

$$X = \lambda + \delta Z + r_{xz}$$

In other words, both of these regressions yield identical estimates $\hat{\beta}_1$:

$$y = \hat{\gamma}_0 + \hat{\beta}_1 \hat{r}_{xz} \quad \text{and} \quad y = \hat{\beta}_0 + \hat{\beta}_1 x + \hat{\beta}_2 z$$

- δ measures the correlation between X and Z .

Origin of the Partial Out Recipe

Assume $Y = \beta_0 + \beta_1 X + \beta_2 Z + u$. Another way to write the OLS estimator is:

$$\hat{\beta}_1 = \frac{\sum_i^n \hat{r}_{xz,i} y_i}{\sum_i^n \hat{r}_{xz,i}^2}$$

where $\hat{r}_{xz,i}$ are the residuals from the regression of X on Z :

$$X = \lambda + \delta Z + r_{xz}$$

In other words, both of these regressions yield identical estimates $\hat{\beta}_1$:

$$y = \hat{\gamma}_0 + \hat{\beta}_1 \hat{r}_{xz} \quad \text{and} \quad y = \hat{\beta}_0 + \hat{\beta}_1 x + \hat{\beta}_2 z$$

- δ measures the correlation between X and Z .
- Residuals $\hat{r}_{xz}$ are the part of X that is uncorrelated with Z . Put differently, $\hat{r}_{xz}$ is X , after the effect of Z on X has been **partialled out** or netted out.

Origin of the Partial Out Recipe

Assume $Y = \beta_0 + \beta_1 X + \beta_2 Z + u$. Another way to write the OLS estimator is:

$$\hat{\beta}_1 = \frac{\sum_i^n \hat{r}_{xz,i} y_i}{\sum_i^n \hat{r}_{xz,i}^2}$$

where $\hat{r}_{xz,i}$ are the residuals from the regression of X on Z :

$$X = \lambda + \delta Z + r_{xz}$$

In other words, both of these regressions yield identical estimates $\hat{\beta}_1$:

$$y = \hat{\gamma}_0 + \hat{\beta}_1 \hat{r}_{xz} \quad \text{and} \quad y = \hat{\beta}_0 + \hat{\beta}_1 x + \hat{\beta}_2 z$$

- δ measures the correlation between X and Z .
- Residuals $\hat{r}_{xz}$ are the part of X that is uncorrelated with Z . Put differently, $\hat{r}_{xz}$ is X , after the effect of Z on X has been **partialled out** or netted out.
- Can use same equation with k explanatory variables; $\hat{r}_{xz}$ will then come from a regression of X on all the other explanatory variables.

Why Should We Care About Partialling Out?

- Won't R just calculate it for us?

Why Should We Care About Partialling Out?

- Won't R just calculate it for us?
- Sure—but the partialling out strategy provides great intuition for what the regression controlling for another variable is doing.

Why Should We Care About Partialling Out?

- Won't R just calculate it for us?
- Sure—but the partialling out strategy provides great intuition for what the regression controlling for another variable is doing.
- This set up will also be the basis of **diagnostic plots** that we will cover in a couple of weeks. It allows us to visualize the **conditional** relationship.

Why Should We Care About Partialling Out?

- Won't R just calculate it for us?
- Sure—but the partialling out strategy provides great intuition for what the regression controlling for another variable is doing.
- This set up will also be the basis of **diagnostic plots** that we will cover in a couple of weeks. It allows us to visualize the **conditional** relationship.
- Finally, it forms the foundation of a number of machine learning strategies including **double machine learning** by breaking down the regression problem.

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

Bootstrapped Sampling Distributions

What if there was a way to replace thinking with computers?

Bootstrapped Sampling Distributions

What if there was a way to replace thinking with computers?

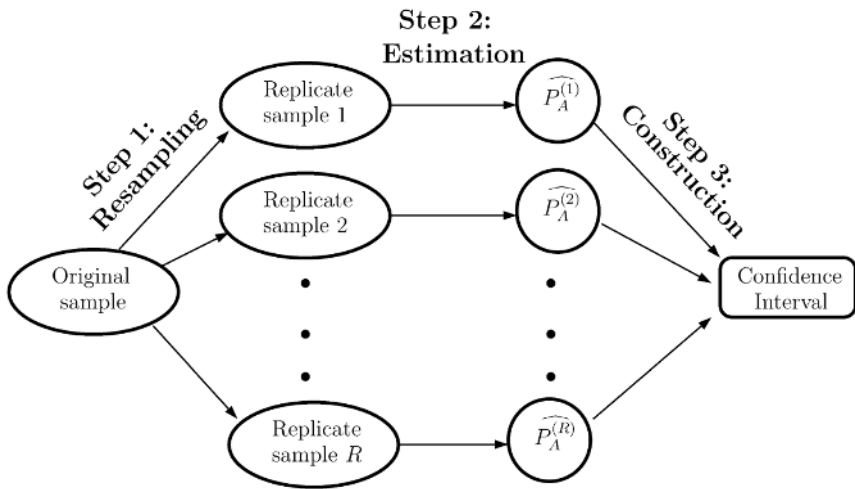
What if there was a way to replacing analytical derivations, which can be hard, with computer simulations which are easy?

Bootstrapped Sampling Distributions

What if there was a way to replace thinking with computers?

What if there was a way to replacing analytical derivations, which can be hard, with computer simulations which are easy?

The **plug-in principle** gives us a way forward.



Source: Salganik (2006)

This works for almost* any estimator

*basically it works when plug-in estimation works

Statistical Science
1988, Vol. 1, No. 1, 54-77

Bootstrap Methods for Standard Errors, Confidence Intervals, and Other Measures of Statistical Accuracy

B. Efron and R. Tibshirani

Efron and Tibshirani (1986), <http://www.jstor.org/stable/2245500>

The Bootstrap

The Bootstrap

Bootstrap: Use the eCDF as a plug-in for the CDF, and **resample** from that. I.e. we are pretending our sample eCDF looks sufficiently close to our true CDF, and so we're sampling from the eCDF as an approximation to repeated sampling from the true CDF. This is called a **resampling** method.

The Bootstrap

Bootstrap: Use the eCDF as a plug-in for the CDF, and **resample** from that. I.e. we are pretending our sample eCDF looks sufficiently close to our true CDF, and so we're sampling from the eCDF as an approximation to repeated sampling from the true CDF. This is called a **resampling** method.

- 1 Take a **with replacement** sample of size n from our sample.

The Bootstrap

Bootstrap: Use the eCDF as a plug-in for the CDF, and **resample** from that. I.e. we are pretending our sample eCDF looks sufficiently close to our true CDF, and so we're sampling from the eCDF as an approximation to repeated sampling from the true CDF. This is called a **resampling** method.

- 1 Take a **with replacement** sample of size n from our sample.
- 2 Calculate our would-be estimate using this **bootstrap sample**.

The Bootstrap

Bootstrap: Use the eCDF as a plug-in for the CDF, and **resample** from that. I.e. we are pretending our sample eCDF looks sufficiently close to our true CDF, and so we're sampling from the eCDF as an approximation to repeated sampling from the true CDF. This is called a **resampling** method.

- 1 Take a **with replacement** sample of size n from our sample.
- 2 Calculate our would-be estimate using this **bootstrap sample**.
- 3 Repeat steps 1 and 2 many (B) times.

The Bootstrap

Bootstrap: Use the eCDF as a plug-in for the CDF, and **resample** from that. I.e. we are pretending our sample eCDF looks sufficiently close to our true CDF, and so we're sampling from the eCDF as an approximation to repeated sampling from the true CDF. This is called a **resampling** method.

- 1 Take a **with replacement** sample of size n from our sample.
- 2 Calculate our would-be estimate using this **bootstrap sample**.
- 3 Repeat steps 1 and 2 many (B) times.
- 4 Using the resulting collection of **bootstrap estimates**, calculate the standard deviation of the **bootstrap distribution** of our estimator. This serves our estimate of the standard deviation of the **sampling distribution**

Example of a Bootstrap

```
samp <- c(9.7, 4.99, 5.9, 3.58, 8.15, 5.54, 4.77, 5.01, 4.89,  
          3.42, 8.63, 7.17, 8.93, 7.5, 4.93, 8.6, 6.26, 7.31,  
          8.96, 3.95)
```

```
obs_mean = mean(samp)
```

```
obs_mean
```

```
## [1] 6.4095
```

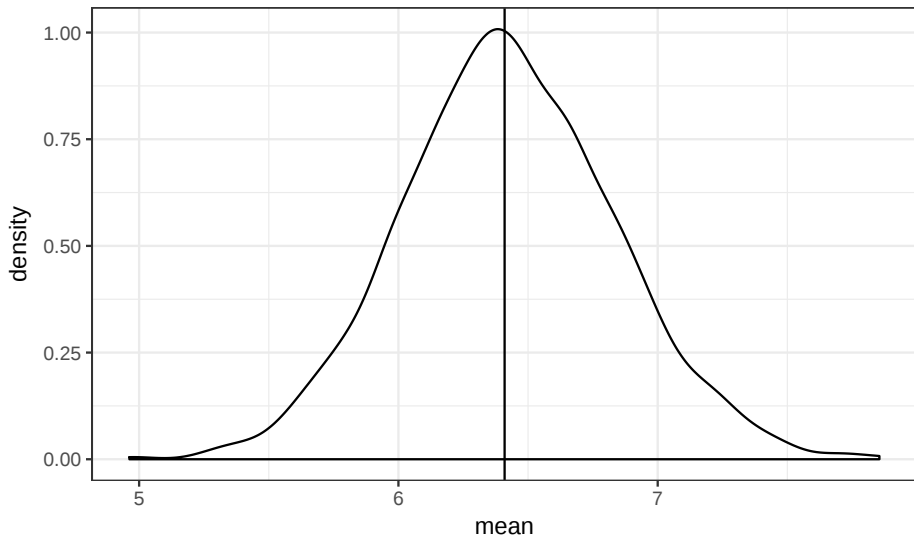
Example of a Bootstrap

```
# resample WITH REPLACEMENT reps times
# recalculate the mean within each bootstrap replicate
boot_samp_dist <- replicate(2000, {
  mean(samp[sample.int(length(samp), replace = TRUE)])
})
```

```
ggplot(tibble(boot_samp_dist = boot_samp_dist),
  aes(x = boot_samp_dist)) +
  geom_density() +
  geom_vline(xintercept = obs_mean) +
  theme_bw() + ggtitle("Bootstrap Sampling Distribution
                        For the Sample Mean") +
  xlab("mean")
```

Example of a Bootstrap

Bootstrap Sampling Distribution
For the Sample Mean



Two ways to calculate bootstrap intervals

Two ways to calculate bootstrap intervals

- 1) Using **normal** approximation intervals, use the estimates from step 4.

$$[\bar{X} - \Phi^{-1}(1 - \alpha/2) * \hat{\sigma}_{\text{boot}}, \bar{X} + \Phi^{-1}(1 - \alpha/2) * \hat{\sigma}_{\text{boot}}]$$

Two ways to calculate bootstrap intervals

- 1) Using **normal** approximation intervals, use the estimates from step 4.

$$[\bar{X} - \Phi^{-1}(1 - \alpha/2) * \hat{\sigma}_{\text{boot}}, \bar{X} + \Phi^{-1}(1 - \alpha/2) * \hat{\sigma}_{\text{boot}}]$$

- Note here that the standard error is just the standard deviation of the bootstrap replicates. There is no square root of n . Why?

Two ways to calculate bootstrap intervals

- 1) Using **normal** approximation intervals, use the estimates from step 4.

$$[\bar{X} - \Phi^{-1}(1 - \alpha/2) * \hat{\sigma}_{\text{boot}}, \bar{X} + \Phi^{-1}(1 - \alpha/2) * \hat{\sigma}_{\text{boot}}]$$

- Note here that the standard error is just the standard deviation of the bootstrap replicates. There is no square root of n . Why?

- 2) **Percentile** method for the CI: Sort B bootstrap estimates from smallest to largest. α interval is constructed as

$$CI_{1-\alpha} = [\alpha/2 * B \text{ sample}, (1 - \alpha/2) * B \text{ sample}]$$

- Percentile method does not rely on normal approximation, and behaves better with small n .

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

Sampling variance for simple linear regression

- Under simple linear regression, we found that the distribution of the slope was the following:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{\sum_{i=1}^n (X_i - \bar{X})^2}$$

Sampling variance for simple linear regression

- Under simple linear regression, we found that the distribution of the slope was the following:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{\sum_{i=1}^n (X_i - \bar{X})^2}$$

- Factors affecting the standard errors (the square root of these sampling variances):

Sampling variance for simple linear regression

- Under simple linear regression, we found that the distribution of the slope was the following:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{\sum_{i=1}^n (X_i - \bar{X})^2}$$

- Factors affecting the standard errors (the square root of these sampling variances):
 - ▶ **The error variance** σ_u^2 (higher conditional variance of Y_i leads to bigger SEs)

Sampling variance for simple linear regression

- Under simple linear regression, we found that the distribution of the slope was the following:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{\sum_{i=1}^n (X_i - \bar{X})^2}$$

- Factors affecting the standard errors (the square root of these sampling variances):
 - ▶ The error variance σ_u^2 (higher conditional variance of Y_i leads to bigger SEs)
 - ▶ The **total variation in X_i** : $\sum_{i=1}^n (X_i - \bar{X})^2$ (lower variation in X_i leads to bigger SEs)

Sampling variation for linear regression with two covariates

- Regression with an additional independent variable:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{(1 - R_1^2) \sum_{i=1}^n (X_i - \bar{X})^2}$$

Sampling variation for linear regression with two covariates

- Regression with an additional independent variable:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{(1 - R_1^2) \sum_{i=1}^n (X_i - \bar{X})^2}$$

- Here, R_1^2 is the R^2 from the regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

Sampling variation for linear regression with two covariates

- Regression with an additional independent variable:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{(1 - R_1^2) \sum_{i=1}^n (X_i - \bar{X})^2}$$

- Here, R_1^2 is the R^2 from the regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- Factors now affecting the standard errors:

Sampling variation for linear regression with two covariates

- Regression with an additional independent variable:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{(1 - R_1^2) \sum_{i=1}^n (X_i - \bar{X})^2}$$

- Here, R_1^2 is the R^2 from the regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- Factors now affecting the standard errors:
 - ▶ The **error variance** (higher conditional variance of Y_i leads to bigger SEs)

Sampling variation for linear regression with two covariates

- Regression with an additional independent variable:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{(1 - R_1^2) \sum_{i=1}^n (X_i - \bar{X})^2}$$

- Here, R_1^2 is the R^2 from the regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- Factors now affecting the standard errors:
 - ▶ The error variance (higher conditional variance of Y_i leads to bigger SEs)
 - ▶ The **total variation of X_i** (lower variation in X_i leads to bigger SEs)

Sampling variation for linear regression with two covariates

- Regression with an additional independent variable:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{(1 - R_1^2) \sum_{i=1}^n (X_i - \bar{X})^2}$$

- Here, R_1^2 is the R^2 from the regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- Factors now affecting the standard errors:
 - ▶ The error variance (higher conditional variance of Y_i leads to bigger SEs)
 - ▶ The total variation of X_i (lower variation in X_i leads to bigger SEs)
 - ▶ The **strength of the linear relationship** between X_i and Z_i (stronger relationships mean higher R_1^2 and thus bigger SEs)

Sampling variation for linear regression with two covariates

- Regression with an additional independent variable:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{(1 - R_1^2) \sum_{i=1}^n (X_i - \bar{X})^2}$$

- Here, R_1^2 is the R^2 from the regression of X_i on Z_i :

$$\hat{X}_i = \hat{\delta}_0 + \hat{\delta}_1 Z_i$$

- Factors now affecting the standard errors:
 - ▶ The error variance (higher conditional variance of Y_i leads to bigger SEs)
 - ▶ The total variation of X_i (lower variation in X_i leads to bigger SEs)
 - ▶ The strength of the linear relationship between X_i and Z_i (stronger relationships mean higher R_1^2 and thus bigger SEs)
- What happens with perfect collinearity? $R_1^2 = 1$ and the variances are infinite.

Multicollinearity

Multicollinearity

Definition

Multicollinearity is defined to be high, but not perfect, correlation between two independent variables in a regression.

Multicollinearity

Definition

Multicollinearity is defined to be high, but not perfect, correlation between two independent variables in a regression.

- With multicollinearity, we'll have $R_1^2 \approx 1$, but not exactly.

Multicollinearity

Definition

Multicollinearity is defined to be high, but not perfect, correlation between two independent variables in a regression.

- With multicollinearity, we'll have $R_1^2 \approx 1$, but not exactly.
- The stronger the relationship between X_i and Z_i , the closer the R_1^2 will be to 1, and the higher the SEs will be:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{(1 - R_1^2) \sum_{i=1}^n (X_i - \bar{X})^2}$$

Multicollinearity

Definition

Multicollinearity is defined to be high, but not perfect, correlation between two independent variables in a regression.

- With multicollinearity, we'll have $R_1^2 \approx 1$, but not exactly.
- The stronger the relationship between X_i and Z_i , the closer the R_1^2 will be to 1, and the higher the SEs will be:

$$\text{var}(\hat{\beta}_1) = \frac{\sigma_u^2}{(1 - R_1^2) \sum_{i=1}^n (X_i - \bar{X})^2}$$

- Given the symmetry, it will also increase $\text{var}(\hat{\beta}_2)$ as well.

Intuition for multicollinearity

- Remember the OLS recipe:

Intuition for multicollinearity

- Remember the OLS recipe:
 - ▶ $\hat{\beta}_1$ from regression of Y_i on $\hat{r}_{xz,i}$

Intuition for multicollinearity

- Remember the OLS recipe:
 - ▶ $\hat{\beta}_1$ from regression of Y_i on $\hat{r}_{xz,i}$
 - ▶ $\hat{r}_{xz,i}$ are the residuals from the regression of X_i on Z_i

Intuition for multicollinearity

- Remember the OLS recipe:
 - ▶ $\hat{\beta}_1$ from regression of Y_i on $\hat{r}_{xz,i}$
 - ▶ $\hat{r}_{xz,i}$ are the residuals from the regression of X_i on Z_i
- Estimated coefficient:

$$\hat{\beta}_1 = \frac{\sum_{i=1}^n \hat{r}_{xz,i} Y_i}{\sum_{i=1}^n \hat{r}_{xz,i}^2}$$

Intuition for multicollinearity

- Remember the OLS recipe:
 - ▶ $\hat{\beta}_1$ from regression of Y_i on $\hat{r}_{xz,i}$
 - ▶ $\hat{r}_{xz,i}$ are the residuals from the regression of X_i on Z_i
- Estimated coefficient:

$$\hat{\beta}_1 = \frac{\sum_{i=1}^n \hat{r}_{xz,i} Y_i}{\sum_{i=1}^n \hat{r}_{xz,i}^2}$$

- When Z_i and X_i have a strong relationship, then the residuals will have low variation

Intuition for multicollinearity

- Remember the OLS recipe:
 - ▶ $\hat{\beta}_1$ from regression of Y_i on $\hat{r}_{xz,i}$
 - ▶ $\hat{r}_{xz,i}$ are the residuals from the regression of X_i on Z_i
- Estimated coefficient:

$$\hat{\beta}_1 = \frac{\sum_{i=1}^n \hat{r}_{xz,i} Y_i}{\sum_{i=1}^n \hat{r}_{xz,i}^2}$$

- When Z_i and X_i have a strong relationship, then the residuals will have low variation
- We explain away a lot of the variation in X_i through Z_i .

Intuition for multicollinearity

- Remember the OLS recipe:
 - ▶ $\hat{\beta}_1$ from regression of Y_i on $\hat{r}_{xz,i}$
 - ▶ $\hat{r}_{xz,i}$ are the residuals from the regression of X_i on Z_i
- Estimated coefficient:

$$\hat{\beta}_1 = \frac{\sum_{i=1}^n \hat{r}_{xz,i} Y_i}{\sum_{i=1}^n \hat{r}_{xz,i}^2}$$

- When Z_i and X_i have a strong relationship, then the residuals will have low variation
- We explain away a lot of the variation in X_i through Z_i .
- Low variation in an independent variable (here, $\hat{r}_{xz,i}$) $\rightsquigarrow$ high SEs

Intuition for multicollinearity

- Remember the OLS recipe:
 - ▶ $\hat{\beta}_1$ from regression of Y_i on $\hat{r}_{xz,i}$
 - ▶ $\hat{r}_{xz,i}$ are the residuals from the regression of X_i on Z_i
- Estimated coefficient:

$$\hat{\beta}_1 = \frac{\sum_{i=1}^n \hat{r}_{xz,i} Y_i}{\sum_{i=1}^n \hat{r}_{xz,i}^2}$$

- When Z_i and X_i have a strong relationship, then the residuals will have low variation
- We explain away a lot of the variation in X_i through Z_i .
- Low variation in an independent variable (here, $\hat{r}_{xz,i}$) $\rightsquigarrow$ high SEs
- Basically, there is less residual variation left in X_i after “partialling out” the effect of Z_i

Effects of multicollinearity

- No effect on the bias of OLS.

Effects of multicollinearity

- No effect on the bias of OLS.
- Only increases the standard errors.

Effects of multicollinearity

- No effect on the bias of OLS.
- Only increases the standard errors.
- Really just a sample size problem:

Effects of multicollinearity

- No effect on the bias of OLS.
- Only increases the standard errors.
- Really just a sample size problem:
 - ▶ If X_i and Z_i are extremely highly correlated, you're going to need a much bigger sample to accurately differentiate between their effects.

Effects of multicollinearity

- No effect on the bias of OLS.
- Only increases the standard errors.
- Really just a sample size problem:
 - ▶ If X_i and Z_i are extremely highly correlated, you're going to need a much bigger sample to accurately differentiate between their effects.

Effects of multicollinearity

- No effect on the bias of OLS.
- Only increases the standard errors.
- Really just a sample size problem:
 - ▶ If X_i and Z_i are extremely highly correlated, you're going to need a much bigger sample to accurately differentiate between their effects.



How Do We Detect Multicollinearity?

- The best practice is to directly compute $\text{Cor}(X_1, X_2)$ before running your regression.

How Do We Detect Multicollinearity?

- The best practice is to directly compute $\text{Cor}(X_1, X_2)$ before running your regression.
- But you might (and probably will) forget to do so. Even then, you can detect multicollinearity from your regression result:

How Do We Detect Multicollinearity?

- The best practice is to directly compute $\text{Cor}(X_1, X_2)$ before running your regression.
- But you might (and probably will) forget to do so. Even then, you can detect multicollinearity from your regression result:
 - ▶ Large changes in the estimated regression coefficients when a predictor variable is added or deleted

How Do We Detect Multicollinearity?

- The best practice is to directly compute $\text{Cor}(X_1, X_2)$ before running your regression.
- But you might (and probably will) forget to do so. Even then, you can detect multicollinearity from your regression result:
 - ▶ Large changes in the estimated regression coefficients when a predictor variable is added or deleted
 - ▶ Lack of statistical significance despite high R^2

How Do We Detect Multicollinearity?

- The best practice is to directly compute $\text{Cor}(X_1, X_2)$ before running your regression.
- But you might (and probably will) forget to do so. Even then, you can detect multicollinearity from your regression result:
 - ▶ Large changes in the estimated regression coefficients when a predictor variable is added or deleted
 - ▶ Lack of statistical significance despite high R^2
 - ▶ Estimated regression coefficients have an opposite sign from predicted

How Do We Detect Multicollinearity?

- The best practice is to directly compute $\text{Cor}(X_1, X_2)$ before running your regression.
- But you might (and probably will) forget to do so. Even then, you can detect multicollinearity from your regression result:
 - ▶ Large changes in the estimated regression coefficients when a predictor variable is added or deleted
 - ▶ Lack of statistical significance despite high R^2
 - ▶ Estimated regression coefficients have an opposite sign from predicted
- A more formal indicator is the **variance inflation factor (VIF)**:

$$VIF(\beta_j) = \frac{1}{1 - R_j^2}$$

which measures how much $V[\hat{\beta}_j | X]$ is inflated compared to a (hypothetical) uncorrelated data. (where R_j^2 is the coefficient of determination from the partialing out equation)

How Do We Detect Multicollinearity?

- The best practice is to directly compute $\text{Cor}(X_1, X_2)$ before running your regression.
- But you might (and probably will) forget to do so. Even then, you can detect multicollinearity from your regression result:
 - ▶ Large changes in the estimated regression coefficients when a predictor variable is added or deleted
 - ▶ Lack of statistical significance despite high R^2
 - ▶ Estimated regression coefficients have an opposite sign from predicted
- A more formal indicator is the **variance inflation factor (VIF)**:

$$VIF(\beta_j) = \frac{1}{1 - R_j^2}$$

which measures how much $V[\hat{\beta}_j | X]$ is inflated compared to a (hypothetical) uncorrelated data. (where R_j^2 is the coefficient of determination from the partialing out equation)

In R, `vif()` in the `car` package.

So How Should I Think about Multicollinearity?

- Multicollinearity does NOT lead to bias; estimates will be unbiased and consistent under Gauss-Markov.

So How Should I Think about Multicollinearity?

- Multicollinearity does NOT lead to bias; estimates will be unbiased and consistent under Gauss-Markov.
- Multicollinearity should in fact be seen as a problem of **micronumerosity**, or “too little data.” You can’t ask the OLS estimator to distinguish the partial effects of X_1 and X_2 if they are essentially the same.

So How Should I Think about Multicollinearity?

- Multicollinearity does NOT lead to bias; estimates will be unbiased and consistent under Gauss-Markov.
- Multicollinearity should in fact be seen as a problem of **micronumerosity**, or “too little data.” You can’t ask the OLS estimator to distinguish the partial effects of X_1 and X_2 if they are essentially the same.
- If X_1 and X_2 are almost the same, why would you want a unique β_1 and a unique β_2 ? Think about how you would interpret that?

So How Should I Think about Multicollinearity?

- Multicollinearity does NOT lead to bias; estimates will be unbiased and consistent under Gauss-Markov.
- Multicollinearity should in fact be seen as a problem of **micronumerosity**, or “too little data.” You can’t ask the OLS estimator to distinguish the partial effects of X_1 and X_2 if they are essentially the same.
- If X_1 and X_2 are almost the same, why would you want a unique β_1 and a unique β_2 ? Think about how you would interpret that?
- Relax, you got way more important things to worry about!

So How Should I Think about Multicollinearity?

- Multicollinearity does NOT lead to bias; estimates will be unbiased and consistent under Gauss-Markov.
- Multicollinearity should in fact be seen as a problem of **micronumerosity**, or “too little data.” You can’t ask the OLS estimator to distinguish the partial effects of X_1 and X_2 if they are essentially the same.
- If X_1 and X_2 are almost the same, why would you want a unique β_1 and a unique β_2 ? Think about how you would interpret that?
- Relax, you got way more important things to worry about!
- If possible, get more data

So How Should I Think about Multicollinearity?

- Multicollinearity does NOT lead to bias; estimates will be unbiased and consistent under Gauss-Markov.
- Multicollinearity should in fact be seen as a problem of **micronumerosity**, or “too little data.” You can’t ask the OLS estimator to distinguish the partial effects of X_1 and X_2 if they are essentially the same.
- If X_1 and X_2 are almost the same, why would you want a unique β_1 and a unique β_2 ? Think about how you would interpret that?
- Relax, you got way more important things to worry about!
- If possible, get more data
- Drop one of the variables, or combine them

So How Should I Think about Multicollinearity?

- Multicollinearity does NOT lead to bias; estimates will be unbiased and consistent under Gauss-Markov.
- Multicollinearity should in fact be seen as a problem of **micronumerosity**, or “too little data.” You can’t ask the OLS estimator to distinguish the partial effects of X_1 and X_2 if they are essentially the same.
- If X_1 and X_2 are almost the same, why would you want a unique β_1 and a unique β_2 ? Think about how you would interpret that?
- Relax, you got way more important things to worry about!
- If possible, get more data
- Drop one of the variables, or combine them
- Or maybe linear regression is not the right tool

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

Why Interaction Terms?

Why Interaction Terms?

- Interaction terms will allow you to let the **slope** on one variable vary as a function of another variable

Why Interaction Terms?

- Interaction terms will allow you to let the **slope** on one variable vary as a function of another variable
- Interaction terms are central in regression analysis to:
 - ▶ Model and test conditional hypothesis (do the returns to education vary by race?)

Why Interaction Terms?

- Interaction terms will allow you to let the **slope** on one variable vary as a function of another variable
- Interaction terms are central in regression analysis to:
 - ▶ Model and test conditional hypothesis (do the returns to education vary by race?)
 - ▶ Make model of the conditional expectation function more realistic by letting coefficients vary across subgroups

Why Interaction Terms?

- Interaction terms will allow you to let the **slope** on one variable vary as a function of another variable
- Interaction terms are central in regression analysis to:
 - ▶ Model and test conditional hypothesis (do the returns to education vary by race?)
 - ▶ Make model of the conditional expectation function more realistic by letting coefficients vary across subgroups
- We can interact:
 - ▶ two or more dummy variables

Why Interaction Terms?

- Interaction terms will allow you to let the **slope** on one variable vary as a function of another variable
- Interaction terms are central in regression analysis to:
 - ▶ Model and test conditional hypothesis (do the returns to education vary by race?)
 - ▶ Make model of the conditional expectation function more realistic by letting coefficients vary across subgroups
- We can interact:
 - ▶ two or more dummy variables
 - ▶ dummy variables and continuous variables

Why Interaction Terms?

- Interaction terms will allow you to let the **slope** on one variable vary as a function of another variable
- Interaction terms are central in regression analysis to:
 - ▶ Model and test conditional hypothesis (do the returns to education vary by race?)
 - ▶ Make model of the conditional expectation function more realistic by letting coefficients vary across subgroups
- We can interact:
 - ▶ two or more dummy variables
 - ▶ dummy variables and continuous variables
 - ▶ two or more continuous variables

Why Interaction Terms?

- Interaction terms will allow you to let the **slope** on one variable vary as a function of another variable
- Interaction terms are central in regression analysis to:
 - ▶ Model and test conditional hypothesis (do the returns to education vary by race?)
 - ▶ Make model of the conditional expectation function more realistic by letting coefficients vary across subgroups
- We can interact:
 - ▶ two or more dummy variables
 - ▶ dummy variables and continuous variables
 - ▶ two or more continuous variables
- Interactions often confuses researchers and mistakes in use and interpretation occur frequently (even in top journals)

Return to the Fish Example

- Data comes from Fish (2002), “Islam and Authoritarianism.”

Return to the Fish Example

- Data comes from Fish (2002), “Islam and Authoritarianism.”
- Basic relationship: does more economic development lead to more democracy?

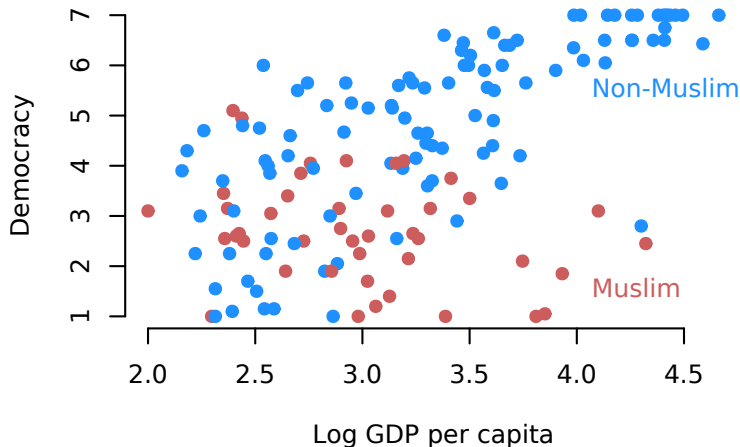
Return to the Fish Example

- Data comes from Fish (2002), “Islam and Authoritarianism.”
- Basic relationship: does more economic development lead to more democracy?
- We measure economic development with log GDP per capita

Return to the Fish Example

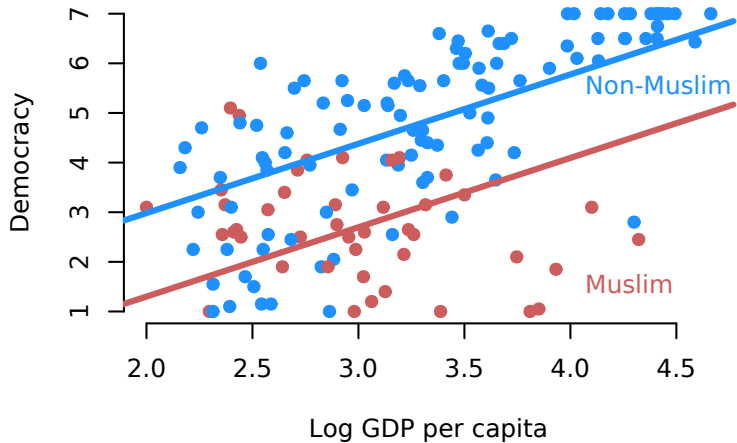
- Data comes from Fish (2002), “Islam and Authoritarianism.”
- Basic relationship: does more economic development lead to more democracy?
- We measure economic development with log GDP per capita
- We measure democracy with a Freedom House score, 1 (less free) to 7 (more free)

Let's See the Data

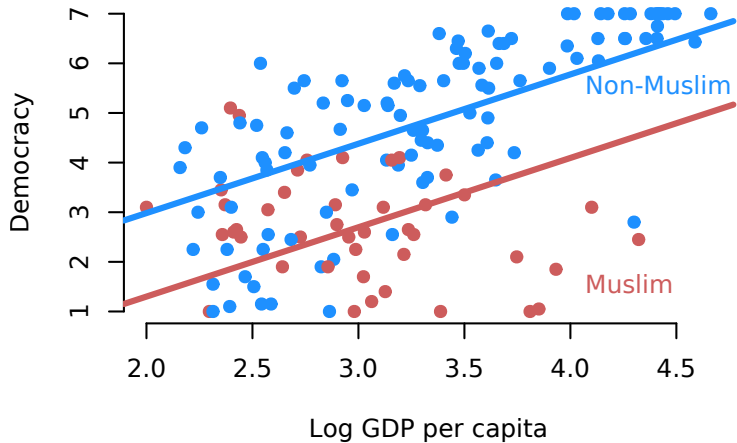


Fish argues that Muslim countries are less likely to be democratic no matter their economic development

Controlling for Religion Additively

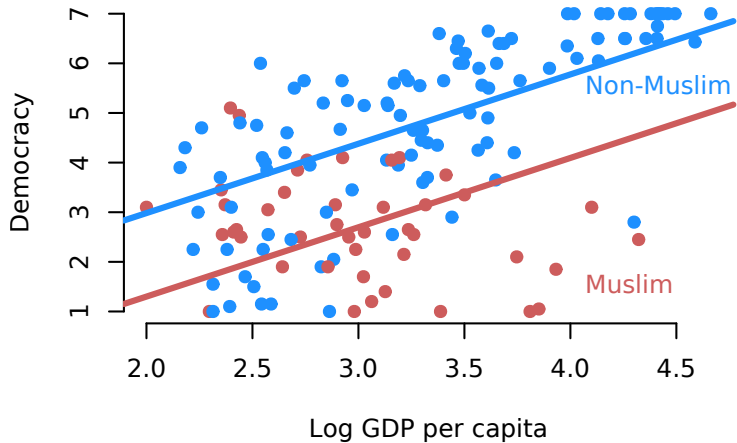


Controlling for Religion Additively



But the regression is a poor fit for Muslim countries

Controlling for Religion Additively



But the regression is a poor fit for Muslim countries

Can we allow for different slopes for each group?

Interactions with a Binary Variable

Interactions with a Binary Variable

- Let Z_i be binary

Interactions with a Binary Variable

- Let Z_i be binary
- In this case, $Z_i = 1$ for the country being Muslim

Interactions with a Binary Variable

- Let Z_i be binary
- In this case, $Z_i = 1$ for the country being Muslim
- We can add another covariate to the baseline model that allows the effect of income to vary by Muslim status.

Interactions with a Binary Variable

- Let Z_i be binary
- In this case, $Z_i = 1$ for the country being Muslim
- We can add another covariate to the baseline model that allows the effect of income to vary by Muslim status.
- This covariate is called an interaction term and it is the product of the two **marginal** variables of interest: $income_i \times muslim_i$

Interactions with a Binary Variable

- Let Z_i be binary
- In this case, $Z_i = 1$ for the country being Muslim
- We can add another covariate to the baseline model that allows the effect of income to vary by Muslim status.
- This covariate is called an interaction term and it is the product of the two **marginal** variables of interest: $income_i \times muslim_i$
- Here is the model with the interaction term:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- How can we interpret this model?

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- How can we interpret this model?
- We can plug in the two possible values of Z_i

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- How can we interpret this model?
- We can plug in the two possible values of Z_i
- When $Z_i = 0$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- How can we interpret this model?
- We can plug in the two possible values of Z_i
- When $Z_i = 0$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- How can we interpret this model?
- We can plug in the two possible values of Z_i
- When $Z_i = 0$:

$$\begin{aligned}\hat{Y}_i &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 \times 0 + \hat{\beta}_3 X_i \times 0\end{aligned}$$

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- How can we interpret this model?
- We can plug in the two possible values of Z_i
- When $Z_i = 0$:

$$\begin{aligned}\hat{Y}_i &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 \times 0 + \hat{\beta}_3 X_i \times 0 \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i\end{aligned}$$

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- How can we interpret this model?
- We can plug in the two possible values of Z_i
- When $Z_i = 0$:

$$\begin{aligned}\hat{Y}_i &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 \times 0 + \hat{\beta}_3 X_i \times 0 \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i\end{aligned}$$

- When $Z_i = 1$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- How can we interpret this model?
- We can plug in the two possible values of Z_i
- When $Z_i = 0$:

$$\begin{aligned}\hat{Y}_i &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 \times 0 + \hat{\beta}_3 X_i \times 0 \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i\end{aligned}$$

- When $Z_i = 1$:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- How can we interpret this model?
- We can plug in the two possible values of Z_i
- When $Z_i = 0$:

$$\begin{aligned}\hat{Y}_i &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 \times 0 + \hat{\beta}_3 X_i \times 0 \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i\end{aligned}$$

- When $Z_i = 1$:

$$\begin{aligned}\hat{Y}_i &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 \times 1 + \hat{\beta}_3 X_i \times 1\end{aligned}$$

Two Lines in One Regression

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

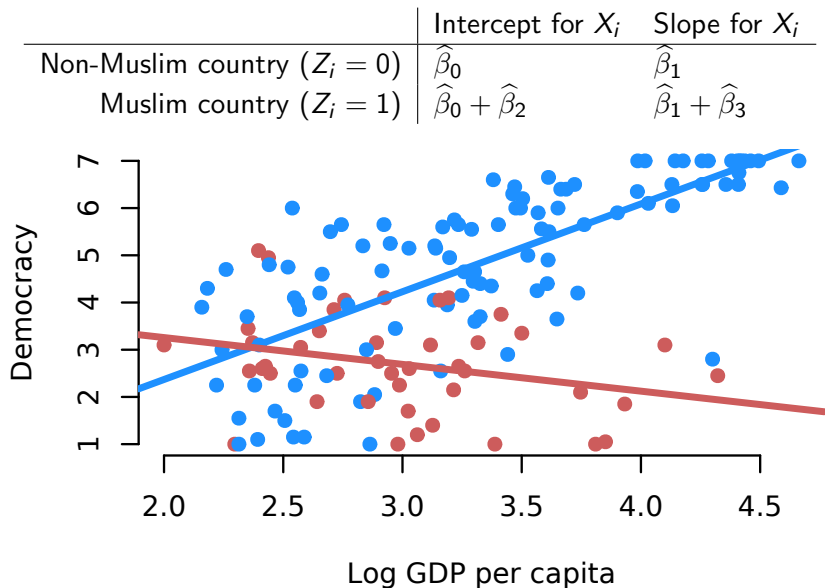
- How can we interpret this model?
- We can plug in the two possible values of Z_i
- When $Z_i = 0$:

$$\begin{aligned}\hat{Y}_i &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 \times 0 + \hat{\beta}_3 X_i \times 0 \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i\end{aligned}$$

- When $Z_i = 1$:

$$\begin{aligned}\hat{Y}_i &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i \\ &= \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 \times 1 + \hat{\beta}_3 X_i \times 1 \\ &= (\hat{\beta}_0 + \hat{\beta}_2) + (\hat{\beta}_1 + \hat{\beta}_3) X_i\end{aligned}$$

Example Interpretation of the Coefficients



General Interpretation of the Coefficients

- $\hat{\beta}_0$: average value of Y_i when both X_i and Z_i are equal to 0

General Interpretation of the Coefficients

- $\hat{\beta}_0$: average value of Y_i when both X_i and Z_i are equal to 0
- $\hat{\beta}_1$: a one-unit change in X_i is associated with a $\hat{\beta}_1$ -unit change in Y_i when $Z_i = 0$

General Interpretation of the Coefficients

- $\hat{\beta}_0$: average value of Y_i when both X_i and Z_i are equal to 0
- $\hat{\beta}_1$: a one-unit change in X_i is associated with a $\hat{\beta}_1$ -unit change in Y_i when $Z_i = 0$
- $\hat{\beta}_2$: average difference in Y_i between $Z_i = 1$ group and $Z_i = 0$ group when $X_i = 0$

General Interpretation of the Coefficients

- $\hat{\beta}_0$: average value of Y_i when both X_i and Z_i are equal to 0
- $\hat{\beta}_1$: a one-unit change in X_i is associated with a $\hat{\beta}_1$ -unit change in Y_i when $Z_i = 0$
- $\hat{\beta}_2$: average difference in Y_i between $Z_i = 1$ group and $Z_i = 0$ group when $X_i = 0$
- $\hat{\beta}_3$: change in the effect of X_i on Y_i between $Z_i = 1$ group and $Z_i = 0$

Lower Order Terms

Lower Order Terms

- Principle of Marginality: Always include the **marginal effects** (sometimes called the **lower order terms**)

Lower Order Terms

- Principle of Marginality: Always include the **marginal effects** (sometimes called the **lower order terms**)
- Imagine we omitted the lower order term for muslim:

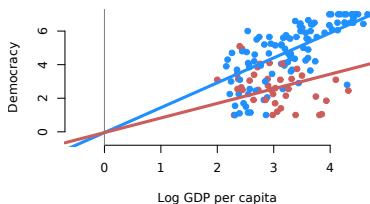
$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + 0 \times Z_i + \hat{\beta}_3 X_i Z_i$$

Omitting Lower Order Terms

Omitting Lower Order Terms

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + 0 \times Z_i + \hat{\beta}_3 X_i Z_i$$

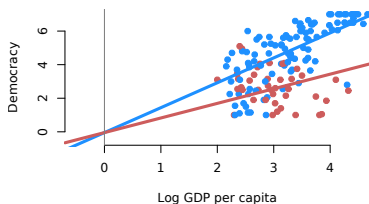
	Intercept for X_i	Slope for X_i
Non-Muslim country ($Z_i = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Muslim country ($Z_i = 1$)	$\hat{\beta}_0 + 0$	$\hat{\beta}_1 + \hat{\beta}_3$



Omitting Lower Order Terms

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + 0 \times Z_i + \hat{\beta}_3 X_i Z_i$$

	Intercept for X_i	Slope for X_i
Non-Muslim country ($Z_i = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Muslim country ($Z_i = 1$)	$\hat{\beta}_0 + 0$	$\hat{\beta}_1 + \hat{\beta}_3$

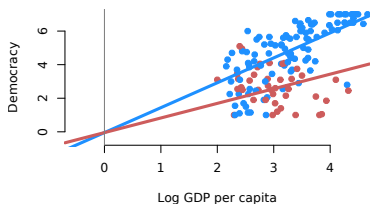


- Model assumption: no difference between Muslim and non-Muslim countries when income is 0

Omitting Lower Order Terms

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + 0 \times Z_i + \hat{\beta}_3 X_i Z_i$$

	Intercept for X_i	Slope for X_i
Non-Muslim country ($Z_i = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Muslim country ($Z_i = 1$)	$\hat{\beta}_0 + 0$	$\hat{\beta}_1 + \hat{\beta}_3$

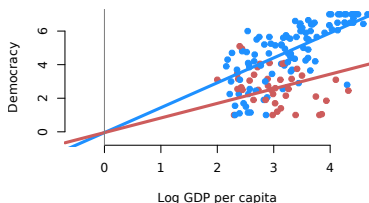


- Model assumption: no difference between Muslim and non-Muslim countries when income is 0
- Distorts slope estimates.

Omitting Lower Order Terms

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + 0 \times Z_i + \hat{\beta}_3 X_i Z_i$$

	Intercept for X_i	Slope for X_i
Non-Muslim country ($Z_i = 0$)	$\hat{\beta}_0$	$\hat{\beta}_1$
Muslim country ($Z_i = 1$)	$\hat{\beta}_0 + 0$	$\hat{\beta}_1 + \hat{\beta}_3$



- Model assumption: no difference between Muslim and non-Muslim countries when income is 0
- Distorts slope estimates.
- Very rarely justified, but for some reason, people keep doing it.

Interactions with Two Continuous Variables

- Now let Z_i be continuous

Interactions with Two Continuous Variables

- Now let Z_i be continuous
- Z_i is the percent growth in GDP per capita from 1975 to 1998

Interactions with Two Continuous Variables

- Now let Z_i be continuous
- Z_i is the percent growth in GDP per capita from 1975 to 1998
- Is the effect of economic development for rapidly developing countries higher or lower than for stagnant economies?

Interactions with Two Continuous Variables

- Now let Z_i be continuous
- Z_i is the percent growth in GDP per capita from 1975 to 1998
- Is the effect of economic development for rapidly developing countries higher or lower than for stagnant economies?
- We can still define the interaction:

$$income_i \times growth_i$$

Interactions with Two Continuous Variables

- Now let Z_i be continuous
- Z_i is the percent growth in GDP per capita from 1975 to 1998
- Is the effect of economic development for rapidly developing countries higher or lower than for stagnant economies?
- We can still define the interaction:

$$income_i \times growth_i$$

- And include it in the regression:

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

Interpretation

- With a continuous Z_i , we can have more than two values that it can take on:

	Intercept for X_i	Slope for X_i
$Z_i = 0$	$\hat{\beta}_0$	$\hat{\beta}_1$

Interpretation

- With a continuous Z_i , we can have more than two values that it can take on:

	Intercept for X_i	Slope for X_i
$Z_i = 0$	$\hat{\beta}_0$	$\hat{\beta}_1$

Interpretation

- With a continuous Z_i , we can have more than two values that it can take on:

	Intercept for X_i	Slope for X_i
$Z_i = 0$	$\hat{\beta}_0$	$\hat{\beta}_1$
$Z_i = 0.5$	$\hat{\beta}_0 + \hat{\beta}_2 \times 0.5$	$\hat{\beta}_1 + \hat{\beta}_3 \times 0.5$

Interpretation

- With a continuous Z_i , we can have more than two values that it can take on:

	Intercept for X_i	Slope for X_i
$Z_i = 0$	$\hat{\beta}_0$	$\hat{\beta}_1$
$Z_i = 0.5$	$\hat{\beta}_0 + \hat{\beta}_2 \times 0.5$	$\hat{\beta}_1 + \hat{\beta}_3 \times 0.5$
$Z_i = 1$	$\hat{\beta}_0 + \hat{\beta}_2 \times 1$	$\hat{\beta}_1 + \hat{\beta}_3 \times 1$

Interpretation

- With a continuous Z_i , we can have more than two values that it can take on:

	Intercept for X_i	Slope for X_i
$Z_i = 0$	$\hat{\beta}_0$	$\hat{\beta}_1$
$Z_i = 0.5$	$\hat{\beta}_0 + \hat{\beta}_2 \times 0.5$	$\hat{\beta}_1 + \hat{\beta}_3 \times 0.5$
$Z_i = 1$	$\hat{\beta}_0 + \hat{\beta}_2 \times 1$	$\hat{\beta}_1 + \hat{\beta}_3 \times 1$
$Z_i = 5$	$\hat{\beta}_0 + \hat{\beta}_2 \times 5$	$\hat{\beta}_1 + \hat{\beta}_3 \times 5$

General Interpretation

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- The coefficient $\hat{\beta}_1$ measures how the predicted outcome varies in X_i when $Z_i = 0$.

General Interpretation

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- The coefficient $\hat{\beta}_1$ measures how the predicted outcome varies in X_i when $Z_i = 0$.
- The coefficient $\hat{\beta}_2$ measures how the predicted outcome varies in Z_i when $X_i = 0$

General Interpretation

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- The coefficient $\hat{\beta}_1$ measures how the predicted outcome varies in X_i when $Z_i = 0$.
- The coefficient $\hat{\beta}_2$ measures how the predicted outcome varies in Z_i when $X_i = 0$
- The coefficient $\hat{\beta}_3$ is the change in the effect of X_i given a one-unit change in Z_i :

$$\frac{\partial E[Y_i | X_i, Z_i]}{\partial X_i} = \beta_1 + \beta_3 Z_i$$

General Interpretation

$$\hat{Y}_i = \hat{\beta}_0 + \hat{\beta}_1 X_i + \hat{\beta}_2 Z_i + \hat{\beta}_3 X_i Z_i$$

- The coefficient $\hat{\beta}_1$ measures how the predicted outcome varies in X_i when $Z_i = 0$.
- The coefficient $\hat{\beta}_2$ measures how the predicted outcome varies in Z_i when $X_i = 0$
- The coefficient $\hat{\beta}_3$ is the change in the effect of X_i given a one-unit change in Z_i :

$$\frac{\partial E[Y_i | X_i, Z_i]}{\partial X_i} = \beta_1 + \beta_3 Z_i$$

- The coefficient $\hat{\beta}_3$ is the change in the effect of Z_i given a one-unit change in X_i :

$$\frac{\partial E[Y_i | X_i, Z_i]}{\partial Z_i} = \beta_2 + \beta_3 X_i$$

Additional Assumptions

Additional Assumptions

Interaction effects are particularly susceptible to model dependence. We are making two assumptions for the estimated effects to be meaningful:

Additional Assumptions

Interaction effects are particularly susceptible to model dependence. We are making two assumptions for the estimated effects to be meaningful:

- 1 Linearity of the interaction effect

Additional Assumptions

Interaction effects are particularly susceptible to model dependence. We are making two assumptions for the estimated effects to be meaningful:

- ① Linearity of the interaction effect
- ② Common support (variation in X throughout the range of Z)

Additional Assumptions

Interaction effects are particularly susceptible to model dependence. We are making two assumptions for the estimated effects to be meaningful:

- 1 Linearity of the interaction effect
- 2 Common support (variation in X throughout the range of Z)

We will talk about checking these assumptions in a few weeks.

PA How Much Should We Trust Estimates from Multiplicative Interaction Models? Simple Tools to Improve Empirical Practice

Jens Hainmueller¹, Jonathan Mummolo² and Yiqing Xu³

¹Professor of Political Science, Stanford University; Department of Political Science, Stanford University, CA
jhainmu@stanford.edu

²Associate Professor of Political Science, Princeton University; Department of Political Science, Princeton University, NJ
jmmummolo@princeton.edu

³Associate Professor of Political Science, Stanford University; Department of Political Science, Stanford University, CA
yiqingxu@stanford.edu

Abstract

Multiplicative interaction models are widely used in social science to examine whether the relationship between an outcome and an independent variable changes with a moderating variable. Current statistical practice tends to overlook two important problems. First, these models assume a linear interaction effect that changes at a constant rate with the moderator. Second, estimates of the conditional effects of the independent variable can be misleading if there is a lack of common support in the moderator. Replicating 46 empirical effects from 22 recent publications in five top political science journals, we find that these common assumptions often fail in practice, suggesting that a large portion of findings across all political science subfields based on interaction models are fragile and model dependent. We propose a checklist of simple diagnostic tests to assess the validity of these assumptions and offer flexible estimation strategies that allow for nonlinear interactions, nonadditive and/or general, nonadditive extrapolation. These statistical routines are available in both R and Stata.

Keywords: interaction effect, linear regression, non-linear regression, interaction models, marginal effects

Example: Common Support

Chapman 2009 analysis

example and reanalysis from Hainmueller, Mummolo, Xu 2019

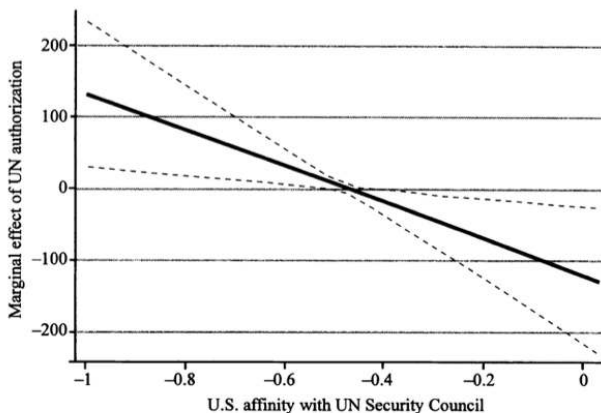
“The interaction term shows a strong negative and statically significant coefficient; suggesting that when UN authorization occurs and the interaction term is ‘switched on’ positive movement in the similarity score (towards more similar) reduces rallies. Rallies with UN authorization are only larger than average when the pivotal member is ideologically distant from the United States. This provides strong support for the informational rationale for IO legitimacy. . .

Clearly, the effect of authorization on rallies decreases as similarity increases: foreign policy actions that receive authorization from a less conservative institution receive similar rallies to those that do not receive authorization from an IO.”

Example: Common Support

Chapman 2009 analysis

example and reanalysis from Hainmueller, Mummolo, Xu 2019

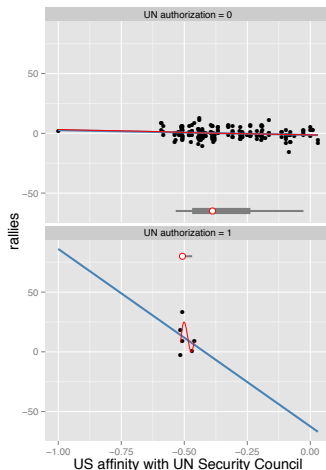


Note: Dashed lines give 95 percent confidence interval.

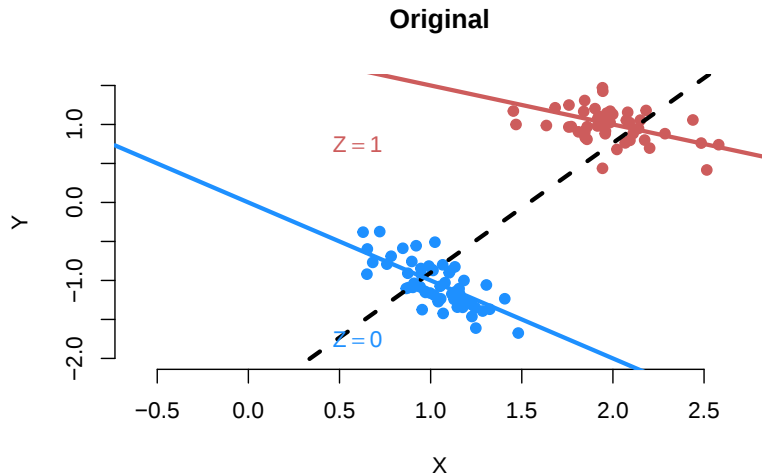
Example: Common Support

Chapman 2009 analysis

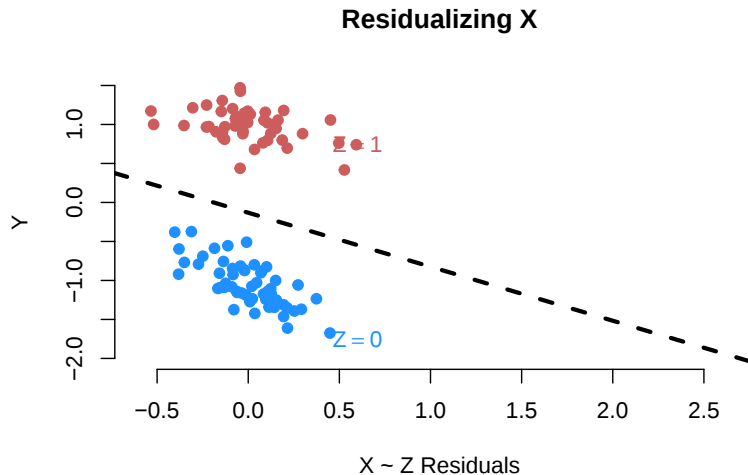
example and reanalysis from Hainmueller, Mummolo, Xu 2019



What Happens Without Interactions

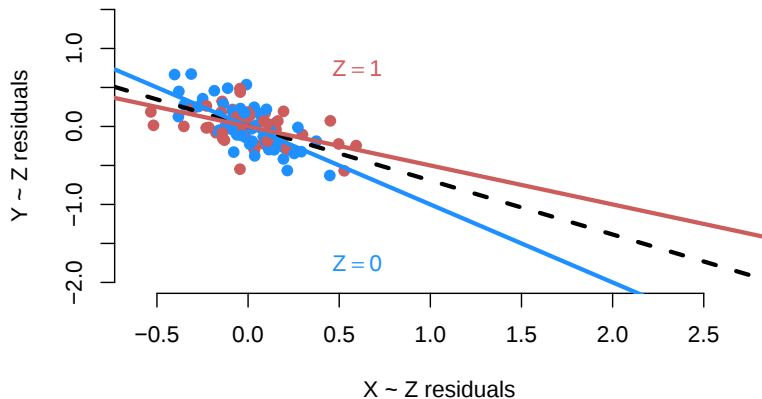


What Happens Without Interactions

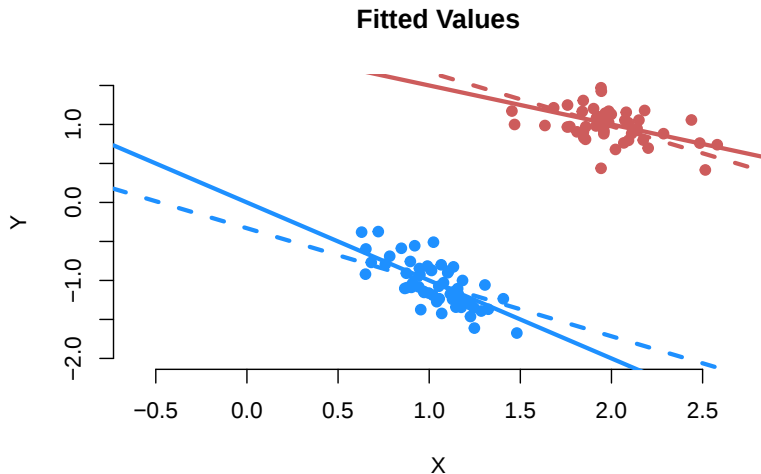


What Happens Without Interactions

Residualizing X and Y



What Happens Without Interactions



Summary for Interactions

Summary for Interactions

- Do not omit lower order terms (unless you have a strong theory that tells you so) because this usually imposes unrealistic restrictions.

Summary for Interactions

- Do not omit lower order terms (unless you have a strong theory that tells you so) because this usually imposes unrealistic restrictions.
- Do not interpret the coefficients on the lower terms as marginal effects (they give the marginal effect only for the case where the other variable is equal to zero)

Summary for Interactions

- Do not omit lower order terms (unless you have a strong theory that tells you so) because this usually imposes unrealistic restrictions.
- Do not interpret the coefficients on the lower terms as marginal effects (they give the marginal effect only for the case where the other variable is equal to zero)
- Produce tables or figures that summarize the conditional marginal effects of the variable of interest at plausible different levels of the other variable; use correct formula to compute variance for these conditional effects (sum of coefficients)

Summary for Interactions

- Do not omit lower order terms (unless you have a strong theory that tells you so) because this usually imposes unrealistic restrictions.
- Do not interpret the coefficients on the lower terms as marginal effects (they give the marginal effect only for the case where the other variable is equal to zero)
- Produce tables or figures that summarize the conditional marginal effects of the variable of interest at plausible different levels of the other variable; use correct formula to compute variance for these conditional effects (sum of coefficients)
- In simple cases the p-value on the interaction term can be used as a test against the null of no interaction, but significant tests for the lower order terms rarely make sense.

Summary for Interactions

- Do not omit lower order terms (unless you have a strong theory that tells you so) because this usually imposes unrealistic restrictions.
- Do not interpret the coefficients on the lower terms as marginal effects (they give the marginal effect only for the case where the other variable is equal to zero)
- Produce tables or figures that summarize the conditional marginal effects of the variable of interest at plausible different levels of the other variable; use correct formula to compute variance for these conditional effects (sum of coefficients)
- In simple cases the p-value on the interaction term can be used as a test against the null of no interaction, but significant tests for the lower order terms rarely make sense.

Further Reading: Brambor, Clark, and Golder. 2006. Understanding Interaction Models: Improving Empirical Analyses. *Political Analysis*.

Hainmueller, Mummolo, Xu. 2019. How Much Should We Trust Estimates from Multiplicative Interaction Models? Simple Tools to Improve Empirical Practice. *Political Analysis*.

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

Polynomial Terms

Polynomial Terms

- Polynomial terms are a special case of the continuous variable interactions.

Polynomial Terms

- Polynomial terms are a special case of the continuous variable interactions.
- For example, when $X_1 = X_2$ in the previous interaction model, we get a quadratic:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_1 X_2 + u$$

$$Y = \beta_0 + (\beta_1 + \beta_2) X_1 + \beta_3 X_1 X_1 + u$$

$$Y = \beta_0 + \tilde{\beta}_1 X_1 + \tilde{\beta}_2 X_1^2 + u$$

- This is called a **second order polynomial** in X_1

Polynomial Terms

- Polynomial terms are a special case of the continuous variable interactions.
- For example, when $X_1 = X_2$ in the previous interaction model, we get a quadratic:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_1 X_2 + u$$

$$Y = \beta_0 + (\beta_1 + \beta_2) X_1 + \beta_3 X_1 X_1 + u$$

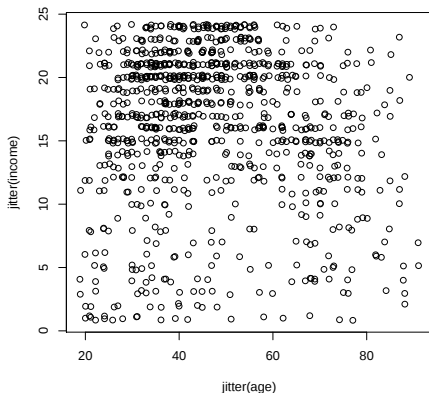
$$Y = \beta_0 + \tilde{\beta}_1 X_1 + \tilde{\beta}_2 X_1^2 + u$$

- This is called a **second order polynomial** in X_1
- A **third order polynomial** is given by:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_1^2 + \beta_3 X_1^3 + u$$

Polynomial Example: Income and Age

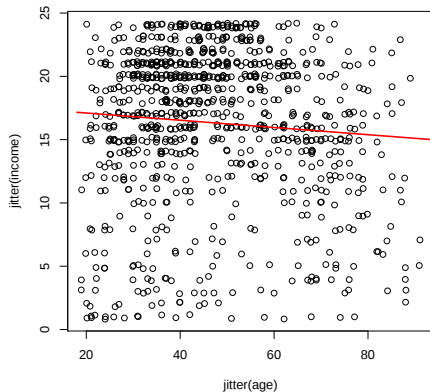
- Let's look at data from the U.S. and examine the relationship between **Y: income** and **X: age**



Polynomial Example: Income and Age

- Let's look at data from the U.S. and examine the relationship between **Y: income** and **X: age**
- We see that a simple linear specification does not fit the data very well:

$$Y = \beta_0 + \beta_1 X_1 + u$$



Polynomial Example: Income and Age

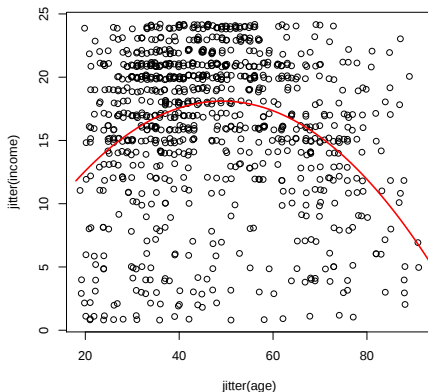
- Let's look at data from the U.S. and examine the relationship between **Y: income** and **X: age**

- We see that a simple linear specification does not fit the data very well:

$$Y = \beta_0 + \beta_1 X_1 + u$$

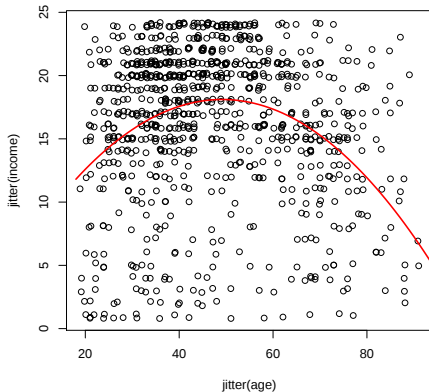
- A second order polynomial in age fits the data a lot better:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_1^2 + u$$



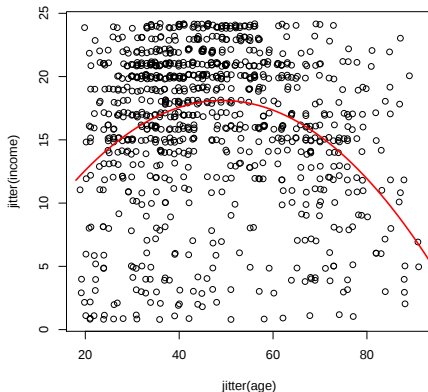
Polynomial Example: Income and Age

- $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_1^2 + u$



Polynomial Example: Income and Age

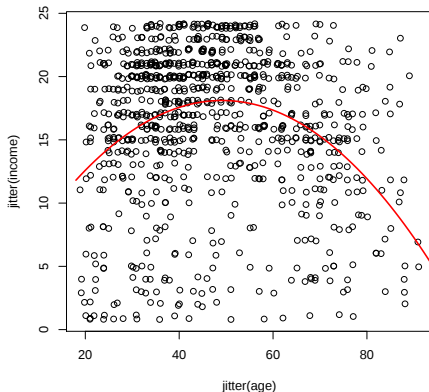
- $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_1^2 + u$
- Is β_1 the marginal effect of age on income?



Polynomial Example: Income and Age

- $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_1^2 + u$
- Is β_1 the marginal effect of age on income?
- No! The marginal effect of age depends on the level of age:
$$\frac{\partial Y}{\partial X_1} = \hat{\beta}_1 + 2\hat{\beta}_2 X_1$$

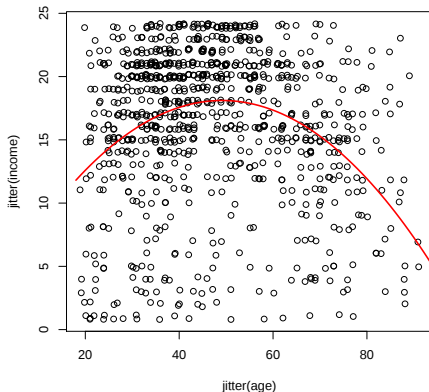
Here the effect of age changes monotonically from positive to negative with income.



Polynomial Example: Income and Age

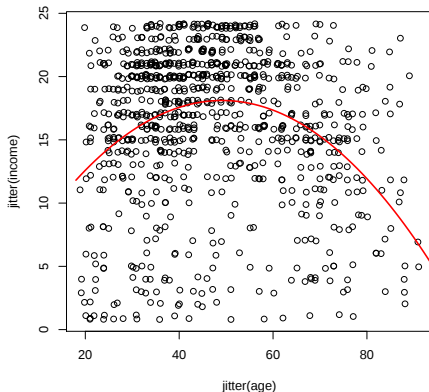
- $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_1^2 + u$
- Is β_1 the marginal effect of age on income?
- No! The marginal effect of age depends on the level of age:
$$\frac{\partial Y}{\partial X_1} = \hat{\beta}_1 + 2\hat{\beta}_2 X_1$$

Here the effect of age changes monotonically from positive to negative with income.
- If $\beta_2 > 0$ we get a U-shape, and if $\beta_2 < 0$ we get an inverted U-shape.

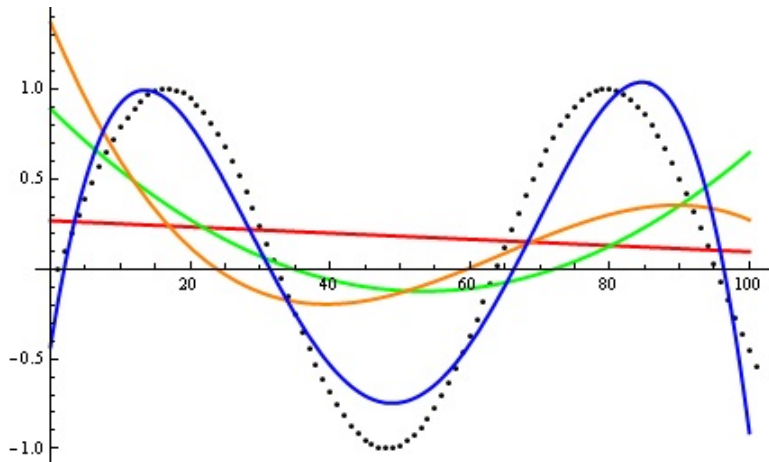


Polynomial Example: Income and Age

- $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_1^2 + u$
- Is β_1 the marginal effect of age on income?
- No! The marginal effect of age depends on the level of age:
 $\frac{\partial Y}{\partial X_1} = \hat{\beta}_1 + 2\hat{\beta}_2 X_1$
Here the effect of age changes monotonically from positive to negative with income.
- If $\beta_2 > 0$ we get a U-shape, and if $\beta_2 < 0$ we get an inverted U-shape.
- Maximum/Minimum occurs at $|\frac{\beta_1}{2\beta_2}|$. Here turning point is at $X_1 = 50$.



Higher Order Polynomials



Approximating data generated with a sine function. Red line is a first degree polynomial, green line is second degree, orange line is third degree and blue is fourth degree

Complex Parameter Interpretation

We can mix and match these model specifications

Complex Parameter Interpretation

We can mix and match these model specifications

- Interactions with higher order terms

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + \beta_3 X_i^2 + \beta_4 X_i * Z_i + \beta_5 X_i^2 * Z_i + u$$

Complex Parameter Interpretation

We can mix and match these model specifications

- Interactions with higher order terms

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + \beta_3 X_i^2 + \beta_4 X_i * Z_i + \beta_5 X_i^2 * Z_i + u$$

- World is your oyster!

Complex Parameter Interpretation

We can mix and match these model specifications

- Interactions with higher order terms

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + \beta_3 X_i^2 + \beta_4 X_i * Z_i + \beta_5 X_i^2 * Z_i + u$$

- World is your oyster!
- But interpretation gets difficult.

Complex Parameter Interpretation

We can mix and match these model specifications

- Interactions with higher order terms

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + \beta_3 X_i^2 + \beta_4 X_i * Z_i + \beta_5 X_i^2 * Z_i + u$$

- World is your oyster!
- But interpretation gets difficult.
- We can take **partial derivatives** as a way to get a sense of the shape of the estimated CEF.

Complex Parameter Interpretation

We can mix and match these model specifications

- Interactions with higher order terms

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + \beta_3 X_i^2 + \beta_4 X_i * Z_i + \beta_5 X_i^2 * Z_i + u$$

- World is your oyster!
- But interpretation gets difficult.
- We can take **partial derivatives** as a way to get a sense of the shape of the estimated CEF.
- An easier approach can be to make predictions and then average over the data

$$\frac{1}{n} \sum_i^n \hat{E}[Y|X = x_i + 1, Z = z_i] - \hat{E}[Y|X = x_i, Z = z_i]$$

Complex Parameter Interpretation

We can mix and match these model specifications

- Interactions with higher order terms

$$Y_i = \beta_0 + \beta_1 X_i + \beta_2 Z_i + \beta_3 X_i^2 + \beta_4 X_i * Z_i + \beta_5 X_i^2 * Z_i + u$$

- World is your oyster!
- But interpretation gets difficult.
- We can take **partial derivatives** as a way to get a sense of the shape of the estimated CEF.
- An easier approach can be to make predictions and then average over the data

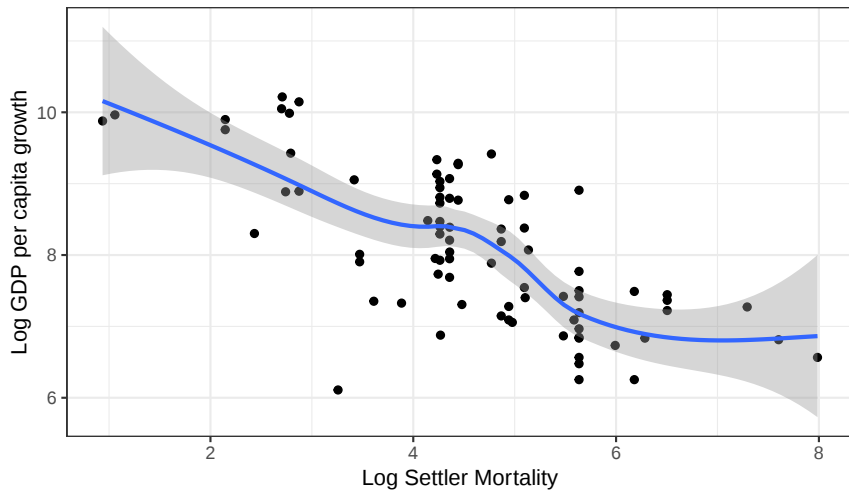
$$\frac{1}{n} \sum_i^n \hat{E}[Y|X = x_i + 1, Z = z_i] - \hat{E}[Y|X = x_i, Z = z_i]$$

- Getting uncertainty on this quantity is a great use case for the bootstrap!

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

- 1 Core Concepts: Why Add a Variable?
 - Two Examples
- 2 How to Add a Variable
 - Adding a Binary Variable
 - Adding a Continuous Covariate
 - Dummy Variables
- 3 Estimation and inference for Two Variable Regression
 - Estimation and Inference
 - Partialling out
 - Bootstrap
 - Multicollinearity
- 4 Interaction Terms
- 5 Non-Linearity
 - Polynomials
 - LOESS
 - Log Transformations
 - Fun With Logs

So what is ggplot2 doing?



LOESS

LOESS

- We can combine the nonparametric kernel method idea of using only **local** data with a **parametric** model

LOESS

- We can combine the nonparametric kernel method idea of using only **local** data with a **parametric** model
- Idea: fit a linear regression within each band

LOESS

- We can combine the nonparametric kernel method idea of using only **local** data with a **parametric** model
- Idea: fit a linear regression within each band
- Locally weighted scatterplot smoothing (**LOWESS** or **LOESS**):

LOESS

- We can combine the nonparametric kernel method idea of using only **local** data with a **parametric** model
- Idea: fit a linear regression within each band
- Locally weighted scatterplot smoothing (**LOWESS** or **LOESS**):
 - 1 Pick a subset of the data that falls in the interval $[x - h, x + h]$

LOESS

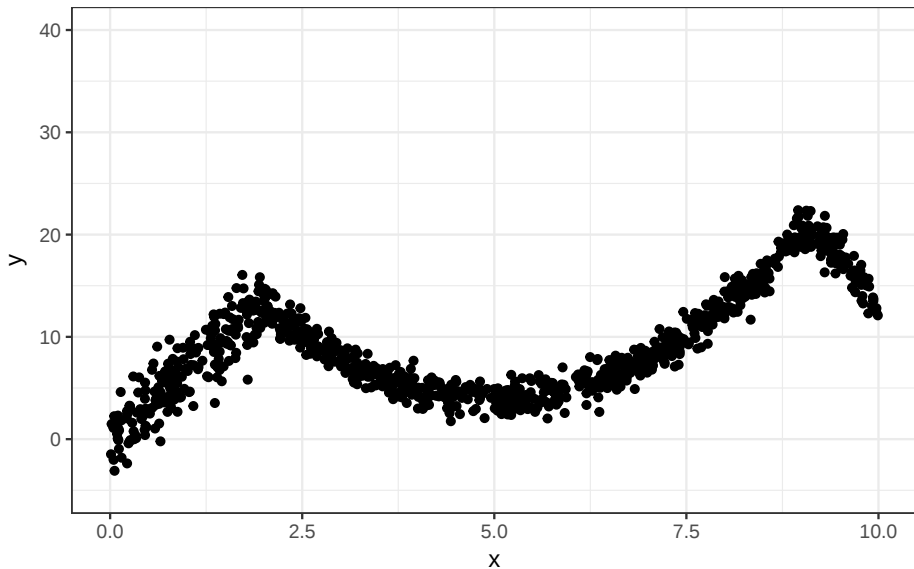
- We can combine the nonparametric kernel method idea of using only **local** data with a **parametric** model
- Idea: fit a linear regression within each band
- Locally weighted scatterplot smoothing (**LOWESS** or **LOESS**):
 - 1 Pick a subset of the data that falls in the interval $[x - h, x + h]$
 - 2 Fit a line to this subset of the data (= **local linear regression**), weighting the points by their distance to x using a kernel function

LOESS

- We can combine the nonparametric kernel method idea of using only **local** data with a **parametric** model
- Idea: fit a linear regression within each band
- Locally weighted scatterplot smoothing (**LOWESS** or **LOESS**):
 - ① Pick a subset of the data that falls in the interval $[x - h, x + h]$
 - ② Fit a line to this subset of the data (= **local linear regression**), weighting the points by their distance to x using a kernel function
 - ③ Use the fitted regression line to predict the expected value of $E[Y|X = x_0]$

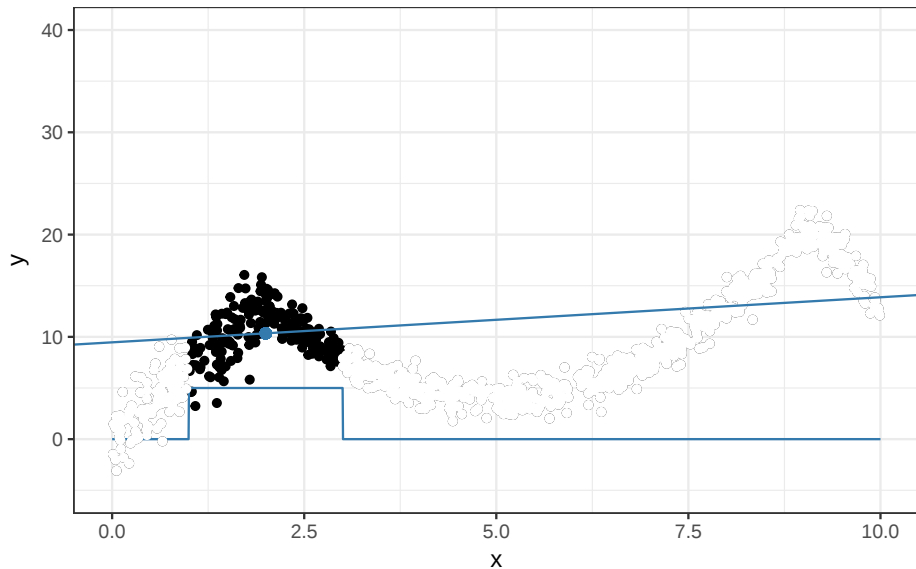
LOESS Example

Uniform Kernel Regression Estimation



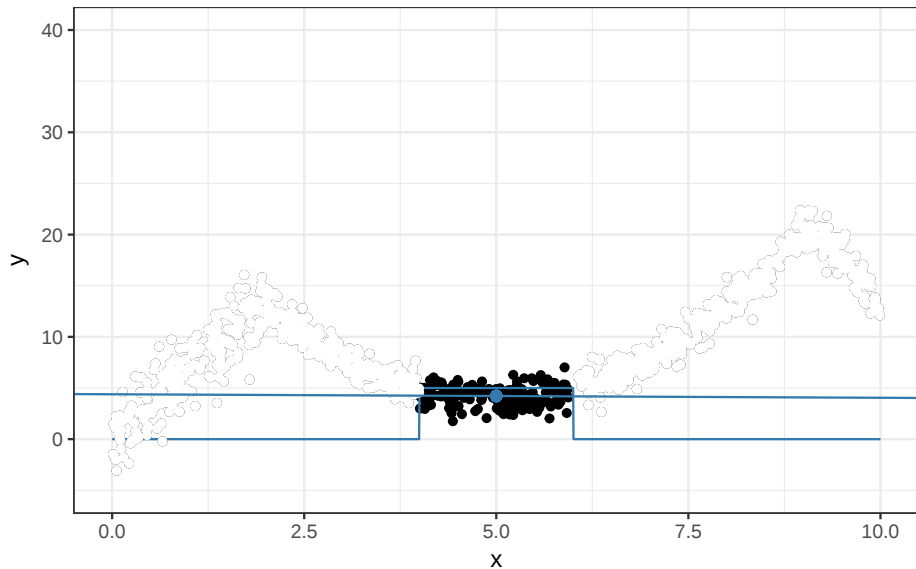
LOESS Example

Uniform Kernel Regression Estimation



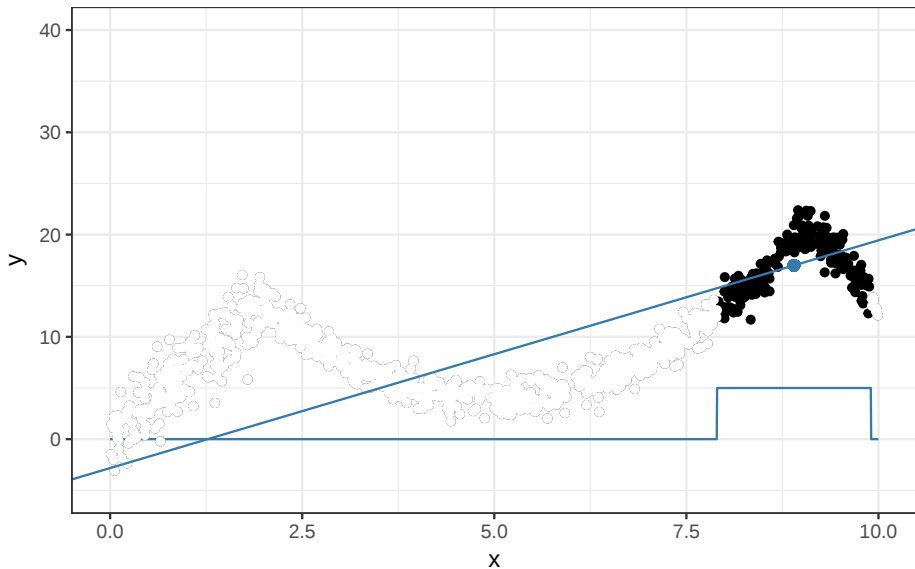
LOESS Example

Uniform Kernel Regression Estimation



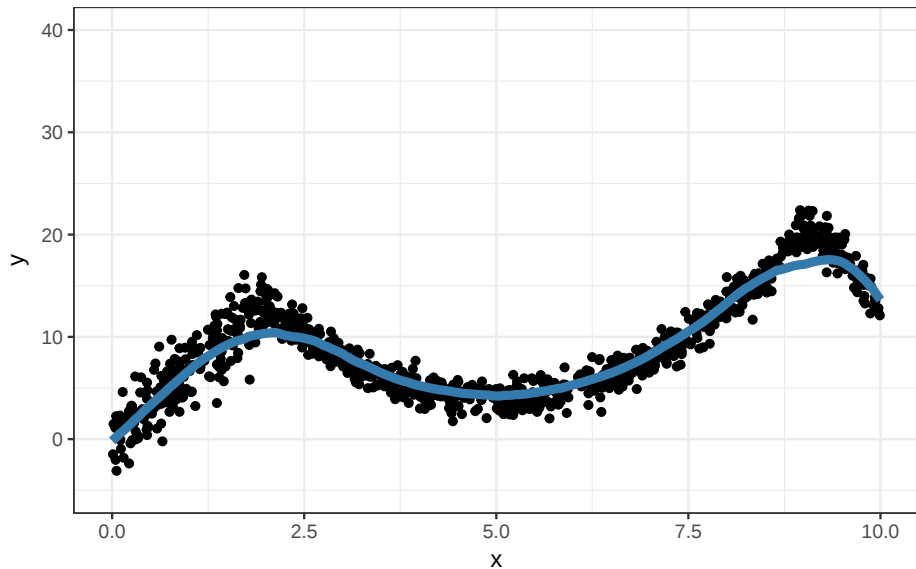
LOESS Example

Uniform Kernel Regression Estimation



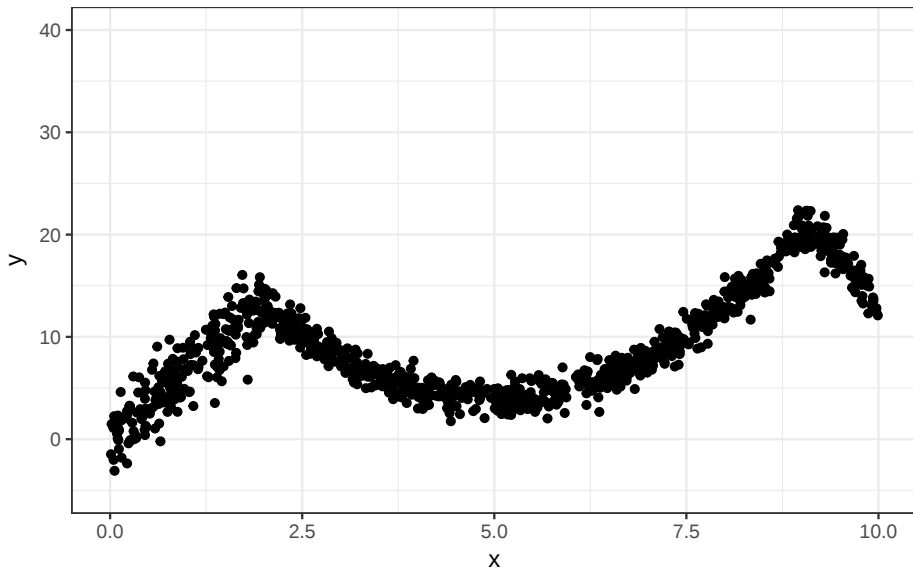
LOESS Example

Uniform Kernel Regression Estimation



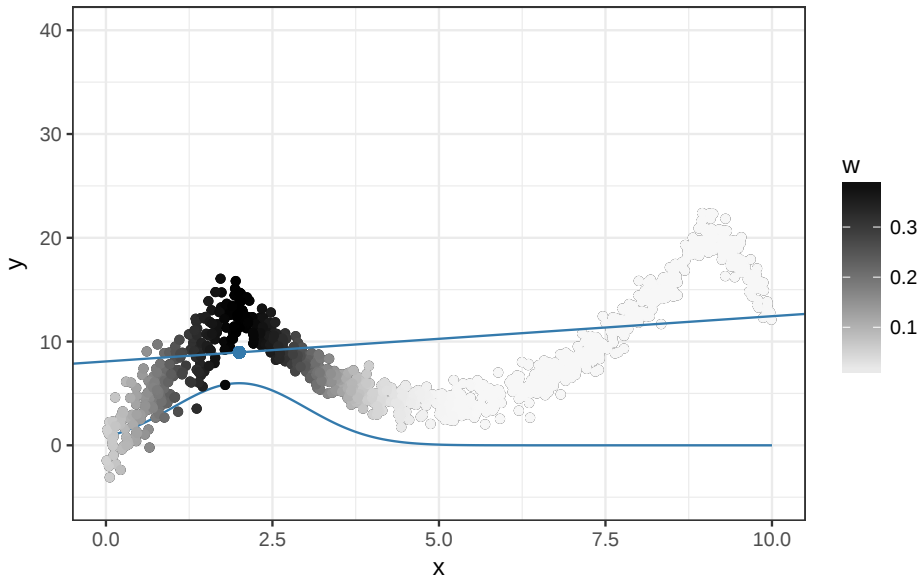
LOESS Example

Gaussian Kernel Regression Estimation



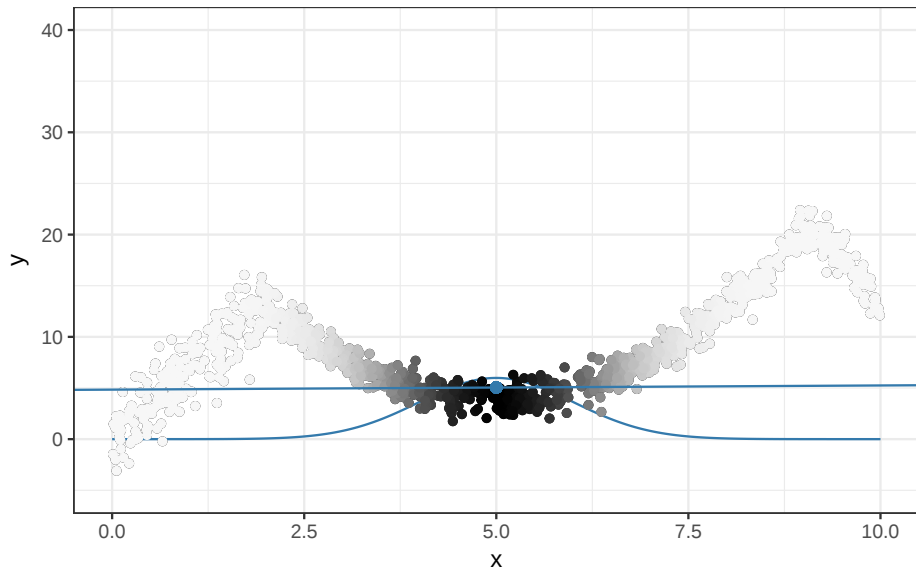
LOESS Example

Gaussian Kernel Regression Estimation



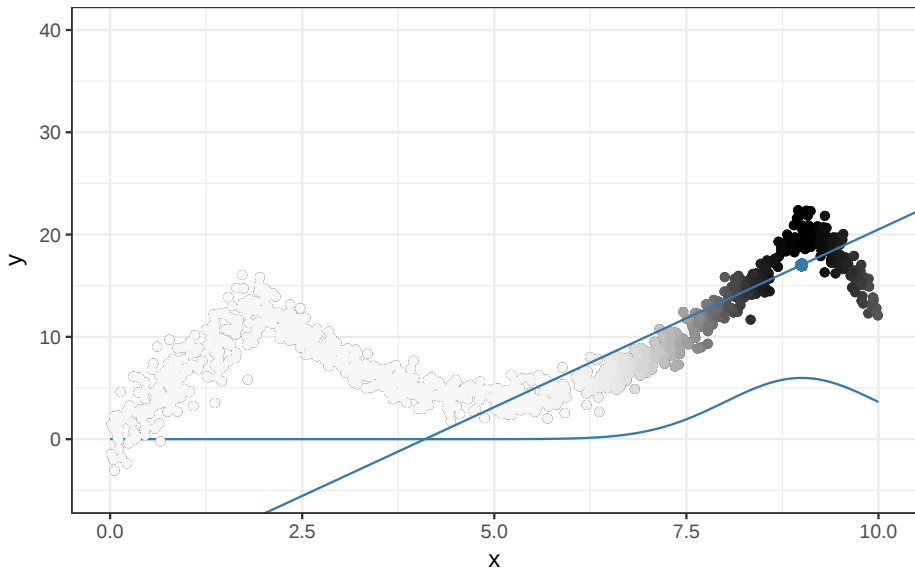
LOESS Example

Gaussian Kernel Regression Estimation



LOESS Example

Gaussian Kernel Regression Estimation



Non-linear CEFs

Non-linear CEFs

- When we say that CEFs are linear with regression, we mean **linear in parameters** but by including transformations of our variables we can make non-linear shapes of pre-specified functional forms.

Non-linear CEFs

- When we say that CEFs are linear with regression, we mean **linear in parameters** but by including transformations of our variables we can make non-linear shapes of pre-specified functional forms.
- Many of these **non-linear transformations** are made by creating multiple variables out of a single X and so will have to wait for future weeks.

Non-linear CEFs

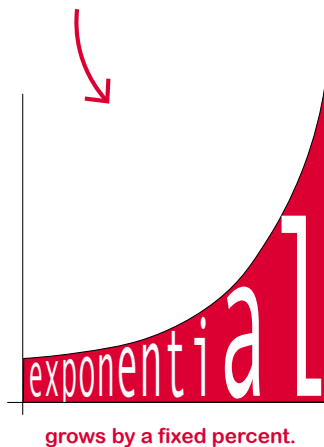
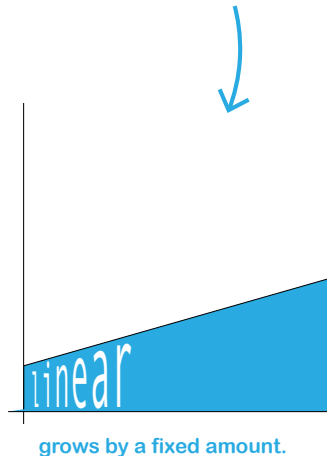
- When we say that CEFs are linear with regression, we mean **linear in parameters** but by including transformations of our variables we can make non-linear shapes of pre-specified functional forms.
- Many of these **non-linear transformations** are made by creating multiple variables out of a single X and so will have to wait for future weeks.
- The function $\log(\cdot)$ is one common transformation that has only one parameter.

Non-linear CEFs

- When we say that CEFs are linear with regression, we mean **linear in parameters** but by including transformations of our variables we can make non-linear shapes of pre-specified functional forms.
- Many of these **non-linear transformations** are made by creating multiple variables out of a single X and so will have to wait for future weeks.
- The function $\log(\cdot)$ is one common transformation that has only one parameter.
- This is particularly useful for **positive** and **right-skewed** variables.

Why does everyone keep logging stuff??

Logs **linearize** **exponential** growth.



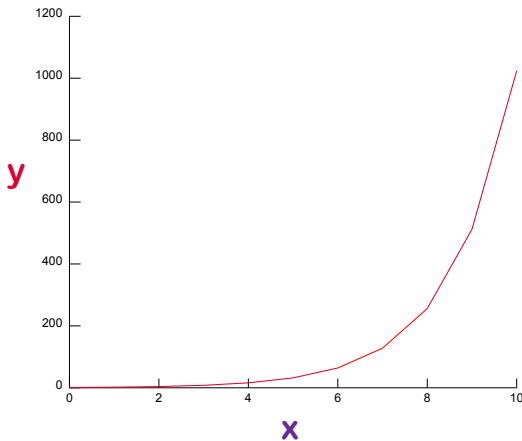
How? Let's look.

First, here's a graph showing **exponential growth**.

We're going to use $y = 2^x$, but any other exponent will work

$$x \quad y = (2^x)$$

0	1
1	2
2	4
3	8
4	16
5	32
6	64
7	128
8	256
9	512
10	1024



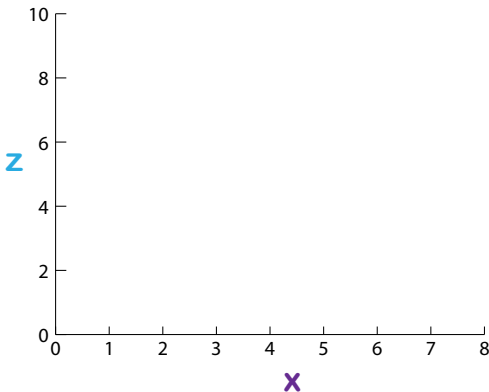
What happens when we take the log of y ?

$$\log y = z$$

$$e^z = y$$

We're going to use $y = 2^x$, but any other exponent will work

x	y
0	1
1	2
2	4
3	8
4	16
5	32
6	64
7	128
8	256
9	512
10	1024

 z 

What happens when we take the log of y ?

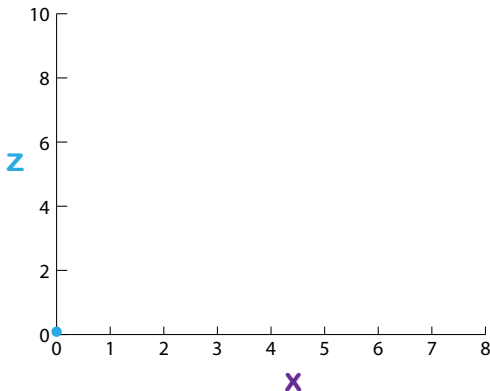
$$\log 1 = 0$$

$$e^0 = 1$$

We're going to use $y = 2^x$, but any other exponent will work

x	y
0	1
1	2
2	4
3	8
4	16
5	32
6	64
7	128
8	256
9	512
10	1024

z
0



What happens when we take the log of y ?

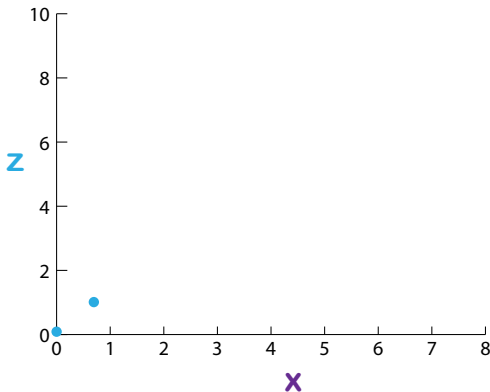
$$\log 2 = .69$$

$$e^{.69} = 2$$

We're going to use $y = 2^x$, but any other exponent will work

X	y
0	1
1	2
2	4
3	8
4	16
5	32
6	64
7	128
8	256
9	512
10	1024

z
0
.69

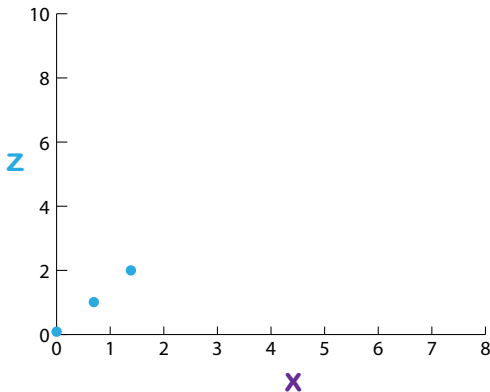


What happens when we take the log of y ?

$$\log 4 = 1.39 \quad e^{1.39} = 4$$

We're going to use $y = 2^x$, but any other exponent will work

X	y	Z
0	1	0
1	2	.69
2	4	1.39
3	8	
4	16	
5	32	
6	64	
7	128	
8	256	
9	512	
10	1024	

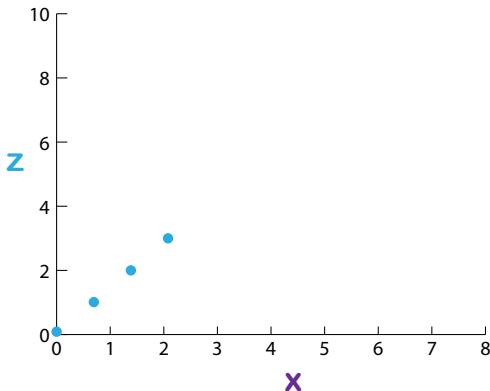


What happens when we take the log of y ?

$$\log 8 = 2.08 \quad e^{2.08} = 8$$

We're going to use $y = 2^x$, but any other exponent will work

X	y	Z
0	1	0
1	2	.69
2	4	1.39
3	8	2.08
4	16	
5	32	
6	64	
7	128	
8	256	
9	512	
10	1024	



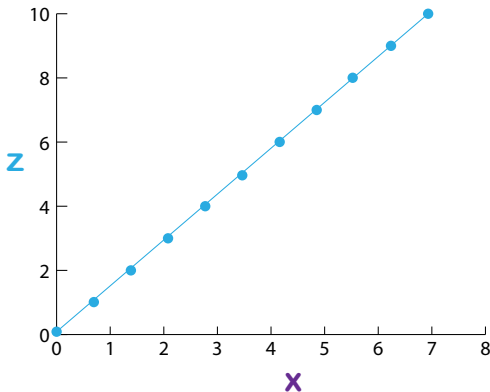
What happens when we take the log of y ?

$$\log y = z$$

$$e^z = y$$

We're going to use $y = 2^x$, but any other exponent will work

x	y	z
0	1	0
1	2	.69
2	4	1.39
3	8	2.08
4	16	2.77
5	32	3.47
6	64	4.16
7	128	4.85
8	256	5.55
9	512	6.24
10	1024	6.93



Interpretation

Interpretation

The log transformation changes the interpretation of β_1 :

Interpretation

The log transformation changes the interpretation of β_1 :

- Regress $\log(Y)$ on $X \rightarrow \beta_1$ approximates **percent increase** in our prediction of Y associated with one unit increase in X .

Interpretation

The log transformation changes the interpretation of β_1 :

- Regress $\log(Y)$ on $X \rightarrow \beta_1$ approximates **percent increase** in our prediction of Y associated with one unit increase in X .
- Regress Y on $\log(X) \rightarrow \beta_1$ approximates increase in Y associated with a **percent increase** in X .

Interpretation

The log transformation changes the interpretation of β_1 :

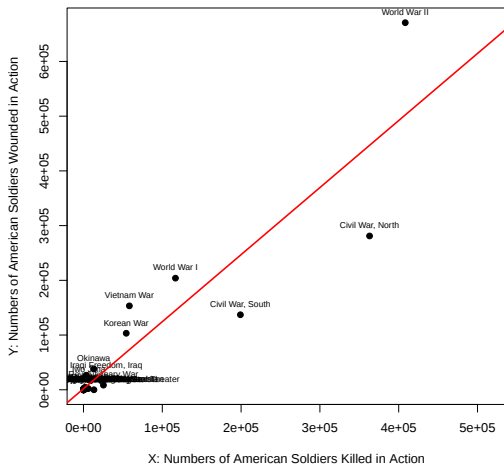
- Regress $\log(Y)$ on $X \rightarrow \beta_1$ approximates **percent increase** in our prediction of Y associated with one unit increase in X .
- Regress Y on $\log(X) \rightarrow \beta_1$ approximates increase in Y associated with a **percent increase** in X .
- Note that these approximations work only for small increments.

Interpretation

The log transformation changes the interpretation of β_1 :

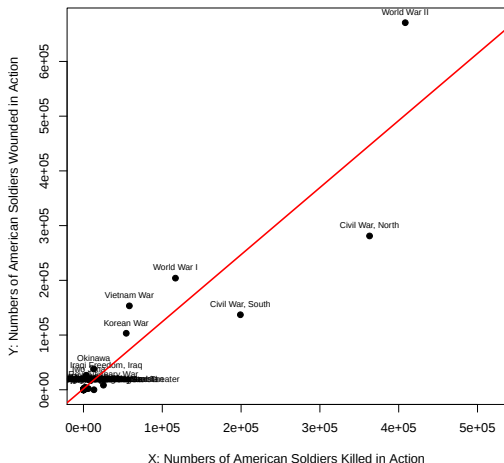
- Regress $\log(Y)$ on $X \rightarrow \beta_1$ approximates **percent increase** in our prediction of Y associated with one unit increase in X .
- Regress Y on $\log(X) \rightarrow \beta_1$ approximates increase in Y associated with a **percent increase** in X .
- Note that these approximations work only for small increments.
- In particular, they do not work when X is a discrete random variable.

Example from the American War Library



$$\hat{\beta}_1 = 1.23 \rightarrow$$

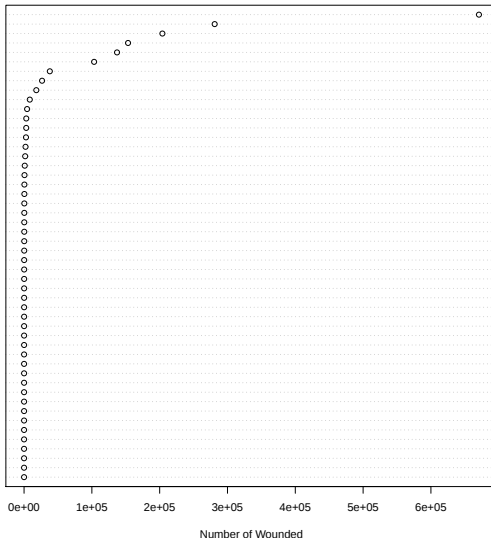
Example from the American War Library



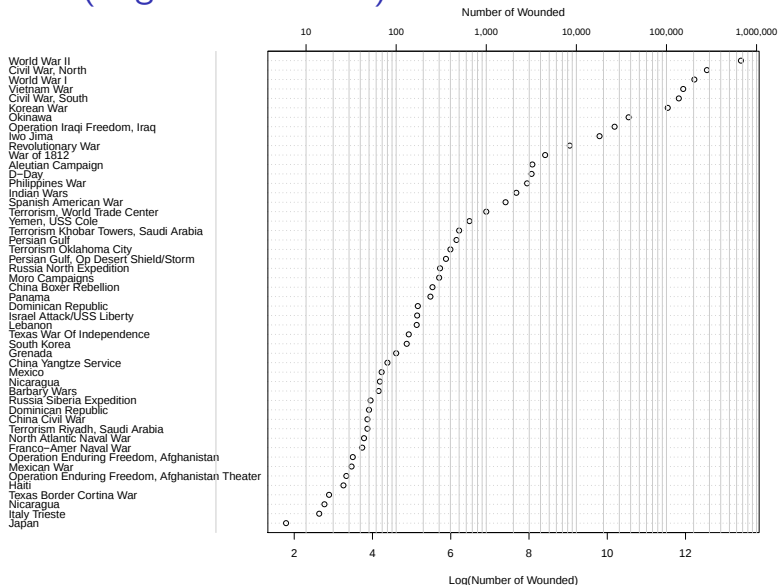
$\hat{\beta}_1 = 1.23 \rightarrow$ One additional soldier killed predicts 1.23 additional soldiers wounded

Wounded (Scale in Levels)

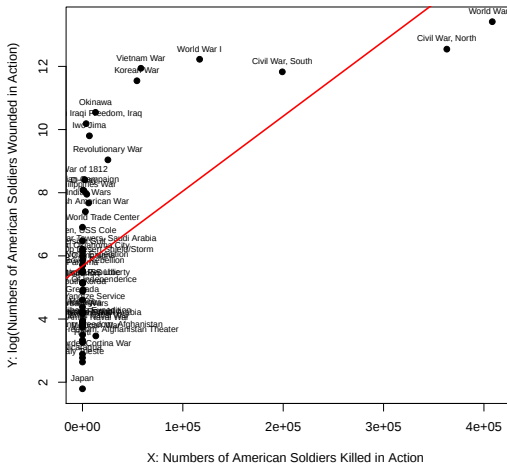
World War II
Civil War, North
World War I
Vietnam War
Civil War, South
Korean War
Okinawa
Operation Iraqi Freedom, Iraq
Iwo Jima
Revolutionary War
War of 1812
Aleutian Campaign
D-Day
Philippines War
Indian Wars
Spanish American War
Terrorism, World Trade Center
Yemen, USS Cole
Terrorism Khobar Towers, Saudi Arabia
Persian Gulf
Terrorism Oklahoma City
Persian Gulf, Op Desert Shield/Storm
Russia North Expedition
Moro Campaigns
China Boxer Rebellion
Panama
Dominican Republic
Israel Attack/USS Liberty
Lebanon
Texas War Of Independence
South Korea
Grenada
China Yangtze Service
Mexico
Nicaragua
Barbary Wars
Russia Siberia Expedition
Dominican Republic
China Civil War
Terrorism Riyadh, Saudi Arabia
North Atlantic Naval War
Franco-Amer Naval War
Operation Enduring Freedom, Afghanistan
Mexican War
Operation Enduring Freedom, Afghanistan Theater
Haiti
Texas Border Cortina War
Nicaragua
Italy Trieste
Japan



Wounded (Logarithmic Scale)

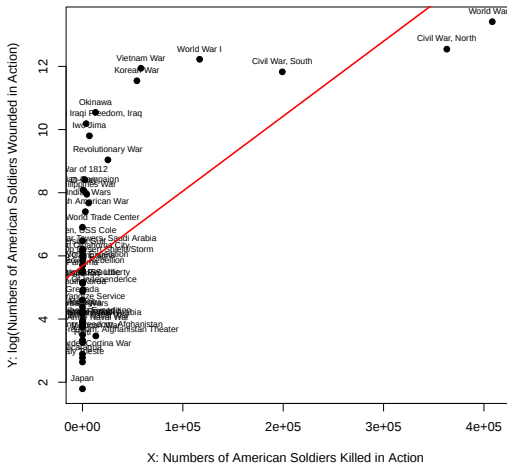


Regression: Log-Level



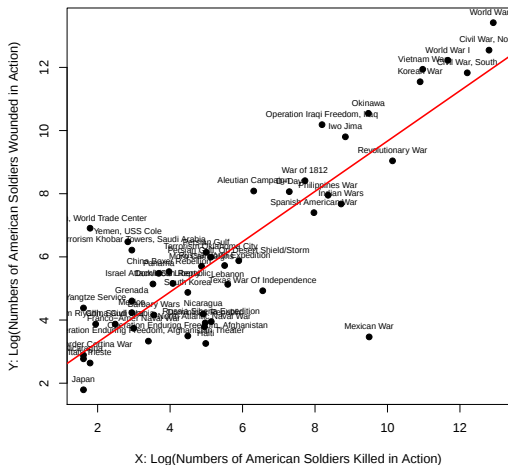
$$\hat{\beta}_1 = 0.0000237 \longrightarrow$$

Regression: Log-Level



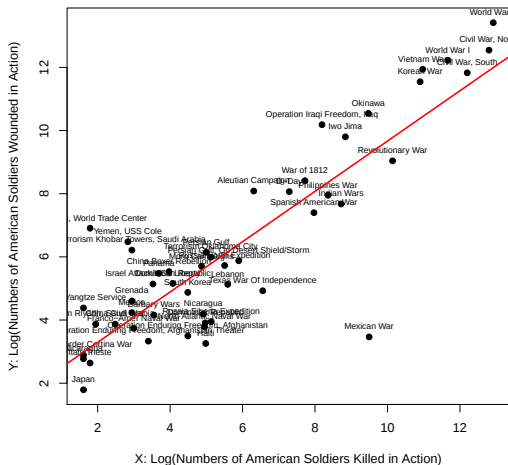
$\hat{\beta}_1 = 0.0000237 \rightarrow$ One additional soldier killed predicts 0.0023 percent increase in the number of soldiers wounded

Regression: Log-Log



$$\hat{\beta}_1 = 0.797 \rightarrow$$

Regression: Log-Log



$\hat{\beta}_1 = 0.797 \rightarrow$ A percent increase in deaths predicts 0.797 percent increase in the wounded

Four Most Commonly Used Models

Model	Equation	β_1 Interpretation
Level-Level	$Y = \beta_0 + \beta_1 X$	$\Delta Y = \beta_1 \Delta X$
Log-Level	$\log(Y) = \beta_0 + \beta_1 X$	$\% \Delta Y = 100 \beta_1 \Delta X$
Level-Log	$Y = \beta_0 + \beta_1 \log(X)$	$\Delta Y = (\beta_1 / 100) \% \Delta X$
Log-Log	$\log(Y) = \beta_0 + \beta_1 \log(X)$	$\% \Delta Y = \beta_1 \% \Delta X$

Why Does This Approximation Work?

Why Does This Approximation Work?

A useful thing to know is that for small x ,

$$\log(1 + x) \approx x$$

$$\exp(x) \approx 1 + x$$

Why Does This Approximation Work?

A useful thing to know is that for small x ,

$$\log(1 + x) \approx x$$

$$\exp(x) \approx 1 + x$$

This can be derived from a series expansion of the log function.
Numerically, when $|x| \leq .1$, the approximation is within 0.001.

Why Does This Approximation Work?

Take two numbers $a > b > 0$. The percentage difference between a and b is

$$p = 100 \left(\frac{a - b}{b} \right)$$

Why Does This Approximation Work?

Take two numbers $a > b > 0$. The percentage difference between a and b is

$$p = 100 \left(\frac{a - b}{b} \right)$$

We can rewrite this as

$$\frac{a}{b} = 1 + \frac{p}{100}$$

Why Does This Approximation Work?

Take two numbers $a > b > 0$. The percentage difference between a and b is

$$p = 100 \left(\frac{a - b}{b} \right)$$

We can rewrite this as

$$\frac{a}{b} = 1 + \frac{p}{100}$$

Taking natural logs

$$\log(a) - \log(b) = \log \left(1 + \frac{p}{100} \right)$$

Why Does This Approximation Work?

Take two numbers $a > b > 0$. The percentage difference between a and b is

$$p = 100 \left(\frac{a - b}{b} \right)$$

We can rewrite this as

$$\frac{a}{b} = 1 + \frac{p}{100}$$

Taking natural logs

$$\log(a) - \log(b) = \log \left(1 + \frac{p}{100} \right)$$

Applying our approximation and multiplying by 100 we find,

$$p \approx 100 (\log(a) - \log(b))$$

Be Careful: Log-Level with binary X

Assume we have: $\log(Y) = \beta_0 + \beta_1 X$ where X is binary with values 1 or 0.

Assume $\beta_1 > .2$. What is the problem with saying that a one unit increase in X is associated with a $\beta_1 \cdot 100$ percent change in Y ?

Be Careful: Log-Level with binary X

Assume we have: $\log(Y) = \beta_0 + \beta_1 X$ where X is binary with values 1 or 0.
Assume $\beta_1 > .2$. What is the problem with saying that a one unit increase in X is associated with a $\beta_1 \cdot 100$ percent change in Y ?

Log approximation is inaccurate for large changes like going from $X = 0$ to $X = 1$.

Be Careful: Log-Level with binary X

Assume we have: $\log(Y) = \beta_0 + \beta_1 X$ where X is binary with values 1 or 0.
Assume $\beta_1 > .2$. What is the problem with saying that a one unit increase in X is associated with a $\beta_1 \cdot 100$ percent change in Y ?

Log approximation is inaccurate for large changes like going from $X = 0$ to $X = 1$. Instead the percent change in Y when X goes from 0 to 1 needs to be computed using:

$$100(\hat{Y}_{X=1} - \hat{Y}_{X=0})/\hat{Y}_{X=0} = 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1)$$

Be Careful: Log-Level with binary X

Assume we have: $\log(Y) = \beta_0 + \beta_1 X$ where X is binary with values 1 or 0.
Assume $\beta_1 > .2$. What is the problem with saying that a one unit increase in X is associated with a $\beta_1 \cdot 100$ percent change in Y ?

Log approximation is inaccurate for large changes like going from $X = 0$ to $X = 1$. Instead the percent change in Y when X goes from 0 to 1 needs to be computed using:

$$\begin{aligned} 100(\hat{Y}_{X=1} - \hat{Y}_{X=0})/Y_{X=0} &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \\ &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \end{aligned}$$

Be Careful: Log-Level with binary X

Assume we have: $\log(Y) = \beta_0 + \beta_1 X$ where X is binary with values 1 or 0.
Assume $\beta_1 > .2$. What is the problem with saying that a one unit increase in X is associated with a $\beta_1 \cdot 100$ percent change in Y ?

Log approximation is inaccurate for large changes like going from $X = 0$ to $X = 1$. Instead the percent change in Y when X goes from 0 to 1 needs to be computed using:

$$\begin{aligned} 100(\hat{Y}_{X=1} - \hat{Y}_{X=0})/Y_{X=0} &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \\ &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \\ &= 100(\exp(\hat{\beta}_1) - 1) \end{aligned}$$

Be Careful: Log-Level with binary X

Assume we have: $\log(Y) = \beta_0 + \beta_1 X$ where X is binary with values 1 or 0.
Assume $\beta_1 > .2$. What is the problem with saying that a one unit increase in X is associated with a $\beta_1 \cdot 100$ percent change in Y ?

Log approximation is inaccurate for large changes like going from $X = 0$ to $X = 1$. Instead the percent change in Y when X goes from 0 to 1 needs to be computed using:

$$\begin{aligned} 100(\hat{Y}_{X=1} - \hat{Y}_{X=0})/\hat{Y}_{X=0} &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \\ &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \\ &= 100(\exp(\hat{\beta}_1) - 1) \end{aligned}$$

Recall: $\log(\hat{Y}_{X=1}) - \log(\hat{Y}_{X=0}) = \log(\hat{Y}_{X=1}/\hat{Y}_{X=0}) = \hat{\beta}_1$.

Be Careful: Log-Level with binary X

Assume we have: $\log(Y) = \beta_0 + \beta_1 X$ where X is binary with values 1 or 0.
Assume $\beta_1 > .2$. What is the problem with saying that a one unit increase in X is associated with a $\beta_1 \cdot 100$ percent change in Y ?

Log approximation is inaccurate for large changes like going from $X = 0$ to $X = 1$. Instead the percent change in Y when X goes from 0 to 1 needs to be computed using:

$$\begin{aligned} 100(\hat{Y}_{X=1} - \hat{Y}_{X=0})/\hat{Y}_{X=0} &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \\ &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \\ &= 100(\exp(\hat{\beta}_1) - 1) \end{aligned}$$

Recall: $\log(\hat{Y}_{X=1}) - \log(\hat{Y}_{X=0}) = \log(\hat{Y}_{X=1}/\hat{Y}_{X=0}) = \hat{\beta}_1$.

Be Careful: Log-Level with binary X

Assume we have: $\log(Y) = \beta_0 + \beta_1 X$ where X is binary with values 1 or 0.
Assume $\beta_1 > .2$. What is the problem with saying that a one unit increase in X is associated with a $\beta_1 \cdot 100$ percent change in Y ?

Log approximation is inaccurate for large changes like going from $X = 0$ to $X = 1$. Instead the percent change in Y when X goes from 0 to 1 needs to be computed using:

$$\begin{aligned} 100(\hat{Y}_{X=1} - \hat{Y}_{X=0})/\hat{Y}_{X=0} &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \\ &= 100((\hat{Y}_{X=1}/\hat{Y}_{X=0}) - 1) \\ &= 100(\exp(\hat{\beta}_1) - 1) \end{aligned}$$

Recall: $\log(\hat{Y}_{X=1}) - \log(\hat{Y}_{X=0}) = \log(\hat{Y}_{X=1}/\hat{Y}_{X=0}) = \hat{\beta}_1$.

A one unit change in X (ie. going from 0 to 1) creates $100(\exp(\hat{\beta}_1) - 1)$ percent increase in our prediction $\hat{Y}$.

Interpreting a Logged Outcome

Interpreting a Logged Outcome

- On the last few slides, there was a bit that was a little dodgy.

Interpreting a Logged Outcome

- On the last few slides, there was a bit that was a little dodgy.
- When we log the **outcome**, we are no longer approximating $E[Y|X]$ we are approximating $E[\log(Y)|X]$.

Interpreting a Logged Outcome

- On the last few slides, there was a bit that was a little dodgy.
- When we log the **outcome**, we are no longer approximating $E[Y|X]$ we are approximating $E[\log(Y)|X]$.
- Jensen's inequality gives us information on this relation:
 $f(E[X]) \leq E[f(X)]$ for any convex function $f()$.

Interpreting a Logged Outcome

- On the last few slides, there was a bit that was a little dodgy.
- When we log the **outcome**, we are no longer approximating $E[Y|X]$ we are approximating $E[\log(Y)|X]$.
- Jensen's inequality gives us information on this relation:
 $f(E[X]) \leq E[f(X)]$ for any convex function $f()$.
- In practice, this means we are no longer characterizing the expectation of Y and it is technically inaccurate to talk about Y 'on average' changing in a certain way.

Interpreting a Logged Outcome

- On the last few slides, there was a bit that was a little dodgy.
- When we log the **outcome**, we are no longer approximating $E[Y|X]$ we are approximating $E[\log(Y)|X]$.
- Jensen's inequality gives us information on this relation:
 $f(E[X]) \leq E[f(X)]$ for any convex function $f()$.
- In practice, this means we are no longer characterizing the expectation of Y and it is technically inaccurate to talk about Y 'on average' changing in a certain way.
- What are we characterizing? The **geometric mean**.

Geometric Mean

$$\exp(E(\log(Y))) = \exp\left(\frac{1}{N} \sum_{i=1}^N \log(Y_i)\right)$$

Geometric Mean

$$\begin{aligned}\exp(E(\log(Y))) &= \exp\left(\frac{1}{N} \sum_{i=1}^N \log(Y_i)\right) \\ &= \exp\left(\frac{1}{N} \log\left(\prod_{i=1}^N Y_i\right)\right)\end{aligned}$$

Geometric Mean

$$\begin{aligned}\exp(E(\log(Y))) &= \exp\left(\frac{1}{N} \sum_{i=1}^N \log(Y_i)\right) \\ &= \exp\left(\frac{1}{N} \log\left(\prod_{i=1}^N Y_i\right)\right) \\ &= \exp\left(\log\left(\left(\prod_{i=1}^N Y_i\right)^{\frac{1}{N}}\right)\right)\end{aligned}$$

Geometric Mean

$$\begin{aligned}\exp(E(\log(Y))) &= \exp\left(\frac{1}{N} \sum_{i=1}^N \log(Y_i)\right) \\ &= \exp\left(\frac{1}{N} \log\left(\prod_{i=1}^N Y_i\right)\right) \\ &= \exp\left(\log\left(\left(\prod_{i=1}^N Y_i\right)^{\frac{1}{N}}\right)\right) \\ &= \left(\prod_{i=1}^N Y_i\right)^{\frac{1}{N}}\end{aligned}$$

Geometric Mean

$$\begin{aligned}\exp(E(\log(Y))) &= \exp\left(\frac{1}{N} \sum_{i=1}^N \log(Y_i)\right) \\&= \exp\left(\frac{1}{N} \log\left(\prod_{i=1}^N Y_i\right)\right) \\&= \exp\left(\log\left(\left(\prod_{i=1}^N Y_i\right)^{\frac{1}{N}}\right)\right) \\&= \left(\prod_{i=1}^N Y_i\right)^{\frac{1}{N}} \\&= \text{Geometric Mean}(Y)\end{aligned}$$

Geometric Mean

$$\begin{aligned}\exp(E(\log(Y))) &= \exp\left(\frac{1}{N} \sum_{i=1}^N \log(Y_i)\right) \\ &= \exp\left(\frac{1}{N} \log\left(\prod_{i=1}^N Y_i\right)\right) \\ &= \exp\left(\log\left(\left(\prod_{i=1}^N Y_i\right)^{\frac{1}{N}}\right)\right) \\ &= \left(\prod_{i=1}^N Y_i\right)^{\frac{1}{N}} \\ &= \text{Geometric Mean}(Y)\end{aligned}$$

The **geometric mean** is a robust measure of central tendency.

THE INTERGENERATIONAL ELASTICITY OF WHAT? THE CASE FOR REDEFINING THE WORKHORSE MEASURE OF ECONOMIC MOBILITY

Pablo A. Mitnik* 

David B. Grusky*

Abstract

The intergenerational elasticity (IGE) has been assumed to refer to the expectation of children's income when in fact it pertains to the geometric mean of children's income. We show that mobility analyses based on the conventional IGE have been widely misinterpreted, are subject to selection bias, and cannot disentangle the different channels for transmitting economic status across generations. The solution to these problems—estimating the IGE of expected income or earnings—returns the field to what it has long meant to estimate. Under this approach, intergenerational persistence is found to be substantially higher, thus raising the possibility that the field's stock results are misleading.

Keywords

intergenerational economic mobility, elasticity of expected income, selection bias, gender, marriage and economic mobility

Core Idea

Classic approach :

$$\underbrace{E(\log(Y) | X)}_{\substack{\text{Mean of log} \\ \text{offspring income } Y \\ \text{given parent income } X}} = \beta_0 + \underbrace{\beta_1}_{\substack{\text{Intergenerational} \\ \text{elasticity} \\ \text{(IGE)}}} \underbrace{\log(X)}_{\substack{\text{Log parent} \\ \text{income}}}$$

Core Idea

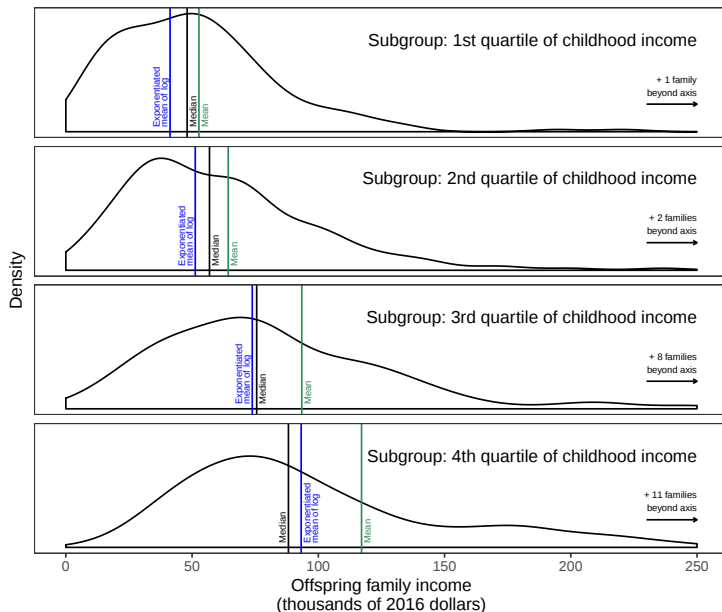
Classic approach :

$$\underbrace{E(\log(Y) | X)}_{\substack{\text{Mean of log} \\ \text{offspring income } Y \\ \text{given parent income } X}} = \beta_0 + \underbrace{\beta_1}_{\substack{\text{Intergenerational} \\ \text{elasticity} \\ \text{(IGE)}}} \underbrace{\log(X)}_{\substack{\text{Log parent} \\ \text{income}}}$$

MG proposal :

$$\underbrace{\log(E(Y | X))}_{\substack{\text{Log of mean} \\ \text{offspring income } Y \\ \text{given parent income } X}} = \alpha_0 + \underbrace{\alpha_1}_{\substack{\text{Intergenerational} \\ \text{elasticity of the} \\ \text{expectation (IGEE)}}} \underbrace{\log(X)}_{\substack{\text{Log parent} \\ \text{income}}}$$

Geometric Mean is Closer to the Median Than the Mean



COMMENT: SUMMARIZING INCOME MOBILITY WITH MULTIPLE SMOOTH QUANTILES INSTEAD OF PARAMETERIZED MEANS

*Ian Lundberg**

*Brandon M. Stewart**

*Department of Sociology and Office of Population Research, Princeton University,
Princeton, NJ, USA

Corresponding Author: Ian Lundberg, ilundberg@princeton.edu

DOI: 10.1177/0081175020931126

Single-number summaries that capture the relationship of socioeconomic outcomes across generations are a cornerstone of economic mobility research. Studies often focus on the intergenerational elasticity (IGE) of income: the coefficient β_1 on parent log income in a model predicting offspring log income (e.g., Aaronson and Mazumder 2008; Björklund and Jäntti 1997; Solon 2004). A large β_1 is often interpreted as evidence that incomes persist to a substantial degree across generations.

Images from this section are from this paper or earlier drafts of it.

Two Implicit Choices

Two Implicit Choices

(1) Summary Statistics for the Conditional Distribution

Two Implicit Choices

(1) Summary Statistics for the Conditional Distribution

and

(2) Assume or Learn a Functional Form

Two Implicit Choices

(1) Summary Statistics for the Conditional Distribution

(gets you down to one number per value of x)

and

(2) Assume or Learn a Functional Form

Two Implicit Choices

(1) Summary Statistics for the Conditional Distribution

(gets you down to one number per value of x)

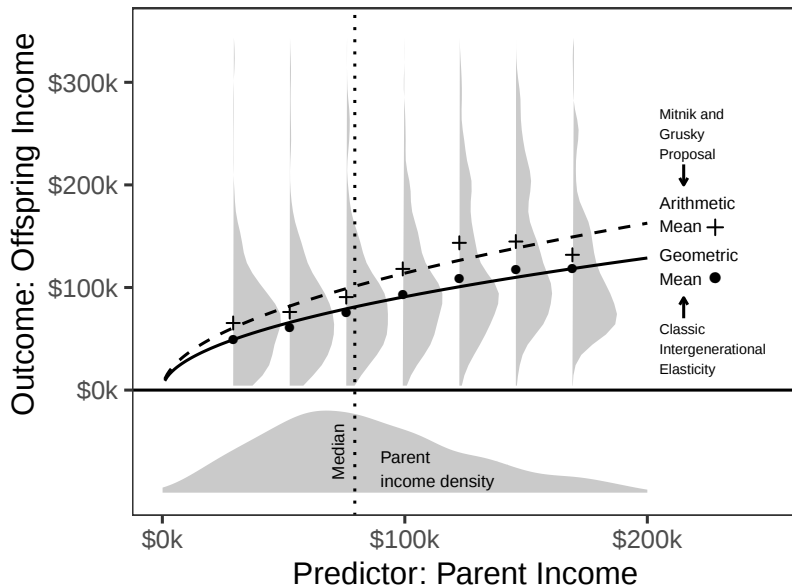
and

(2) Assume or Learn a Functional Form

(potentially simplifies the set of summary statistics to a single number)

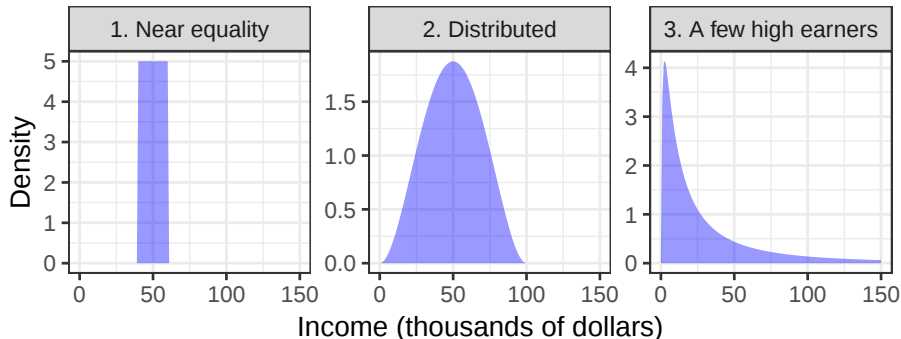
Visualizing the MG Proposal

Visualizing the MG Proposal



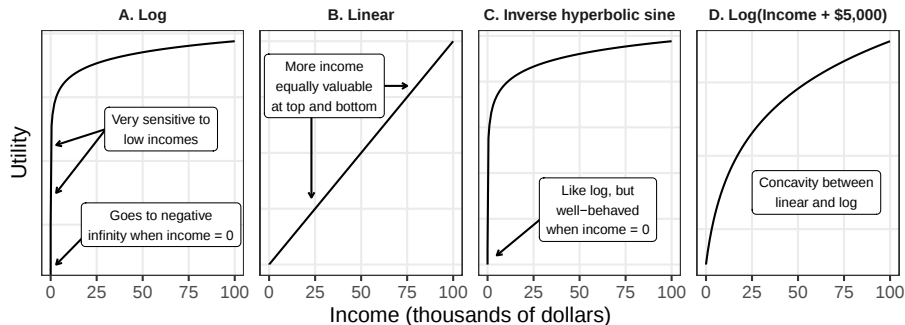
Single Summary Statistics Necessarily Mask Information

Single Summary Statistics Necessarily Mask Information



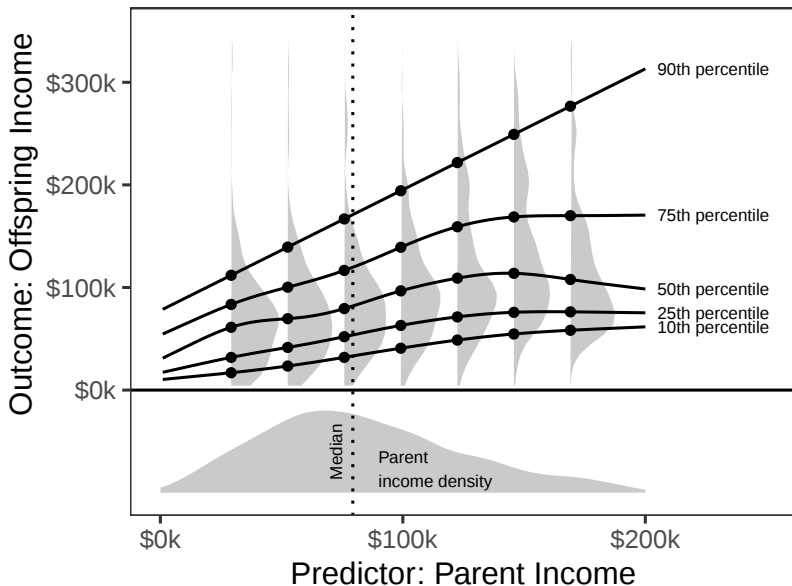
The Mean is a Normative Choice

The Mean is a Normative Choice



A New Proposal

A New Proposal



Single Number Summaries

Single Number Summaries

- A key selling point of the conventional IGE, the MG proposal and regression more broadly is the **single-number summary**.

Single Number Summaries

- A key selling point of the conventional IGE, the MG proposal and regression more broadly is the **single-number summary**.
- Any such summary necessitates a **loss of information**.

Single Number Summaries

- A key selling point of the conventional IGE, the MG proposal and regression more broadly is the **single-number summary**.
- Any such summary necessitates a **loss of information**.
- Even with more complex functional forms, we can always calculate such a summary. For instance here, median is (on average) \$4k higher when parent income is \$10k higher.

Single Number Summaries

- A key selling point of the conventional IGE, the MG proposal and regression more broadly is the **single-number summary**.
- Any such summary necessitates a **loss of information**.
- Even with more complex functional forms, we can always calculate such a summary. For instance here, median is (on average) \$4k higher when parent income is \$10k higher.
- We obtain this by simply plugging in the 50th percentile at each offspring income, adding \$10k to each parent income and taking the average.

Single Number Summaries

- A key selling point of the conventional IGE, the MG proposal and regression more broadly is the **single-number summary**.
- Any such summary necessitates a **loss of information**.
- Even with more complex functional forms, we can always calculate such a summary. For instance here, median is (on average) \$4k higher when parent income is \$10k higher.
- We obtain this by simply plugging in the 50th percentile at each offspring income, adding \$10k to each parent income and taking the average.
- If you are willing to commit to a **quantity of interest**, you can usually estimate it directly.

Single Number Summaries

- A key selling point of the conventional IGE, the MG proposal and regression more broadly is the **single-number summary**.
- Any such summary necessitates a **loss of information**.
- Even with more complex functional forms, we can always calculate such a summary. For instance here, median is (on average) \$4k higher when parent income is \$10k higher.
- We obtain this by simply plugging in the 50th percentile at each offspring income, adding \$10k to each parent income and taking the average.
- If you are willing to commit to a **quantity of interest**, you can usually estimate it directly.
- At their best, single-number summaries are a way that the reader can calculate any approximation to a variety of quantities they are interested in. At their worst, they are a way for authors to abdicate responsibility for choosing a clear quantity of interest.

This Week in Review

In this brave new world with 2 independent variables:

This Week in Review

In this brave new world with 2 independent variables:

- 1 β 's have slightly different interpretations

This Week in Review

In this brave new world with 2 independent variables:

- ① β 's have slightly different interpretations
- ② OLS still minimizing the sum of the squared residuals

This Week in Review

In this brave new world with 2 independent variables:

- ① β 's have slightly different interpretations
- ② OLS still minimizing the sum of the squared residuals
- ③ Small adjustments to OLS assumptions and inference

This Week in Review

In this brave new world with 2 independent variables:

- ① β 's have slightly different interpretations
- ② OLS still minimizing the sum of the squared residuals
- ③ Small adjustments to OLS assumptions and inference
- ④ We can model some non-linearities using polynomials, logs, or other transformations

This Week in Review

In this brave new world with 2 independent variables:

- ① β 's have slightly different interpretations
- ② OLS still minimizing the sum of the squared residuals
- ③ Small adjustments to OLS assumptions and inference
- ④ We can model some non-linearities using polynomials, logs, or other transformations
- ⑤ We can optionally consider interactions, but must take care to interpret them correctly

This Week in Review

In this brave new world with 2 independent variables:

- 1 β 's have slightly different interpretations
- 2 OLS still minimizing the sum of the squared residuals
- 3 Small adjustments to OLS assumptions and inference
- 4 We can model some non-linearities using polynomials, logs, or other transformations
- 5 We can optionally consider interactions, but must take care to interpret them correctly

Next week: Linear Regression in its Full Glory!