

ASTA

The ASTA team

Contents

1	The regression problem	2
1.1	We want to predict	2
1.2	Initial graphics	2
1.3	Simple linear regression	3
1.4	Model for linear regression	4
1.5	Least squares	5
1.6	The prediction equation and residuals	5
1.7	Estimation of conditional standard deviation	5
1.8	Example in R	6
1.9	Test for independence	6
1.10	Example	7
1.11	Confidence interval for slope	7
1.12	Correlation	8
2	R-squared: Reduction in prediction error	9
2.1	R-squared: Reduction in prediction error	9
2.2	Graphical illustration of sums of squares	9
2.3	r^2 : Reduction in prediction error	10
3	Multiple regression model	10
3.1	Multiple regression model	10
3.2	Example	10
3.3	Correlations	11
3.4	Several predictors	12
3.5	Example	12
3.6	Simpsons paradox	13
4	The general model	13
4.1	Regression model	13
4.2	Interpretation of parameters	13
5	Estimation	14
5.1	Estimation of model	14
6	Multiple R-squared	14
6.1	Multiple R^2	14
6.2	Example	15
6.3	Example	16
7	F-test for effect of predictors	16
7.1	F-test	16
7.2	Example	17

8	Test for interaction	18
8.1	Interaction between effects of predictors	18

1 The regression problem

1.1 We want to predict

- We will study the dataset `trees`, which is on the course website:

```
import pandas as pd

trees = pd.read_csv("https://asta.math.aau.dk/datasets?file=trees.txt", sep='\t')
trees.head()
```

```
##      Girth  Height  Volume
## 0      8.3      70    10.3
## 1      8.6      65    10.3
## 2      8.8      63    10.2
## 3     10.5      72    16.4
## 4     10.7      81    18.8
```

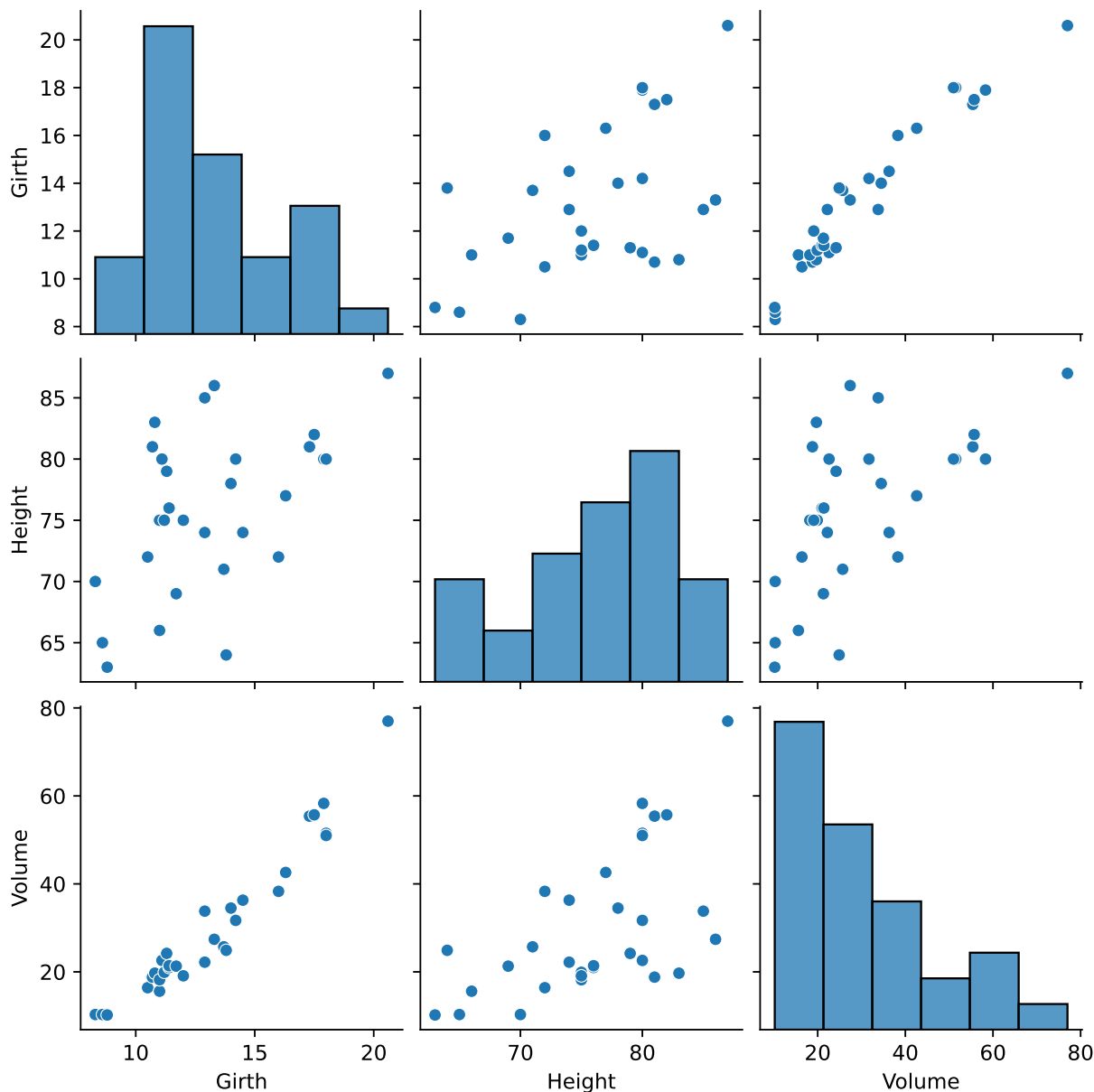
- In this experiment we have measurements of 3 variables for 31 randomly chosen trees:
- **Girth** numeric. Tree diameter in inches.
- **Height** numeric. Height in ft.
- **Volume** numeric. Volume of timber in cubic ft.
- We want to predict the tree volume, if we measure the tree height and/or the tree girth (diameter).
- This type of problem is called **regression**.
- Relevant terminology:
 - We measure a quantitative **response** y , e.g. **Volume**.
 - In connection with the response value y we also measure one (later we will consider several) potential **explanatory** variable x . Another name for the explanatory variable is **predictor**.

1.2 Initial graphics

- Any analysis starts with relevant graphics.

```
import seaborn as sns
import matplotlib.pyplot as plt

p = sns.pairplot(trees)
```



- For each combination of the variables we plot the (x, y) values.
- It looks like **Girth** is a good predictor for **Volume**.
- If we only are interested in the association between two (and not three or more) variables we use the usual `lmpplot` function.

1.3 Simple linear regression

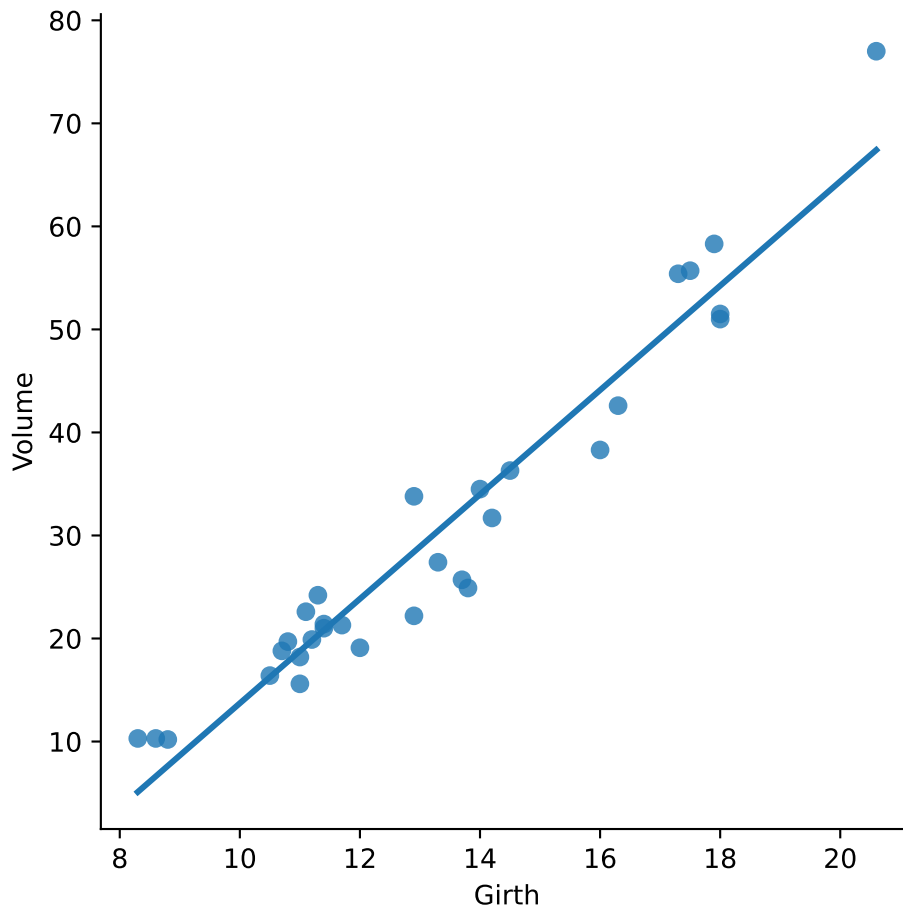
- We choose to use $x=\text{Girth}$ as predictor for $y=\text{Volume}$. When we only use one predictor we are doing **simple regression**.
- The simplest **model** to describe an association between **response** y and a **predictor** x is **simple linear regression**.
- I.e. ideally we see the picture

$$y(x) = \alpha + \beta x$$

where

- α is called the **Intercept** - the line's intercept with the y -axis, corresponding to the response for $x = 0$.
- β is called **Slope** - the line's slope, corresponding to the change in response, when we increase the predictor by one unit.

```
p = sns.lmplot(x='Girth', y='Volume', data=trees, ci=None)
```



1.4 Model for linear regression

- Assume we have a sample with joint measurements (x, y) of predictor and response.
- Ideally the model states that

$$y(x) = \alpha + \beta x,$$

but due to random variation there are deviations from the line.

- What we observe can then be described by

$$y = \alpha + \beta x + \varepsilon,$$

where ε is a **random error**, which causes deviations from the line.

- We will continue under the following **fundamental assumption**:
 - The errors ε are normally distributed with mean zero and standard deviation $\sigma_{y|x}$.
- We call $\sigma_{y|x}$ the **conditional standard deviation** given x , since it describes the variation in y around the regression line, when we know x .

1.5 Least squares

- In summary, we have a model with 3 parameters:
 - (α, β) which determine the line
 - $\sigma_{y|x}$ which is the standard deviation of the deviations from the line.
- How are these estimated, when we have a sample $(x_1, y_1) \dots (x_n, y_n)$ of (x, y) values??
- To do this we focus on the errors

$$\varepsilon_i = y_i - \alpha - \beta x_i$$

which should be as close to 0 as possible in order to fit the data best possible.

- We will choose the line, which minimizes the sum of squares of the errors:

$$\sum_{i=1}^n \varepsilon_i^2 = \sum_{i=1}^n (y_i - \alpha - \beta x_i)^2.$$

- If we set the partial derivatives to zero we obtain two linear equations for the unknowns (α, β) , where the solution (a, b) is given by:

$$b = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2} \quad \text{and} \quad a = \bar{y} - b\bar{x}$$

1.6 The prediction equation and residuals

- The equation for the estimates $(\hat{\alpha}, \hat{\beta}) = (a, b)$,

$$\hat{y} = a + bx$$

is called **the prediction equation**, since it can be used to predict y for any value of x .

- Note: The prediction equation is determined by the current sample. I.e. there is an uncertainty attached to it. A new sample would without any doubt give a different prediction equation.
- Our best estimate of the errors is

$$e_i = y_i - \hat{y} = y_i - a - bx_i,$$

i.e. the vertical deviations from the prediction line.

- These quantities are called **residuals**.
- We have that
 - The prediction line passes through the point $(\bar{x}, \bar{y})$.
 - The sum of the residuals is zero.

1.7 Estimation of conditional standard deviation

- To estimate $\sigma_{y|x}$ we need **Sum of Squared Errors**

$$SSE = \sum_{i=1}^n e_i^2 = \sum_{i=1}^n (y_i - \hat{y}_i)^2,$$

which is the squared distance between the model and data.

- We then estimate $\sigma_{y|x}$ by the quantity

$$s_{y|x} = \sqrt{\frac{SSE}{n-2}}$$

- Instead of n we divide SSE with **the degrees of freedom** $df = n - 2$. Theory shows, that this is reasonable.
- The degrees of freedom df are determined as the sample size minus the number of parameters in the regression equation.
- In the current setup we have 2 parameters: (α, β) .

1.8 Example in R

```
import numpy as np
import statsmodels.formula.api as smf

model = smf.ols('Volume ~ Girth', data=trees).fit()
model.summary(slim = True) # text output

## <class 'statsmodels.iolib.summary.Summary'>
## """
##                               OLS Regression Results
## =====
## Dep. Variable:                Volume    R-squared:                0.935
## Model:                      OLS      Adj. R-squared:          0.933
## No. Observations:            31        F-statistic:             419.4
## Covariance Type:            nonrobust   Prob (F-statistic):       8.64e-19
## =====
##               coef      std err          t      P>|t|      [0.025      0.975]
## -----
## Intercept      -36.9435      3.365     -10.978      0.000     -43.826     -30.061
## Girth           5.0659      0.247      20.478      0.000       4.560       5.572
## =====
##
## Notes:
## [1] Standard Errors assume that the covariance matrix of the errors is correctly specified.
## """

[model.resid.min(), model.resid.max(), model.resid.median()]

## [np.float64(-8.065359510665438), np.float64(9.586816813996933), np.float64(0.15196668471900665)]
np.sqrt(model.mse_resid)

## np.float64(4.2519875217291965)
```

- The estimated residuals vary from -8.065 to 9.578 with median 0.152.
- The estimate of **Intercept** is $a = -36.9435$
- The estimate of slope of **Girth** is $b = 5.0659$
- The estimate of the conditional standard deviation (also called residual standard error) is $s_{y|x} = 4.252$ with $31 - 2 = 29$ degrees of freedom.

1.9 Test for independence

- We consider the regression model

$$y = \alpha + \beta x + \varepsilon$$

where we use a sample to obtain estimates (a, b) of (α, β) , an estimate $s_{y|x}$ of $\sigma_{y|x}$ and the degrees of freedom $df = n - 2$.

- We are going to test

$$H_0 : \beta = 0 \quad \text{against} \quad H_a : \beta \neq 0$$

- The null hypothesis specifies, that y **doesn't** depend linearly on x .
- In other words the question is: Is the value of b far away from zero?
- It can be shown that b has standard error

$$se_b = \frac{s_{y|x}}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2}}$$

with df degrees of freedom.

- So, we want to use the test statistic

$$t_{\text{obs}} = \frac{b}{se_b}$$

which has to be evaluated in a t-distribution with df degrees of freedom.

1.10 Example

- Recall the summary of our example:

```
coef_table = model.summary2().tables[1] # data output
coef_table

##              Coef.  Std.Err.      t      P>|t|      [0.025      0.975]
## Intercept -36.943459  3.365145 -10.978267  7.621449e-12 -43.825953 -30.060965
## Girth      5.065856  0.247377  20.478288  8.644334e-19  4.559914  5.571799
np.sqrt(model.mse_resid)

## np.float64(4.2519875217291965)
model.df_resid

## np.float64(29.0)
```

- As we noted previously $b = 5.0659$ and $s_{y|x} = 4.252$ with $df = 29$ degrees of freedom.
- In the second column(Std. Error) of the **Coefficients** table we find $se_b = 0.2474$.
- The observed t-score (test statistic) is then

$$t_{\text{obs}} = \frac{b}{se_b} = \frac{5.0659}{0.2474} = 20.48$$

which also can be found in the third column (**t value**).

- The corresponding p-value is found in the usual way by using the t-distribution with 29 degrees of freedom.
- In the fourth column($\text{Pr}(>|t|)$) we see that the p-value is less than 2×10^{-16} . This is no surprise since the t-score was way above 3.

1.11 Confidence interval for slope

- When we have both the standard error and the reference distribution, we can construct a confidence interval in the usual way:

$$b \pm t_{\text{crit}} se_b,$$

where the t-score is determined by the confidence level.

- The t-score can be found using **t.ppf** in **Python**: In our example we have 29 degrees of freedom and with a confidence level of 95% we get $t_{\text{crit}} = \text{from scipy import stats; stats.t.ppf}(0.975, df=29) = 2.045$.
- If you are lazy (like most statisticians are):

```
model.conf_int(alpha = 0.05)

##              0              1
## Intercept -43.825953 -30.060965
## Girth      4.559914  5.571799
```

- i.e. (4.56, 5.57) is a 95% confidence interval for the slope of **Girth**.

1.12 Correlation

- The estimated slope b in a linear regression doesn't say anything about the strength of association between y and x .
- **Girth** was measured in inches, but if we rather measured it in kilometers the slope is much larger: An increase of 1km in **Girth** yield an enormous increase in **Volume**.
- Let s_y and s_x denote the sample standard deviation of y and x , respectively.
- The corresponding t-scores

$$y_t = \frac{y}{s_y} \quad \text{and} \quad x_t = \frac{x}{s_x}$$

are independent of the chosen measurement scale.

- The corresponding prediction equation is then

$$\hat{y}_t = \frac{a}{s_y} + \frac{s_x}{s_y} b x_t$$

- i.e. **the standardized regression coefficient** (slope) is

$$r = \frac{s_x}{s_y} b$$

which also is called **the correlation** between y and x .

-
- It can be shown that:
 - $-1 \leq r \leq 1$
 - The absolute value of r measures the (linear) strength of dependence between y and x .
 - When $r = 1$ all the points are on the prediction line, which has positive slope.
 - When $r = -1$ all the points are on the prediction line, which has negative slope.
 - To calculate the correlation:

```
trees.corr()
```

```
##           Girth   Height   Volume
## Girth    1.000000  0.51928  0.967119
## Height   0.519280  1.00000  0.598250
## Volume   0.967119  0.59825  1.000000
```

- There is a strong positive correlation between **Volume** and **Girth** ($r=0.967$).
- Note, it only works when all columns are numeric. Otherwise the relevant numeric columns should be extracted like this:

```
trees[['Height', 'Girth', 'Volume']].corr()
```

```
##           Height   Girth   Volume
## Height    1.00000  0.51928  0.59825
## Girth     0.51928  1.00000  0.96711
## Volume    0.59825  0.96711  1.00000
```

which produces the same output as above.

- Alternatively, one can calculate the correlation between two variables of interest like:

```
trees['Height'].corr(trees['Volume'])
```

```
## np.float64(0.5982496519917821)
```

2 R-squared: Reduction in prediction error

2.1 R-squared: Reduction in prediction error

- We want to compare two different models used to predict the response y .
- Model 1: We do not use the knowledge of x , and use $\bar{y}$ to predict any y -measurement. The corresponding prediction error is defined as

$$TSS = \sum_{i=1}^n (y_i - \bar{y})^2$$

and is called the **Total Sum of Squares**.

- Model 2: We use the prediction equation $\hat{y} = a + bx$ to predict y_i . The corresponding prediction error is then the Sum of Squared Errors

$$SSE = \sum_{i=1}^n (y_i - \hat{y}_i)^2.$$

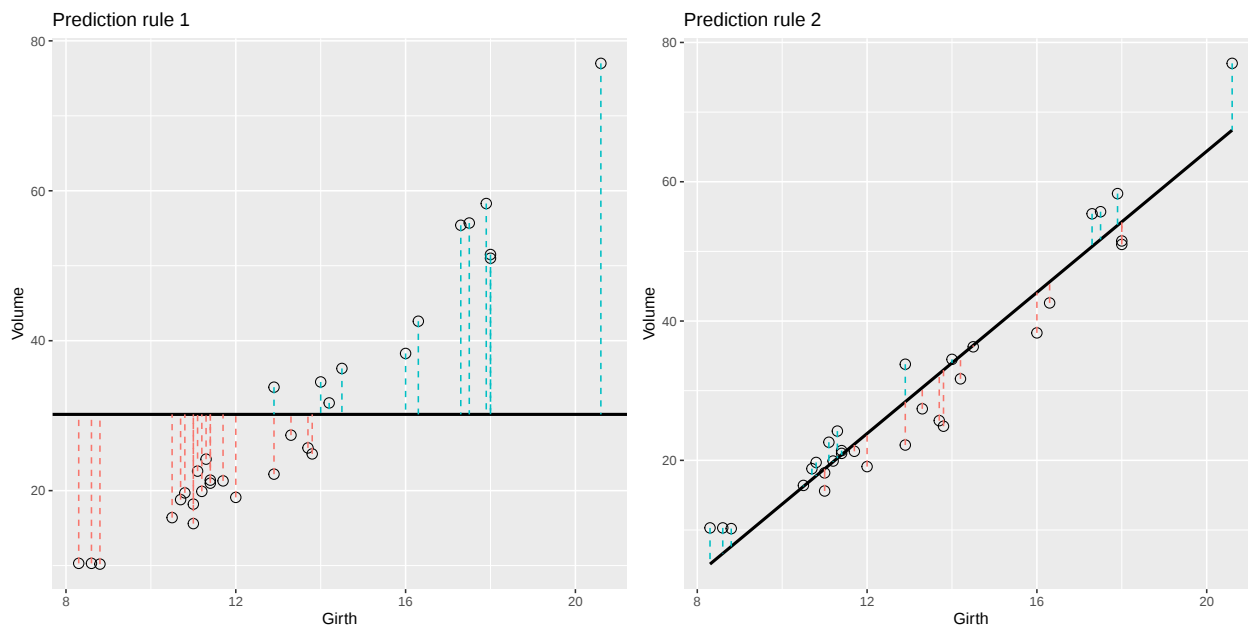
- We then define

$$r^2 = \frac{TSS - SSE}{TSS}$$

which can be interpreted as the relative reduction in the prediction error, when we include x as explanatory variable.

- This is also called the **fraction of explained variation, coefficient of determination** or simply **r-squared**.
- For example if $r^2 = 0.65$, the interpretation is that x explains about 65% of the variation in y , whereas the rest is due to other sources of random variation.

2.2 Graphical illustration of sums of squares



- Note the data points are the same in both plots. Only the prediction rule changes.
- The error of using Rule 1 is the total sum of squares $E_1 = TSS = \sum_{i=1}^n (y_i - \bar{y})^2$.
- The error of using Rule 2 is the residual sum of squares (sum of squared errors) $E_2 = SSE = \sum_{i=1}^n (y_i - \hat{y}_i)^2$.

2.3 r^2 : Reduction in prediction error

- For the simple linear regression we have that

$$r^2 = \frac{TSS - SSE}{TSS}$$

is equal to the square of the correlation between y and x , so it makes sense to denote it r^2 .

- At the top of the output below we can read off the value $R\text{-squared} = r^2 = 0.9353 = 93.53\%$, which is a large fraction of explained variation.

```
model.summary(slim = True)
```

```
## <class 'statsmodels.iolib.summary.Summary'>
## """
##                                     OLS Regression Results
## =====
## Dep. Variable:                    Volume    R-squared:                0.935
## Model:                            OLS      Adj. R-squared:          0.933
## No. Observations:                  31      F-statistic:             419.4
## Covariance Type:                  nonrobust Prob (F-statistic):       8.64e-19
## =====
##                coef      std err          t      P>|t|      [0.025      0.975]
## -----
## Intercept      -36.9435      3.365     -10.978      0.000     -43.826     -30.061
## Girth           5.0659      0.247      20.478      0.000       4.560       5.572
## =====
##
## Notes:
## [1] Standard Errors assume that the covariance matrix of the errors is correctly specified.
## """
```

3 Multiple regression model

3.1 Multiple regression model

- We look at data from Table 9.15 in Agresti. The data are measurements in the 67 counties of Florida.
- Our focus is on
 - The response y : **Crime** which is the crime rate
 - The predictor x_1 : **Education** which is proportion of the population with high school exam
 - The predictor x_2 : **Urbanisation** which is proportion of the population living in urban areas

3.2 Example

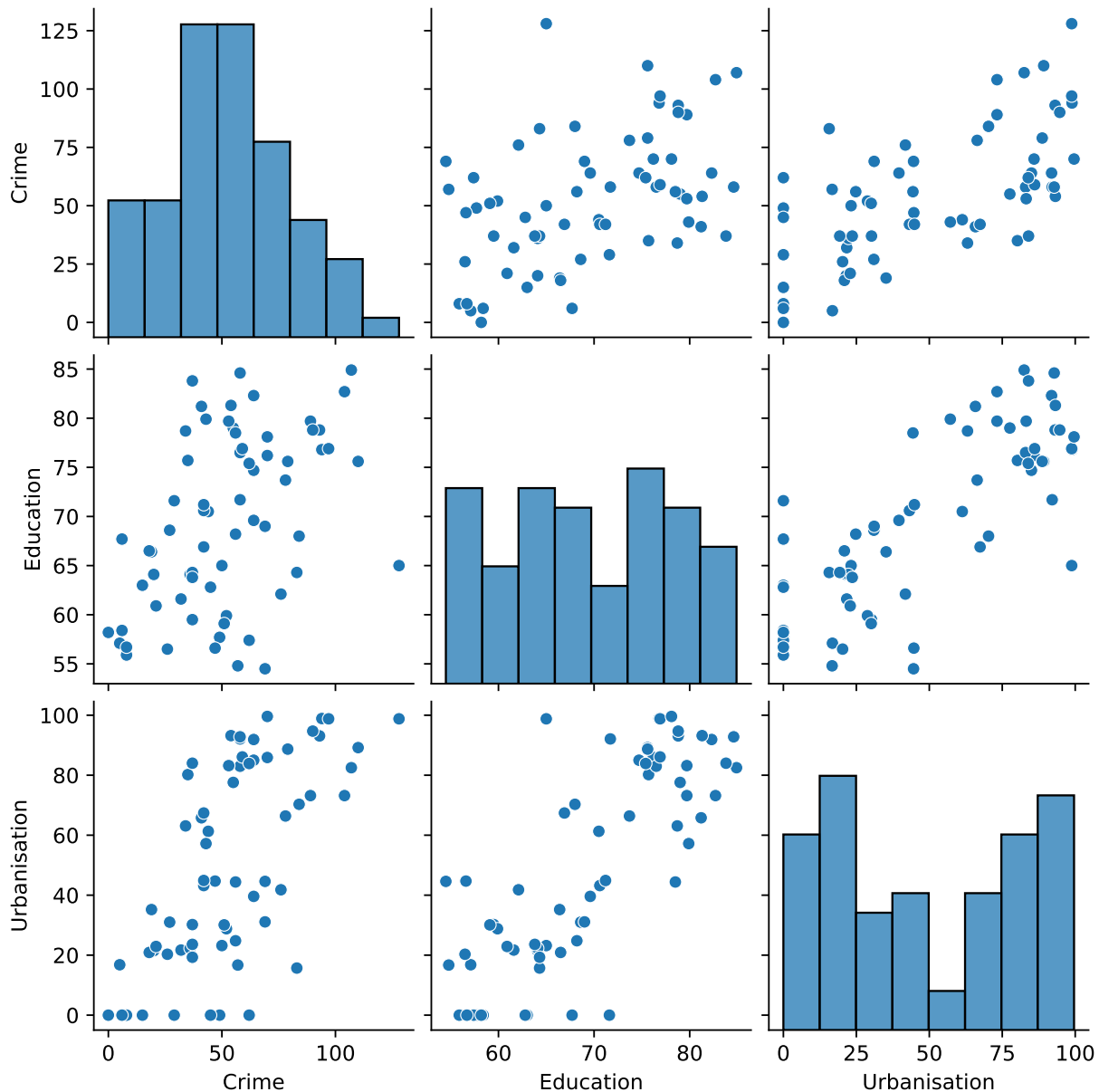
```
import pandas as pd

FL = pd.read_csv("https://asta.math.aau.dk/datasets?file=fl-crime.txt", sep='\t')
FL.head(n = 3)

##      Crime  Education  Urbanisation
## 0     104         82.7          73.2
## 1      20         64.1          21.5
## 2      64         74.7          85.0
```

```
import seaborn as sns
import matplotlib.pyplot as plt

p = sns.pairplot(FL)
```



3.3 Correlations

- There is significant ($p \approx 7 \times 10^{-5}$) positive correlation ($r=0.47$) between **Crime** and **Education**
- Then there is also significant positive correlation ($r=0.68$) between **Crime** and **Urbanisation**

```
FL.corr()
```

```
##           Crime  Education  Urbanisation
## Crime      1.000000   0.466912   0.677368
## Education   0.466912   1.000000   0.790719
```

```
## Urbanisation  0.677368    0.790719        1.000000
```

```
import pingouin as pg
```

```
pg.corr(FI['Crime'], FI['Education'], method='pearson')
```

```
##          n          r          CI95%          p-val          BF10          power
## pearson  67  0.466912  [0.26, 0.64]  0.000068  357.897  0.982794
```

3.4 Several predictors

- Both Education (x_1) and Urbanisation (x_2) are pretty good predictors for Crime (y).
- We therefore want to consider the model

$$y = \alpha + \beta_1 x_1 + \beta_2 x_2 + \epsilon$$

- The errors ϵ are random noise with mean zero and standard deviation $\sigma_{y|x}$.
- The graph for the mean response is in other words a 2-dimensional plane in a 3-dimensional space.
- We determine estimates (a, b_1, b_2) for (α, β_1, β_2) via the least squares method, i.e. deviations from the plane.

3.5 Example

```
import numpy as np
```

```
import statsmodels.formula.api as smf
```

```
model = smf.ols('Crime ~ Education + Urbanisation', data=FI).fit()
```

```
model.summary(slim = True)
```

```
## <class 'statsmodels.iolib.summary.Summary'>
```

```
## """
```

```
##                                OLS Regression Results
```

```
## =====
## Dep. Variable:                Crime    R-squared:                0.471
## Model:                      OLS      Adj. R-squared:          0.455
## No. Observations:           67       F-statistic:              28.54
## Covariance Type:            nonrobust  Prob (F-statistic):       1.38e-09
```

```
## =====
##               coef      std err          t      P>|t|      [0.025      0.975]
```

```
## -----
## Intercept          59.1181      28.365        2.084      0.041        2.452      115.784
## Education         -0.5834       0.472       -1.235      0.221       -1.527        0.360
## Urbanisation        0.6825       0.123        5.539      0.000        0.436        0.929
```

```
## =====
##
```

```
## Notes:
```

```
## [1] Standard Errors assume that the covariance matrix of the errors is correctly specified.
```

```
## """
```

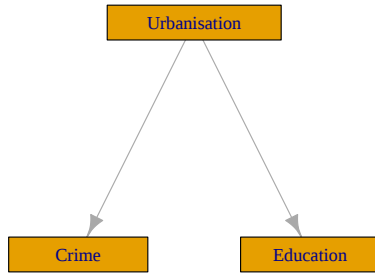
- From the output we find the prediction equation

$$\hat{y} = 59 - 0.58x_1 + 0.68x_2$$

- Not exactly what we expected based on the correlation.
- Now there appears to be a negative association between y and x_1 (Simpsons Paradox)!
- We can also find the standard error (0.4725) and the corresponding t-score (-1.235) for the the slope of Education
- This yields a p-value of 22%, i.e. the slope is not significantly different from zero.

3.6 Simpson's paradox

- The example illustrates **Simpson's paradox**.
- When considered alone **Education** is a good predictor for **Crime** (with positive correlation).
- When we add **Urbanisation**, then **Education** has a negative effect on **Crime** (but not significant).



- A possible explanation is illustrated by the graph above.
 - **Urbanisation** has positive effect on both **Education** and **Crime**.
 - For a given level of **urbanisation** there is a (non-significant) negative association between **Education** and **Crime**.
 - Viewed alone **Education** is a good predictor for **Crime**. If **Education** has a large value, then this indicates a large value of **Urbanisation** and thereby a large value of **Crime**.

4 The general model

4.1 Regression model

- We have a sample of size n , where we have measured
 - the response y .
 - k potential predictors $x_1, x_2, \dots, x_k$.
- Multiple regression model:
$$y = \alpha + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k + \epsilon.$$
- The errors ϵ are a sample from a population with mean zero and standard deviation $\sigma_{y|x}$.
- The **systematic** part of the model, i.e. when all errors are zero, says that **the mean response** is a linear function of the predictors:

$$E(y|x_1, x_2, \dots, x_k) = \alpha + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k$$

- The symbol E is used here to denote expectation, i.e., mean value.

4.2 Interpretation of parameters

- In the multiple linear regression model

$$E(y|x_1, x_2, \dots, x_k) = \alpha + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k$$

- The parameter α is the **Intercept**, corresponding to the mean response, when all predictors are equal to zero.
- The parameters $(\beta_1, \beta_2, \dots, \beta_k)$ are called **partial regression coefficients**.
- Imagine that all predictors but x_1 are held fixed. Then $y = \tilde{\alpha} + \beta_1 x_1$ is a line with slope β_1 , which describes the rate of change in the mean response, when x_1 is changed one unit. Here

$$\tilde{\alpha} = \alpha + \beta_2 x_2 + \dots + \beta_k x_k$$

is a constant number since we assumed all predictors but x_1 was held fixed.

- The rate of change β_1 does not depend on the value of the remaining predictors. In this case we say that there is **no interaction** between the effects of the predictors on the response.

- The above holds similarly for the other partial regression coefficients.
- An example of a model with interaction is

$$E(y|x_1, x_2) = \alpha + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_1 x_2 = \alpha + \beta_2 x_2 + (\beta_1 + \beta_3 x_2) x_1$$

- When we fix x_2 the line has slope $\beta_1 + \beta_3 x_2$, which depends on the chosen value of x_2 .

5 Estimation

5.1 Estimation of model

- The estimate $(a, b_1, b_2, \dots, b_k)$ for $(\alpha, \beta_1, \beta_2, \dots, \beta_k)$ is determined by minimizing the sum of squared errors.
- Based on this estimate we write the prediction equation as

$$\hat{y} = a + b_1 x_1 + b_2 x_2 + \dots + b_k x_k$$

- The distance between model and data is measured by the sum of squared errors

$$SSE = \sum_{i=1}^n e_i^2 = \sum_{i=1}^n (y_i - \hat{y}_i)^2.$$

- We estimate $\sigma_{y|x}$ by the quantity

$$s_{y|x} = \sqrt{\frac{SSE}{n - k - 1}}.$$

- Rather than n we divide SSE by **the degrees of freedom** $df = n - k - 1$. Theory shows, that this is reasonable.
- The degrees of freedom df are determined by the sample size minus the number of parameters in the regression equation.
- Currently we have $k + 1$ parameters: 1 intercept and k slopes.

6 Multiple R-squared

6.1 Multiple R^2

- We want to compare two models to predict the response y . Analogous to simple linear regression we have the following setup:
- Model 1: We do not use the predictors, and use $\bar{y}$ to predict any y -measurement. The corresponding prediction error is
 - $TSS = \sum_{i=1}^n (y_i - \bar{y})^2$ and is called the **Total Sum of Squares**.
- Model 2: We use the multiple prediction equation to predict y , i.e. the prediction error is
 - $SSE = \sum_{i=1}^n (y_i - \hat{y}_i)^2$ and is called **Sum of Squared Errors**.
- We then define **the multiple coefficient of determination**

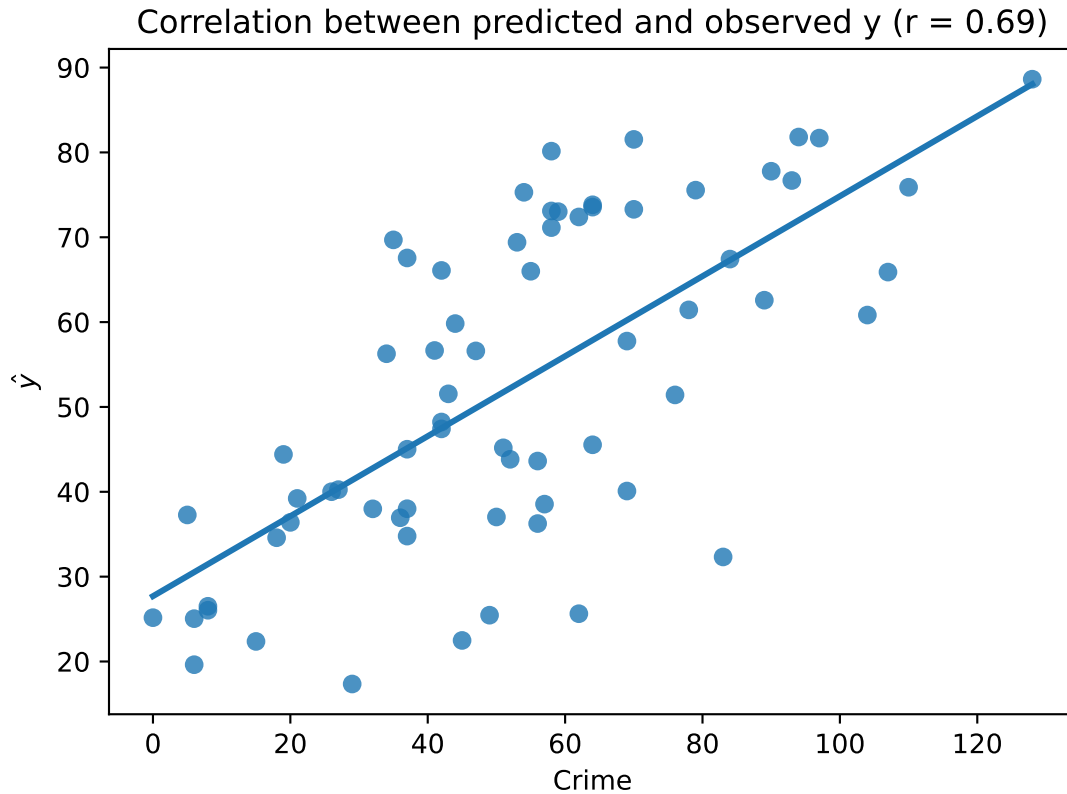
$$R^2 = \frac{TSS - SSE}{TSS}.$$

- Thus, R^2 is the relative reduction in prediction error, when we use $x_1, x_2, \dots, x_k$ as explanatory variables.
- It can be shown that the **multiple correlation** $R = \sqrt{R^2}$ is the correlation between y and $\hat{y}$.

```
FL['CrimePred'] = model.fittedvalues
r_val = model.rsquared**0.5
sns.regplot(x='Crime', y='CrimePred', data=FL, ci=None)

# Labels and title
```

```
plt.title(f"Correlation between predicted and observed y (r = {r_val:.2f})")
plt.xlabel("Crime")
plt.ylabel(r'$\hat{y}$') #r for raw string, no need to escape backslash in math
```



6.2 Example

```
model.summary(slim = True)
```

```
## <class 'statsmodels.iolib.summary.Summary'>
## """
##                      OLS Regression Results
##  =====
## Dep. Variable:          Crime    R-squared:                0.471
## Model:                  OLS      Adj. R-squared:           0.455
## No. Observations:       67      F-statistic:            28.54
## Covariance Type:        nonrobust Prob (F-statistic):      1.38e-09
##  =====
##                  coef    std err          t      P>|t|      [0.025    0.975]
##  -----
## Intercept            59.1181    28.365      2.084    0.041     2.452    115.784
## Education           -0.5834     0.472     -1.235    0.221    -1.527     0.360
## Urbanisation         0.6825     0.123      5.539    0.000     0.436     0.929
##  =====
##
## Notes:
```

```
## [1] Standard Errors assume that the covariance matrix of the errors is correctly specified.
## """
```

```
np.sqrt(model.mse_resid)
```

```
## np.float64(20.81558240979863)
```

```
model.df_resid
```

```
## np.float64(64.0)
```

- The prediction equation is $\hat{y} = 59 - 0.58x_1 + 0.68x_2$
- The estimate for $\sigma_{y|x}$ is $s_{y|x} = 20.82$ (Residual standard error) with $df = 67 - 3 = 64$ degrees of freedom.
- Multiple $R^2 = 0.4714$, i.e. 47% of the variation in the response is explained by including the predictors in the model.
- The estimate $b_1 = -0.5834$ has standard error (Std. Error) $se = 0.4725$ with corresponding t -score (t value) $t_{\text{obs}} = \frac{-0.5834}{0.4725} = -1.235$.
- The hypothesis $H_0 : \beta_1 = 0$ has the t -score $t_{\text{obs}} = -1.235$, which means that b_1 isn't significantly different from zero, since the p -value ($\Pr(>|t|)$) is 22%. That means that we should exclude Education as a predictor.

6.3 Example

- Our final model is then a simple linear regression:

```
model2 = smf.ols('Crime ~ Urbanisation', data=FL).fit()
model2.summary(slim = True)
```

```
## <class 'statsmodels.iolib.summary.Summary'>
## """
##                                OLS Regression Results
## =====
## Dep. Variable:                Crime    R-squared:                0.459
## Model:                      OLS      Adj. R-squared:           0.451
## No. Observations:           67       F-statistic:              55.11
## Covariance Type:            nonrobust Prob (F-statistic):       3.08e-10
## =====
##              coef    std err          t      P>|t|      [0.025     0.975]
## -----
## Intercept          24.5412      4.539      5.406      0.000      15.476      33.607
## Urbanisation         0.5622      0.076      7.424      0.000         0.411         0.713
## =====
##
## Notes:
## [1] Standard Errors assume that the covariance matrix of the errors is correctly specified.
## """
```

- The coefficient of determination always decreases, when the model is simpler. Now we have $R^2 = 46\%$, where before we had 47%. But the decrease is not significant.

7 F-test for effect of predictors

7.1 F-test

- We consider the hypothesis

$$H_0 : \beta_1 = \beta_2 = \dots = \beta_k = 0$$

against the alternative, that at least one of these are non-zero.

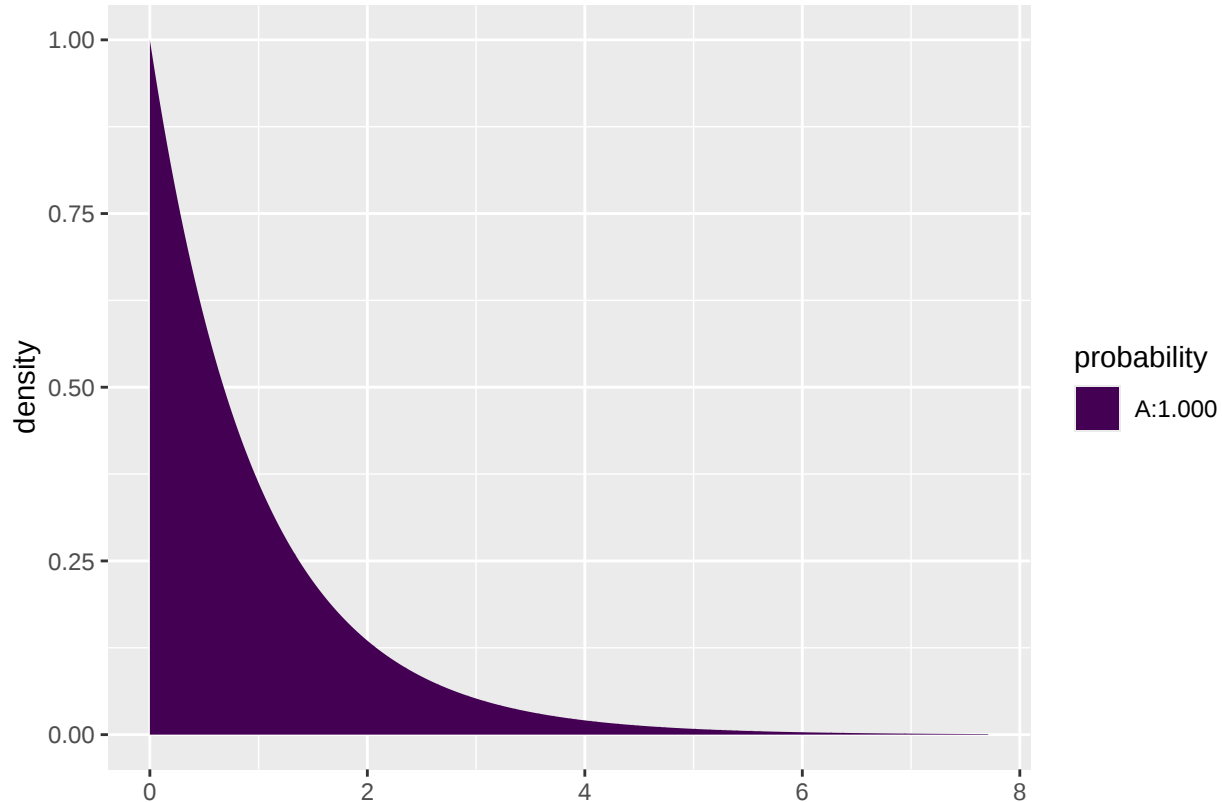
- As test statistic we use

$$F_{obs} = \frac{(n - k - 1)R^2}{k(1 - R^2)}$$

- Large values of R^2 implies large values of F , which points to the alternative hypothesis.
- I.e. when we have calculated the observed value F_{obs} , then we have to find the probability that a new experiment would result in a larger value.
- It can be shown that the reference distribution is (can be approximated by) a so-called **F-distribution** with **degrees of freedom** $df_1 = k$ and $df_2 = n - k - 1$.

7.2 Example

- We return to **Crime** and the prediction equation $\hat{y} = 59 - 0.58x_1 + 0.68x_2$, where $n = 67$ and $R^2 = 0.4714$. We have
 - $df_1 = k = 2$ since we have 2 predictors.
 - $df_2 = n - k - 1 = 67 - 2 - 1 = 64$.
 - Then we can calculate $F_{obs} = \frac{(n-k-1)R^2}{k(1-R^2)} = 28.54$
- To evaluate the value 28.54 in the relevant F-distribution:



```
from scipy.stats import f
```

```
1 - f.cdf(28.54, 2, 64)
```

```
## np.float64(1.3786117802894182e-09)
```

- So $p\text{-value} = 1.38 \times 10^{-9}$ (notice we don't multiply by 2 since this is a one-sided test; only large values point more towards the alternative than the null hypothesis).
- All this can be found in the summary output we already have:

```
model.summary(slim = True)
```

```
## <class 'statsmodels.iolib.summary.Summary'>
## """
##                               OLS Regression Results
## =====
## Dep. Variable:                Crime   R-squared:                0.471
## Model:                      OLS     Adj. R-squared:          0.455
## No. Observations:           67      F-statistic:             28.54
## Covariance Type:            nonrobust Prob (F-statistic):       1.38e-09
## =====
##                               coef      std err          t      P>|t|      [0.025      0.975]
## -----
## Intercept                   59.1181     28.365        2.084    0.041     2.452    115.784
## Education                   -0.5834      0.472       -1.235    0.221    -1.527     0.360
## Urbanisation                 0.6825      0.123        5.539    0.000     0.436     0.929
## =====
##
## Notes:
## [1] Standard Errors assume that the covariance matrix of the errors is correctly specified.
## """
```

8 Test for interaction

8.1 Interaction between effects of predictors

- Could it be possible that a combination of Education and Urbanisation is good for prediction? We want to investigate this using the model

$$E(y|x_1, x_2) = \alpha + \beta_1 x_1 + \beta_2 x_2 + \beta_3 x_1 x_2,$$

where we have extended with a possible effect of the product $x_1 x_2$:

```
model3 = smf.ols('Crime ~ Education * Urbanisation', data=FL).fit()
model3.summary(slim = True)
```

```
## <class 'statsmodels.iolib.summary.Summary'>
## """
##                               OLS Regression Results
## =====
## Dep. Variable:                Crime   R-squared:                0.479
## Model:                      OLS     Adj. R-squared:          0.454
## No. Observations:           67      F-statistic:             19.32
## Covariance Type:            nonrobust Prob (F-statistic):       5.37e-09
## =====
##                               coef      std err          t      P>|t|      [0.025      0.975]
## -----
## Intercept                   19.3175     49.959        0.387    0.700    -80.517    119.152
## Education                   0.0340      0.794        0.043    0.966    -1.552     1.620
## Urbanisation                 1.5143      0.868        1.744    0.086    -0.220     3.249
## Education:Urbanisation     -0.0120      0.012       -0.968    0.337    -0.037     0.013
## =====
##
## Notes:
## [1] Standard Errors assume that the covariance matrix of the errors is correctly specified.
## """
```

```
## [2] The condition number is large, 8.97e+04. This might indicate that there are
## strong multicollinearity or other numerical problems.
## ""
```

- When we look at the p -values in the table nothing is significant at the 5% level!
- But the F-statistic tells us that the predictors collectively have a significant prediction ability.
- Why has the highly significant effect of x_2 disappeared? Because the predictors x_1 and x_1x_2 are able to explain the same as x_2 .
- Previously we only had x_1 as alternative explanation to x_2 - and that wasn't enough.
- The phenomenon is called **multicollinearity** and illustrates that we can have different models with equally good predictive properties.
- In this case we will choose the model with x_2 since it is simpler.
- However, in general it can be difficult to choose between models. For example, if both height and weight are good predictors of some response, but one of them can be left out, which one do we choose?