

Vorkurs Ökonometrie

für den Masterstudiengang
Wirtschaftswissenschaften

Vorlesung

Übersicht:

Definition des einfachen Regressionsmodells

Herleitung der gewöhnlichen Kleinste-Quadrate-Schätzer

OLS: Eigenschaften bei beliebigen Stichproben

Maßeinheiten und die funktionale Form

Erwartungswerte und Varianzen der OLS-Schätzer

Definition des einfachen Regressionsmodells

Definition des einfachen Regressionsmodells

Das einfache Regressionsmodell erklärt die Variable y durch die Variable x .

Drei Aspekte:

- ▶ Wie modellieren wir die Wirkung von anderen Faktoren auf y ?
- ▶ Was ist die funktionale Beziehung zwischen y und x ?
- ▶ Wie können wir ein Ceteris-Paribus Zusammenhang sicherstellen?

$$y = \beta_0 + \beta_1 x + u$$

- y : Abhängige/erklärte Variable, Regressand
- β_0 : Achsenabschnitt (unbekannter Parameter)
- β_1 : Steigung (unbekannter Parameter)
- x : Unabhängige/erklärende Variable, Regressor
- u : Störterm (u : unbeobachtbar)

Interpretation des einfachen Regressionsmodells

$$y = \beta_0 + \beta_1 x + u$$

Das einfache Regressionsmodell erklärt inwiefern y in Veränderungen von x variiert.

Ceteris Paribus-Annahme:

Nehme an, dass sich x verändert ($\Delta x \neq 0$) und alle anderen erklärenden aber unbeobachtbaren Größen konstant bleiben ($\Delta u = 0$).

Dann gilt

$$\frac{\Delta y}{\Delta x} = \beta_1$$

Beispiel: Sojabohnen und Dünger

„Wie wirkt sich die Düngemenge auf den Ernteerfolg aus?“

$$yield = \beta_0 + \beta_1 fertilizer + u ,$$

sodass $y = yield$ und $x = fertilizer$.

Der Störterm u beinhaltet andere Größen wie Landqualität, Niederschlag, Parasiten.

Unter der Ceteris Paribus-Annahme gilt:

$$\Delta yield = \beta_1 \Delta fertilizer$$

Beispiel: Eine einfache Lohngleichung

„Wie wirkt sich Bildung auf den Stundenlohn aus?“

$$wage = \beta_0 + \beta_1 educ + u ,$$

sodass $y = wage$ und $x = educ$.

Der Störterm u beinhaltet andere Größen wie Berufserfahrung, Talent, etc.

Der unbekannte Parameter β_1 misst (unter der Ceteris Paribus-Annahme), wie sich ein zusätzliches Bildungsjahr auf den Stundenlohn auswirkt.

Diskussion der Linearität des einfachen Regressionsmodells

Die Gleichung

$$y = \beta_0 + \beta_1 x + u$$

impliziert, dass die Variable x immer den gleichen (linearen) Effekt auf y hat.

Diese Unzulänglichkeit kann durch Transformation der Variablen geheilt werden.

Der schwierigste Aspekt ist die Ceteris Paribus Annahme.

Um diese Annahme zu rechtfertigen interpretieren wir zunächst x und u (und damit auch y) als Zufallsvariablen.

Nun können wir eine statistische Annahme an x und u treffen.

Strikte Exogenität

In Kenntnis der erklärenden Variable x gibt es keine zusätzliche Information über den Erwartungswert der unbeobachtbaren Variable u :

$$E[u|x] = E[u]$$

Im Düngerbeispiel:

Die unterschiedlichen Düngermengen werden zufällig auf die einzelnen Ackerflächen verteilt. Demnach enthält die Düngermenge keine Information über die unbeobachtbare Bodenqualität der einzelnen Ackerfläche. (✓)

Im Bildungsbeispiel:

Die unterschiedlichen Bildungsniveaus können durchaus von den zugrundeliegenden intellektuellen Fähigkeiten abhängen. Daher enthält das Bildungsniveau (statistische) Informationen über die zu erwartenden unbeobachtbaren Fähigkeiten. (!)

Annahme: Strikte Exogenität

Im ersten Teil der Vorlesung werden wir die Annahme der **strikten Exogenität der Regressoren** (auch: mean independence) aufrechterhalten:

$$E[u|x] = E[u]$$

Erst in späteren Kapiteln gehen wir genauer auf die Implikationen und Lösungen bei einer Verletzungen dieser Annahme ein.

Ohne Beschränkung der Allgemeinheit: $E[u] = 0$

$$y = \beta_0 + \beta_1 x + u$$

Wir nehmen des Weiteren an, dass der Erwartungswert des unbeobachtbaren Störterms null beträgt:

$$E[u] = 0$$

Diese Annahme erfolgt jedoch ohne Beschränkung der Allgemeinheit!

Wäre $E[u] \neq 0$, so könnten wir das Modell mit $\tilde{u} = u - E[u]$ und $\tilde{\beta}_0 = \beta_0 + E[u]$ modifizieren.

Für das modifizierte Modell

$$y = \tilde{\beta}_0 + \beta_1 x + \tilde{u}$$

wäre $E[\tilde{u}] = E[u - E[u]] = E[u] - E[u] = 0$ erfüllt.

Annahme: Bedingter Null-Erwartungswert

Engl.: zero conditional mean assumption

Mit den beiden Annahmen

$$E[u] = 0$$

und

$$E[u|x] = E[u]$$

folgt direkt:

$$E[u|x] = 0$$

Oft wird direkt diese Gleichung (ohne Vordiskussion) angenommen.

Population Regression Function (PRF)

Bilden wir den bedingten Erwartungswert des Regressionsmodells

$$y = \beta_0 + \beta_1 x + u$$

erhalten wir

$$\begin{aligned} E[y|x] &= E[\beta_0 + \beta_1 x + u|x] \\ &= E[\beta_0|x] + E[\beta_1 x|x] + E[u|x] \end{aligned}$$

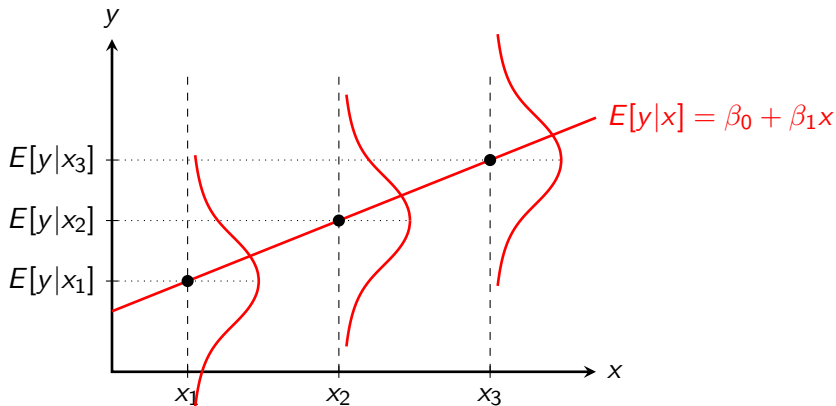
und mit $E[u|x] = 0$ also

$$E[y|x] = \beta_0 + \beta_1 x$$

Eine Erhöhung von x um eine Einheit erhöht den **bedingten Erwartungswert** von y um β_1 Einheiten.

Population Regression Function (PRF)

Für einen gegebenen Wert x ist $E[y|x] = \beta_0 + \beta_1 x$ der Gravitationspunkt der bedingten Dichte von y .



Die Gleichung $E[y|x] = \beta_0 + \beta_1 x$ beschreibt den erwarteten Wert von y . Der tatsächliche Wert von y streut um $E[y|x]$.

Herleitung der gewöhnlichen Kleinste-Quadrate-Schätzer

Herleitung von gewöhnlichen Schätzern

Um die unbekannten Parameter β_0 und β_1 des einfachen linearen Regressionsmodells zu schätzen, benötigen wir neben den Annahmen

$$y = \beta_0 + \beta_1 x + u$$

und

$$E[u|x] = 0$$

Daten: die **Zufallsstichprobe**

(x_1, y_1) : Erste Beobachtung

(x_2, y_2) : Zweite Beobachtung

$\vdots$

(x_n, y_n) : n -te Beobachtung

Herleitung von gewöhnlichen Schätzern

Sei

$$\{(x_i, y_i) \in \mathbb{R}^2 : i = 1, \dots, n\}$$

eine Zufallsstichprobe des Umfangs $n \in \mathbb{N}$.

Für $i = 1, \dots, n$:

x_i : Wert der erklärenden Variable der i -ten Beobachtung

y_i : Wert der erklärten Variable der i -ten Beobachtung

Das Regressionsmodell lässt sich nun schreiben als

$$y_i = \beta_0 + \beta_1 x_i + u_i, \quad i = 1, \dots, n$$

u_i : Störterm für Beobachtung i .

Enthält alle Einflussfaktoren auf y_i , die nicht in x_i enthalten sind.

Die Momentenmethode (Method of Moments, MM)

Seien n Paare von Zufallsvariablen (A_i, B_i) unabhängig und identisch verteilt mit theoretischen Momenten $E[A|\theta]$, $E[B|\theta]$, $E[AB|\theta]$, $E[A^2|\theta]$ und $E[B^2|\theta]$, wobei θ ein Vektor von unbekannten Parametern darstellt.

Dann gilt mit dem Gesetz der großen Zahlen

$$\frac{1}{n} \sum_{i=1}^n A_i \xrightarrow[n \rightarrow \infty]{} E[A|\theta], \quad \frac{1}{n} \sum_{i=1}^n A_i B_i \xrightarrow[n \rightarrow \infty]{} E[AB|\theta], \quad \text{usw.}$$

Die Momentenmethode approximiert nun die theoretischen Momente durch die empirischen Momente und schätzt hier die unbekannten Parameter θ durch die implizit definierten Schätzer $\hat{\theta}$:

$$\frac{1}{n} \sum_{i=1}^n A_i = E[A|\hat{\theta}], \quad \frac{1}{n} \sum_{i=1}^n A_i B_i = E[AB|\hat{\theta}], \quad \text{usw.}$$

Die Momentenmethode mit $E[u|x] = 0$

Für $y_i = \beta_0 + \beta_1 x_i + u_i$, $i = 1, \dots, n$, sei angenommen, dass u_i unabhängig und identisch verteilt seien mit $E[u_i|x_i] = 0$.

Zunächst schreiben wir $u_i = y_i - \beta_0 - \beta_1 x_i$.

Das theoretische Moment $E[u_i|x_i] = E[y_i|x_i] - \beta_0 - \beta_1 x_i$ enthält die unbekannten Parameter β_0 und β_1 . Für diese definieren wir die Schätzer $\hat{\beta}_0$ und $\hat{\beta}_1$ implizit durch die Schätzung des theoretischen Moments durch das empirische Moment

$$\frac{1}{n} \sum_{i=1}^n u_i \approx \frac{1}{n} \sum_{i=1}^n y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i = 0 = E[u_i|x_i]$$

Mit $\bar{y} := \frac{1}{n} \sum_{i=1}^n y_i$ und $\bar{x} := \frac{1}{n} \sum_{i=1}^n x_i$:

$$\Rightarrow \bar{y} = \hat{\beta}_0 + \hat{\beta}_1 \bar{x}$$

Die Momentenmethode mit $E[u|x] = 0$

Für das Moment $E[xu]$ gilt aufgrund der Annahme $E[u|x] = 0$:

$$E[xu] = E[x \underbrace{E[u|x]}_{=0}] = 0$$

Die Eigenschaft $E[xu] = 0$ wird auch mit **Orthogonalität** von x und u bezeichnet.

Mit $u_i = y_i - \beta_0 - \beta_1 x_i$ schätzen wir $E[xu]$ durch

$$\frac{1}{n} \sum_{i=1}^n x_i u_i \approx \frac{1}{n} \sum_{i=1}^n x_i (y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i) = E[xu] = 0$$

Mit $\overline{xy} := \frac{1}{n} \sum_{i=1}^n x_i y_i$ und $\overline{xx} := \frac{1}{n} \sum_{i=1}^n x_i^2$:

$$\Rightarrow \overline{xy} = \hat{\beta}_0 \bar{x} + \hat{\beta}_1 \overline{xx}$$

Die Momentenmethode mit $E[u|x] = 0$

Die Momentenmethode liefert also über die Annahme $E[u|x] = 0$ zwei Gleichungen für die beiden Schätzer $\hat{\beta}_0$ und $\hat{\beta}_1$:

$$\bar{y} = \hat{\beta}_0 + \hat{\beta}_1 \bar{x}$$

und

$$\overline{xy} = \hat{\beta}_0 \bar{x} + \hat{\beta}_1 \overline{xx}$$

Unter der Voraussetzung $\overline{xx} \neq \bar{x}\bar{x}$ ergibt sich:

$$\hat{\beta}_0 = \bar{y} - \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}} \bar{x}$$

und

$$\hat{\beta}_1 = \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}}$$

Die gewöhnliche Kleinste-Quadrate-Methode (OLS)

Eine andere Schätzmöglichkeit besteht in der Vorgabe einer Zielfunktion.

Die gebräuchlichste ist:

Minimiere die Summe der quadrierten Abweichungen zwischen der erklärten Variablen und ihren durch die Schätzung vorhergesagten Werten.

Dies ist das gewöhnliche Kleinste-Quadrate-Prinzip (**OLS, ordinary least squares**).

Die gewöhnliche Kleinste-Quadrate-Methode (OLS)

Schätze die unbekannten Parameter β_0 und β_1 des Regressionsmodells

$$y_i = \beta_0 + \beta_1 x_i + u_i, \quad i = 1, \dots, n$$

zunächst durch beliebige Zahlen b_0 und b_1 .

→ Prognostizierte Werte $\hat{y}_i$ (**fitted values**) mit

$$\hat{y}_i = b_0 + b_1 x_i, \quad i = 1, \dots, n$$

Definiere **Residuen** $\hat{u}_i$ mit

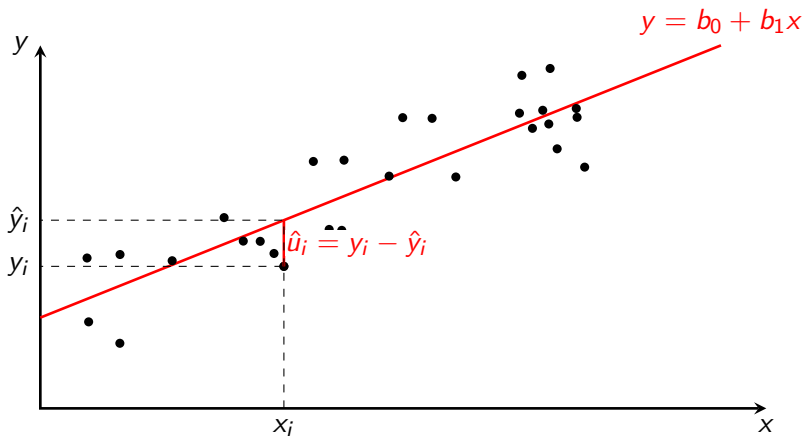
$$\hat{u}_i = y_i - \hat{y}_i = y_i - b_0 - b_1 x_i$$

Je besser die Zahlen b_0 und b_1 , desto „kleiner“ die Residuen.

Prognose und Residuen

Vermutete Parameter: b_0, b_1

Schätzfehler (hängt von b_0, b_1 ab): $\hat{u}_i = y_i - \hat{y}_i = y_i - b_0 - b_1 x_i$



Die gewöhnliche Kleinste-Quadrate-Methode (OLS)

Der Vektor $\hat{u} = (\hat{u}_1, \hat{u}_2, \dots, \hat{u}_n)$ hat die Länge

$$\|\hat{u}\| = \sqrt{\hat{u}_1^2 + \hat{u}_2^2 + \dots + \hat{u}_n^2}$$

Wähle nun b_0 und b_1 , sodass die Norm $\|\hat{u}\|$ minimal wird!

$\Rightarrow$ Zielfunktion $S(b_0, b_1) = \sum_{i=1}^n (y_i - b_0 - b_1 x_i)^2$

$$\min_{b_0, b_1 \in \mathbb{R}} \sum_{i=1}^n (y_i - b_0 - b_1 x_i)^2$$

Die gewöhnliche Kleinste-Quadrate-Methode (OLS)

$$\min_{b_0, b_1 \in \mathbb{R}} \sum_{i=1}^n (y_i - b_0 - b_1 x_i)^2$$

Die optimalen Schätzer $\hat{\beta}_0$ und $\hat{\beta}_1$ erfüllen die Bedingungen erster Ordnung:

$$\sum_{i=1}^n 2(y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i)(-1) = 0$$

$$\sum_{i=1}^n 2(y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i)(-x_i) = 0$$

Hessematrix:

$$H = \begin{pmatrix} 2n & 2 \sum_{i=1}^n x_i \\ 2 \sum_{i=1}^n x_i & 2 \sum_{i=1}^n x_i^2 \end{pmatrix} = 2n \begin{pmatrix} 1 & \bar{x} \\ \bar{x} & \overline{xx} \end{pmatrix}$$

Die gewöhnliche Kleinste-Quadrate-Methode (OLS)

Die Hessematrix ist positiv definit, falls $\overline{xx} > \bar{x}\bar{x}$.

Die Bedingungen erster Ordnung sind äquivalent zu

$$\begin{aligned}\bar{y} &= \hat{\beta}_0 + \hat{\beta}_1 \bar{x} \\ \overline{xy} &= \hat{\beta}_0 \bar{x} + \hat{\beta}_1 \overline{xx}\end{aligned}$$

Unter der Voraussetzung $\overline{xx} > \bar{x}\bar{x}$ ergibt sich also auch hier:

$$\begin{aligned}\hat{\beta}_0 &= \bar{y} - \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}} \bar{x} \\ \text{und} \\ \hat{\beta}_1 &= \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}}\end{aligned}$$

Orthogonale Projektion

Das Regressionsmodell

$$y_i = \beta_0 + \beta_1 x_i + u_i, \quad i = 1, \dots, n$$

liefert mit $E[u|x] = 0$ die funktionale Beziehung

$$E[y_i|x] = \beta_0 + \beta_1 x_i, \quad i = 1, \dots, n$$

Geometrische Interpretation:

$$\begin{pmatrix} E[y_1|x] \\ E[y_2|x] \\ \vdots \\ E[y_n|x] \end{pmatrix} = \beta_0 \begin{pmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{pmatrix} + \beta_1 \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$$

Der Vektor $E[\mathbf{y}|x]$ ist eine Linearkombination des Einsvektors $\mathbf{1}$ und des Regressorvektors $\mathbf{x}$.

Orthogonale Projektion

Falls die beiden n -Vektoren $\boldsymbol{\iota}$ und $\mathbf{x}$ linear unabhängig sind, spannen sie eine Ebene im n -Dimensionalen Raum auf.

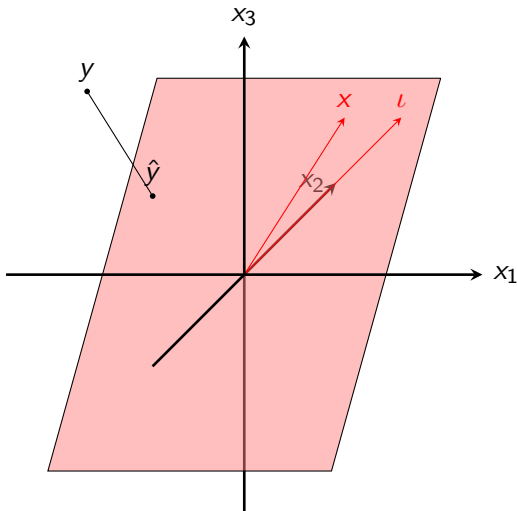
Die n -Vektoren $\boldsymbol{\iota}$ und $\mathbf{x}$ sind genau dann linear unabhängig, falls es kein $\lambda \in \mathbb{R}$ gibt, sodass $\mathbf{x} = \lambda \boldsymbol{\iota}$. Dies ist äquivalent zu $\overline{\mathbf{x}\mathbf{x}} \neq \overline{\mathbf{x}}\overline{\mathbf{x}}$.

Jedes Paar (b_0, b_1) definiert einen Punkt $b_0\boldsymbol{\iota} + b_1\mathbf{x}$ auf der von $\boldsymbol{\iota}$ und $\mathbf{x}$ aufgespannten Ebene.

Gesucht ist derjenige Punkt $\hat{\mathbf{y}} = \hat{\beta}_0\boldsymbol{\iota} + \hat{\beta}_1\mathbf{x}$, welcher in der Ebene am nächsten zum Vektor $\mathbf{y}$ ist.

Dieser Punkt $\hat{\mathbf{y}}$ ist die orthogonale Projektion von $\mathbf{y}$ auf die durch $\boldsymbol{\iota}$ und $\mathbf{x}$ aufgespannte Ebene.

Orthogonale Projektion



Orthogonale Projektion

Der Vektor von $\hat{\mathbf{y}}$ zu $\mathbf{y}$, also $\mathbf{y} - \hat{\mathbf{y}}$ muss demnach orthogonal zu den Vektoren $\boldsymbol{\iota}$ und $\mathbf{x}$ sein.

Es muss also gelten:

$$\boldsymbol{\iota}'(\mathbf{y} - \hat{\mathbf{y}}) = \sum_{i=1}^n 1(y_i - \hat{y}_i) = \sum_{i=1}^n \hat{u}_i = 0$$

$$\mathbf{x}'(\mathbf{y} - \hat{\mathbf{y}}) = \sum_{i=1}^n x_i(y_i - \hat{y}_i) = \sum_{i=1}^n x_i \hat{u}_i = 0$$

Wie wir wissen implizieren diese beiden Bedingungen

$$\hat{\beta}_0 = \bar{y} - \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}} \bar{x}$$

und

$$\hat{\beta}_1 = \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}}$$

Diskussion: $\overline{xx} > \bar{x}\bar{x}$

Der Ausdruck $\overline{xx} - \bar{x}\bar{x}$ ist die **Stichprobenvarianz von x** :

$$\overline{xx} - \bar{x}\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i^2 - \left(\frac{1}{n} \sum_{i=1}^n x_i \right)^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$$

Die Stichprobenvarianz von x_i ist positiv, falls mindestens eine Beobachtung x_i vom Stichprobendurchschnitt $\bar{x}$ abweicht, falls also nicht alle Beobachtungen x_i identisch sind.

Wären alle Beobachtungen x_i identisch, hätte $\mathbf{x}$ keinen Erklärungsgehalt für die Variation von $\mathbf{y}$.

Diskussion: $\overline{xy} - \bar{x}\bar{y}$

$\overline{xy} - \bar{x}\bar{y}$ ist die **Stichprobenkovarianz von x und y**:

$$\overline{xy} - \bar{x}\bar{y} = \frac{1}{n} \sum_{i=1}^n x_i y_i - \left(\frac{1}{n} \sum_{i=1}^n x_i \right) \left(\frac{1}{n} \sum_{i=1}^n y_i \right) = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})$$

Wäre y_i eine lineare Funktion von x_i , also z.B.

$$y_i = \beta_0 + \beta_1 x_i ,$$

dann würde gelten: $\bar{y} = \beta_0 + \beta_1 \bar{x}$ und

$$\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) = \beta_1 \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = \beta_1 (\overline{xx} - \bar{x}\bar{x})$$

und

$$\hat{\beta}_1 = \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}} = \beta_1 \frac{\overline{xx} - \bar{x}\bar{x}}{\overline{xx} - \bar{x}\bar{x}} = \beta_1$$

Der Korrelationskoeffizient

Der (theoretische) Korrelationskoeffizient:

$$\rho_{xy} = \frac{\text{Cov}(x, y)}{\sqrt{\text{Var}(x)}\sqrt{\text{Var}(y)}}$$

Der empirische Korrelationskoeffizient ersetzt die theoretischen Momente durch die empirischen Momente:

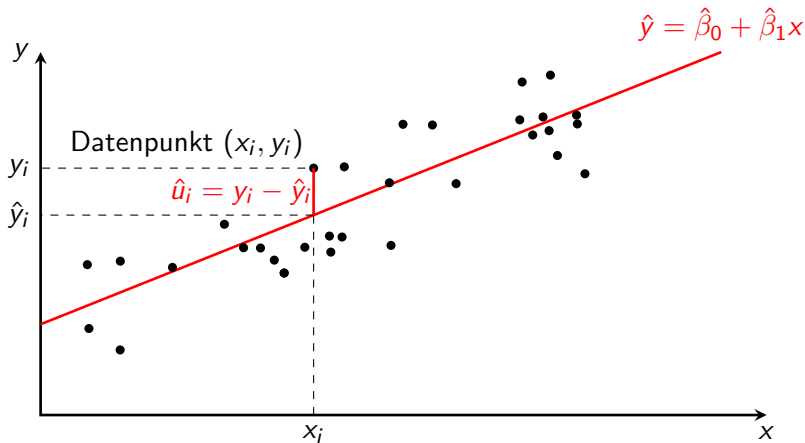
$$\hat{\rho}_{x,y} := \frac{\overline{xy} - \bar{x}\bar{y}}{\sqrt{(\overline{xx} - \bar{x}\bar{x})}\sqrt{(\overline{yy} - \bar{y}\bar{y})}}$$

Mit $\hat{\beta}_1 = \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}}$ ergibt sich

$$\hat{\beta}_1 = \hat{\rho}_{x,y} \frac{\sqrt{\overline{yy} - \bar{y}\bar{y}}}{\sqrt{\overline{xx} - \bar{x}\bar{x}}}$$

Die gewöhnliche Regressionsanalyse misst also nur Korrelation und nicht Kausalität!

Die OLS Regressionsgerade



Die Regressionsgerade $\hat{y} = \hat{\beta}_0 + \hat{\beta}_1 x$ schätzt die unbekannte Population Regression Function $E[y|x] = \beta_0 + \beta_1 x$.

Beispiele der einfachen Regression

CEO-Gehälter und Eigenkapitalrentabilität

Population Regression Function

$$salary = \beta_0 + \beta_1 roe + u$$

salary : Gehalt in tausend \$

roe : durchschnittliche Eigenkapitalrentabilität (return on equity)

Regressionsgerade:

$$\widehat{salary} = 963.191 + 18.501 roe$$

Datensatz: CEOSAL1 ($n = 209$)

Die Regressionsgerade hängt von der Stichprobe ab.

Die Population Regression Function ist unbekannt!



Beispiele der einfachen Regression

Stundenlohn und Bildung

Population Regression Function

$$wage = \beta_0 + \beta_1 educ + u$$

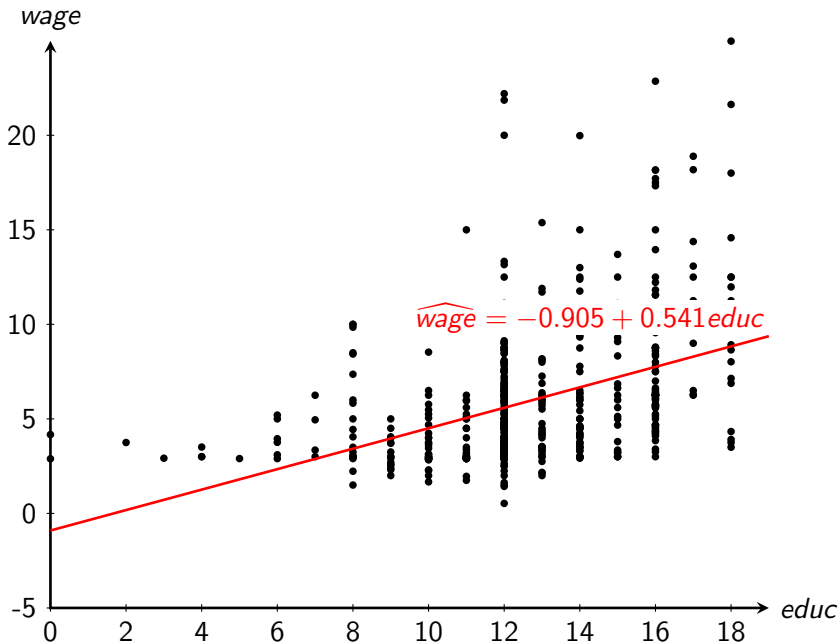
wage : Stundenlohn in \$

educ : Bildungsjahre

Regressionsgerade:

$$\widehat{wage} = -0.90 + 0.54educ$$

Datensatz: WAGE1 ($n = 526$)



Beispiele der einfachen Regression

Wahlergebnisse und Wahlkampfausgaben

Population Regression Function

$$voteA = \beta_0 + \beta_1 shareA + u$$

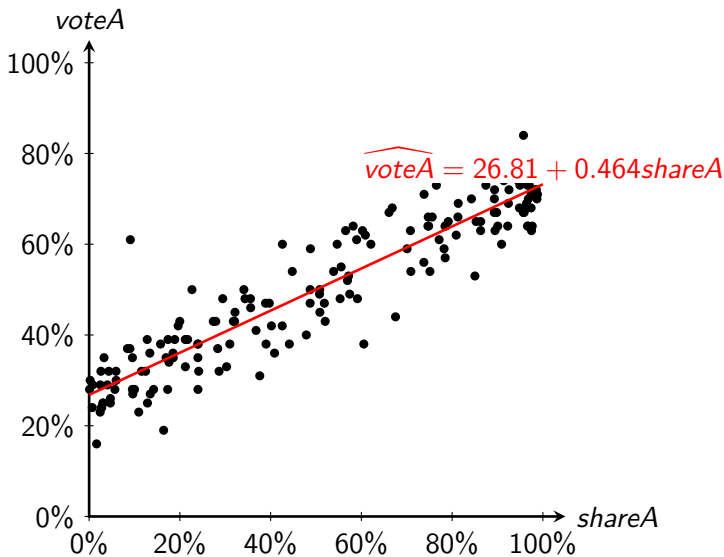
voteA : % Kandidat:in A

shareA : Anteil Wahlkampfausgaben für Kandidat:in A

Regressionsgerade:

$$\widehat{voteA} = 26.81 + 0.464 shareA$$

Datensatz: VOTE1 ($n = 173$)



OLS: Eigenschaften bei beliebigen Stichproben

Prognose und Residuen

Für eine beliebige¹ Stichprobe $\{(x_i, y_i)\}_{i=1}^n$ liefern die OLS-Schätzer $\hat{\beta}_0$ und $\hat{\beta}_1$ die Prognose

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

und damit die Residuen

$$\hat{u}_i = y_i - \hat{y}_i = y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i$$

für $i = 1, \dots, n$

¹Die Stichprobe muss verschiedene Werte x_i enthalten.

Beispiel: Managergehälter (Daten: CEOSAL1)

obsno	roe	salary	$\widehat{\text{salary}}$	$\hat{u}$
1	14.1	1095	1224.058	-129.0581
2	10.9	1001	1164.854	-163.8542
3	23.5	1122	1397.969	-275.9692
4	5.9	578	1072.348	-494.3484
5	13.8	1368	1218.508	149.4923
6	20.0	1145	1333.215	-188.2151
7	16.4	1078	1266.611	-188.6108
8	16.3	1094	1264.761	-170.7606
9	10.5	1237	1157.454	79.54626
10	26.3	833	1449.773	-616.7726
11	25.9	567	1442.372	-875.3721
12	26.8	933	1459.023	-526.0231
13	14.8	1339	1237.009	101.9911
14	22.3	937	1375.768	-438.7678
15	56.3	2011	2004.808	6.191895

Algebraische Eigenschaften der OLS-Statistiken

1. Die Summe der OLS-Residuen ergibt null:

$$\sum_{i=1}^n \hat{u}_i = 0$$

2. Die Stichprobenkovarianz des Regressors und der OLS-Residuen ist null:

$$\sum_{i=1}^n x_i \hat{u}_i = 0$$

3. Die Durchschnitte des Regressors und des Regressanden liegen auf der Regressionsgeraden:

$$\bar{y} = \hat{\beta}_0 + \hat{\beta}_1 \bar{x}$$

Diese drei Gleichungen folgen direkt aus der Definition der OLS-Schätzer $\hat{\beta}_0$ und $\hat{\beta}_1$.

OLS-Zerlegung

Wir können die Beobachtungen y_i in einen erklärten und einen unerklärten Teil zerlegen:

$$y_i = \hat{y}_i + \hat{u}_i$$

Aus den Eigenschaften

$$\sum_{i=1}^n \hat{u}_i = 0 \text{ und } \sum_{i=1}^n x_i \hat{u}_i = 0$$

folgt, dass der erklärte Teil $\hat{y}_i$ und der unerklärte Teil $\hat{u}_i$ unkorreliert sind:

$$\sum_{i=1}^n (\hat{y}_i - \bar{y})(\hat{u}_i - \bar{\hat{u}}) = \sum_{i=1}^n \hat{y}_i \hat{u}_i = \hat{\beta}_0 \sum_{i=1}^n \hat{u}_i + \hat{\beta}_1 \sum_{i=1}^n x_i \hat{u}_i = 0$$

OLS-Zerlegung

Ebenso können wir die Stichprobenvarianz der Beobachtungen in einen erklärten und in einen unerklärten Teil zerlegen.

Total Sum of Squares

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

Explained Sum of Squares

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

Residual Sum of Squares

$$SSR = \sum_{i=1}^n (\hat{u}_i - \bar{\hat{u}})^2$$

Es gilt:

$$SST = SSE + SSR$$

Das Bestimmtheitsmaß / Coefficient of Determination

Die Güte der Anpassung einer OLS-Regression ist der erklärte Anteil der Stichprobenvarianz der erklärten Variable.

Je größer SSE im Vergleich zu SSR , desto besser die Güte der Anpassung.

Mit $SST = SSE + SSR$ definieren wir das **Bestimmtheitsmaß** R^2 :

$$R^2 = SSE/SST = 1 - SSR/SST$$

Das Bestimmtheitsmaß nimmt Werte zwischen null und eins an und kann wie folgt interpretiert werden:

Der Regressor erklärt $100 \cdot R^2\%$ der Streuung des Regressanden.

Das Bestimmtheitsmaß der Beispiele

CEO-Gehälter und Eigenkapitalrentabilität

$$\widehat{salary} = 963.191 + 18.501roe$$

$$n = 209, R^2 = 0.0132$$

Die Regression erklärt nur 1.3% der Streuung der Gehälter!

Wahlergebnisse und Wahlkampfausgaben

$$\widehat{voteA} = 26.81 + 0.464shareA$$

$$n = 173, R^2 = 0.856$$

Der Anteil der Wahlkampfausgaben erklärt 86% der Streuung der Wahlergebnisse!

Maßeinheiten und die funktionale Form

Lineare Transformationen (Veränderung der Maßeinheiten)

Multiplikation der Variablen des linearen Regressionsmodells mit Konstanten impliziert eine entsprechende Transformation der Schätzer:

Seien $\tilde{x}_i = c \cdot x_i$ und $\tilde{y}_i = d \cdot y_i$ für $c, d \in \mathbb{R}$, $c, d \neq 0$, $i = 1, \dots, n$.

Dann gilt:

$$\tilde{y}_i = \tilde{\beta}_0 + \tilde{\beta}_1 \tilde{x}_i$$

$\Leftrightarrow$

$$d y_i = \tilde{\beta}_0 + \tilde{\beta}_1 c x_i$$

$\Leftrightarrow$

$$y_i = \frac{1}{d} \tilde{\beta}_0 + \frac{c}{d} \tilde{\beta}_1 x_i$$

dass

$$\tilde{\beta}_0 = d \hat{\beta}_0 \text{ und } \tilde{\beta}_1 = \frac{d}{c} \hat{\beta}_1$$

Das Bestimmtheitsmaß R^2 ist invariant!

Die Semi-logarithmische Form

Die Semi-logarithmische Form misst die prozentuale Veränderung des Regressanden bei absoluter Veränderung des Regressors.

Das ökonometrische Modell

$$\ln(wage) = \beta_0 + \beta_1 educ + u$$

misst (falls $\Delta u = 0$):

$$\frac{\Delta wage}{wage} \approx \beta_1 \Delta educ$$

Eine Erhöhung der Bildung um ein Jahr geht einher mit einem Anstieg des Lohnes um $100 \cdot \beta_1$ %.

Log-Stundenlohn und Bildung

Population Regression Function

$$\ln(wage) = \beta_0 + \beta_1 educ + u$$

$\ln(wage)$: natürlicher Logarithmus des Stundenlohns in \$

$educ$: Bildungsjahre

Regressionsgerade:

$$\widehat{\ln(wage)} = 0.584 + 0.083educ$$

$$n = 526, R^2 = 0.186$$

Im Vergleich:

Wird $wage$ anstelle von $\ln(wage)$ erklärt, beträgt $R^2 = 0,165$.

Die Log-Logarithmische Form

Die Log-logarithmische Form misst die prozentuale Veränderung des Regressanden bei prozentualer Veränderung des Regressanden.

Im ökonometrische Modell

$$\ln(\text{salary}) = \beta_0 + \beta_1 \ln(\text{sales}) + u$$

ist β_1 die Elastizität von *salary* in Bezug auf *sales*:

Wenn sich *sales* um 1% ändert, ändert sich *salary* um $\beta_1\%$.

Die OLS-Regression ergibt

$$\widehat{\ln(\text{salary})} = 4.822 + 0.257 \ln(\text{sales})$$

$$n = 209, R^2 = 0.211$$

Diese Form wird auch **Modell mit konstanten Elastizitäten** genannt.

Nichtlineare Transformationen

Zusammenfassung:

Modell	Regressand	Regressor	Interpretation β_1
Level-Level	y	x	$\Delta y = \beta_1 \Delta x$
Level-Log	y	$\ln(x)$	$\Delta y \approx \beta_1 \frac{\Delta x}{x}$
Log-Level	$\ln(y)$	x	$\frac{\Delta y}{y} \approx \beta_1 \Delta x$
Log-Log	$\ln(y)$	$\ln(x)$	$\frac{\Delta y}{y} \approx \beta_1 \frac{\Delta x}{x}$

Erwartungswerte und Varianzen der OLS-Schätzer

Die Schätzer $\hat{\beta}_0$ und $\hat{\beta}_1$ des einfachen linearen Regressionsmodells sind Funktionen der Daten der Stichprobe.

Wenn wir die Stichprobe als Zufallsgröße betrachten, so sind auch $\hat{\beta}_0$ und $\hat{\beta}_1$ Zufallsvariablen und haben statistische Eigenschaften.

Um diese Eigenschaften herzuleiten, benötigen wir einige Annahmen.

Die Abkürzung *SLR* steht hierbei für „simple linear regression“.

Linearität in den Parametern

Im ökonometrischen Modell steht die abhängige Variable y in Relation zu der unabhängigen Variable x und dem Störterm u :

$$y = \beta_0 + \beta_1 x + u ,$$

wobei β_0 und β_1 unbekannte aber feste Parameter sind.

Bemerkungen:

- ▶ Erlaubt, dass wir die Variablen nichtlinear transformieren.
- ▶ Fordert additive Separabilität.

Annahme SLR 3

Variation der erklärenden Variable

Die Stichprobenwerte $x_1, x_2, \dots, x_n$ sind nicht alle identisch.

Bemerkung:

Mit SLR 3 folgt: $\overline{xx} - \bar{x}\bar{x} > 0$

Bedingter Nullerwartungswert

Bei gegebenen Werten x beträgt der bedingte Erwartungswert des Störterms $u_i = y_i - \beta_0 - \beta_1 x_i$, $i = 1, \dots, n$, null:

$$E[u_i|x] = 0, \quad i = 1, \dots, n$$

Bemerkung:

Mit SLR 1 folgt: $E[y_i|x] = \beta_0 + \beta_1 x_i$, $i = 1, \dots, n$.

Erwartungstreue

Die beiden Schätzer

$$\begin{aligned}\hat{\beta}_0 &= \bar{y} - \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}}\bar{x} \\ \hat{\beta}_1 &= \frac{\overline{xy} - \bar{x}\bar{y}}{\overline{xx} - \bar{x}\bar{x}}\end{aligned}$$

sind erwartungstreu:

Theorem 2.1

Unter den Annahmen SLR 1, SLR 3 und SLR 4 gilt

$$E[\hat{\beta}_0] = \beta_0 \text{ und } E[\hat{\beta}_1] = \beta_1$$

für beliebige Werte β_0 und β_1 .

Beweis Theorem 2.1 für $E[\hat{\beta}_1] = \beta_1$

$$\begin{aligned} E[\hat{\beta}_1|x] &\stackrel{SLR\ 3}{=} E\left[\frac{\sum_{i=1}^n (x_i - \bar{x})y_i}{\sum_{i=1}^n (x_i - \bar{x})^2} \middle| x\right] \\ &= \frac{\sum_{i=1}^n (x_i - \bar{x})E[y_i|x]}{\sum_{i=1}^n (x_i - \bar{x})^2} \\ &\stackrel{SLR\ 1,4}{=} \frac{\sum_{i=1}^n (x_i - \bar{x})(\beta_0 + \beta_1 x_i)}{\sum_{i=1}^n (x_i - \bar{x})^2} \\ &= \beta_0 \frac{\sum_{i=1}^n (x_i - \bar{x})}{\sum_{i=1}^n (x_i - \bar{x})^2} + \beta_1 \frac{\sum_{i=1}^n (x_i - \bar{x})x_i}{\sum_{i=1}^n (x_i - \bar{x})^2} = \beta_1 \end{aligned}$$

Mit der Teleskopeigenschaft für Erwartungswerte gilt

$$E[\hat{\beta}_1] = E[E[\hat{\beta}_1|x]] = \beta_1$$

Beweis Theorem 2.1 für $E[\hat{\beta}_0] = \beta_0$

$$\begin{aligned} E[\hat{\beta}_0|x] &= E[\bar{y} - \hat{\beta}_1\bar{x} | x] \\ &= E[\bar{y}|x] - E[\hat{\beta}_1|x]\bar{x} \\ &= \frac{1}{n} \sum_{i=1}^n E[y_i|x] - \beta_1\bar{x} \\ &\stackrel{SLR\ 1,4}{=} \beta_0 + \beta_1\bar{x} - \beta_1\bar{x} = \beta_0 \end{aligned}$$

Mit der Teleskopeigenschaft für Erwartungswerte gilt

$$E[\hat{\beta}_0] = E[E[\hat{\beta}_0|x]] = \beta_0$$

Annahmen

Annahme SLR 5

Homoskedastizität

Für beliebige gegebene Regressordaten x ist die bedingte Varianz des Störterms u_i , $i = 1, \dots, n$ identisch:

$$\text{Var}(u_i|x) = \sigma^2 > 0, \quad i = 1, \dots, n$$

Bemerkungen:

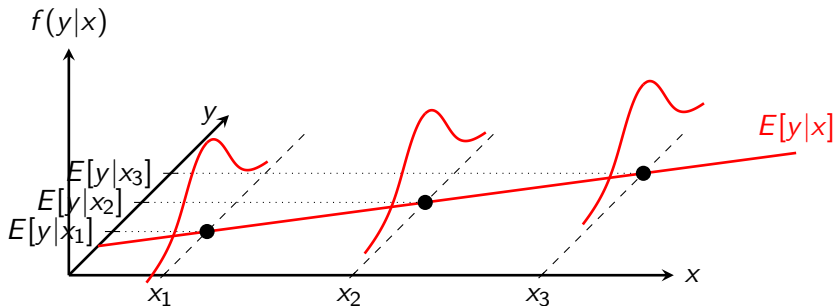
Da $\text{Var}(u_i|x)$ nicht von x abhängt, gilt außerdem $\text{Var}(u_i) = \sigma^2$.

Ist SLR 5 verletzt, heißt der Störterm **heteroskedastisch**.

Annahme SLR 5b

Wir fordern zusätzlich: $\text{Cov}(u_i, u_j|x) = 0$ für $i, j = 1, \dots, n$, $i \neq j$

Bedingte Verteilung von y unter Homoskedastizität



Aus $\text{Var}(u_i|x) = \sigma^2$ für $i = 1, \dots, n$ und SLR 1 folgt

$$\text{Var}(y_i|x) = \text{Var}(\beta_0 + \beta_1 x_i + u_i|x_i) = \text{Var}(u_i|x_i) = \sigma^2 \text{ für } i = 1, \dots, n$$

Theorem 2.2

Stichprobenvarianzen der OLS-Schätzer

Unter den Annahmen SLR 1, SLR 3, SLR 4 und SLR 5+b gilt:

$$\text{Var}(\hat{\beta}_0|x) = \sigma^2 \frac{1}{n} \frac{\overline{xx}}{\overline{xx} - \bar{x}\bar{x}}$$

$$\text{Var}(\hat{\beta}_1|x) = \sigma^2 \frac{1}{n} \frac{1}{\overline{xx} - \bar{x}\bar{x}}$$

Beweis Theorem 2.2 für $Var(\hat{\beta}_1|x)$

$$\begin{aligned} Var(\hat{\beta}_1|x) &\stackrel{SLR\ 3}{=} Var\left(\frac{\sum_{i=1}^n (x_i - \bar{x})y_i}{\sum_{i=1}^n (x_i - \bar{x})^2}\right) \\ &\stackrel{SLR\ 1}{=} Var\left(\frac{\sum_{i=1}^n (x_i - \bar{x})(\beta_0 + \beta_1 x_i)}{\sum_{i=1}^n (x_i - \bar{x})^2} + \frac{\sum_{i=1}^n (x_i - \bar{x})u_i}{\sum_{i=1}^n (x_i - \bar{x})^2}\right) \\ &\stackrel{SLR\ 2,5b}{=} \frac{\sum_{i=1}^n (x_i - \bar{x})^2 Var(u_i|x)}{(\sum_{i=1}^n (x_i - \bar{x})^2)^2} \\ &\stackrel{SLR\ 5}{=} \frac{\sum_{i=1}^n (x_i - \bar{x})^2 \sigma^2}{(\sum_{i=1}^n (x_i - \bar{x})^2)^2} \\ &= \frac{\sigma^2}{\sum_{i=1}^n (x_i - \bar{x})^2} = \frac{1}{n} \sigma^2 \frac{1}{\overline{xx} - \bar{x}\bar{x}} \end{aligned}$$

Je stärker der Regressor streut und je größer die Stichprobe ist, desto genauer ist $\hat{\beta}_1$.

Beweis Theorem 2.2 für $\text{Var}(\hat{\beta}_0|x)$ Schritt 1

$$\begin{aligned}\text{Var}(\bar{y}|x) &\stackrel{SLR\ 1}{=} \text{Var}\left(\frac{1}{n}\sum_{i=1}^n \beta_0 + \beta_1 x_i + u_i \middle| x\right) \\&= \text{Var}\left(\frac{1}{n}\sum_{i=1}^n u_i \middle| x\right) \\&\stackrel{SLR\ 5b}{=} \frac{1}{n^2} \sum_{i=1}^n \text{Var}(u_i|x) \\&\stackrel{SLR\ 5}{=} \frac{1}{n} \sigma^2\end{aligned}$$

Beweis Theorem 2.2 für $\text{Var}(\hat{\beta}_0|x)$ Schritt 2

$$\begin{aligned}\text{Cov}(\bar{y}, \hat{\beta}_1|x) &\stackrel{SLR\ 1}{=} \text{Cov}\left(\frac{1}{n} \sum_{i=1}^n y_i, \frac{\sum_{i=1}^n (x_i - \bar{x}) y_i}{\sum_{i=1}^n (x_i - \bar{x})^2} \middle| x\right) \\&= \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^n \frac{x_j - \bar{x}}{\sum_{k=1}^n (x_k - \bar{x})^2} \text{Cov}(y_i, y_j|x) \\&\stackrel{SLR\ 5b}{=} \frac{1}{n} \sum_{i=1}^n \frac{x_i - \bar{x}}{\sum_{k=1}^n (x_k - \bar{x})^2} \text{Cov}(y_i, y_i|x) \\&\stackrel{SLR\ 5}{=} \sigma^2 \frac{1}{n} \frac{\sum_{i=1}^n x_i - \bar{x}}{\sum_{i=1}^n (x_i - \bar{x})^2} = 0\end{aligned}$$

Beweis Theorem 2.2 für $\text{Var}(\hat{\beta}_0|x)$ Schritt 3

$$\begin{aligned}\text{Var}(\hat{\beta}_0|x) &= \text{Var}(\bar{y} - \hat{\beta}_1 \bar{x}|x) \\&= \text{Var}(\bar{y}|x) - 2\text{Cov}(\bar{y}, \hat{\beta}_1|x)\bar{x} + \text{Var}(\hat{\beta}_1|x)\bar{x}\bar{x} \\&= \frac{1}{n}\sigma^2 - 2 \cdot 0 + \frac{1}{n}\sigma^2 \frac{1}{\overline{xx} - \bar{x}\bar{x}} \bar{x}\bar{x} \\&= \frac{1}{n}\sigma^2 \frac{\overline{xx}}{\overline{xx} - \bar{x}\bar{x}}\end{aligned}$$

Schätzung der Varianz des Störterms

Bei einer gegebenen Stichprobe sind die Ausdrücke

$$\overline{xx} = \frac{1}{n} \sum_{i=1}^n x_i^2$$

und

$$\bar{x}\bar{x} = \left(\frac{1}{n} \sum_{i=1}^n x_i \right)^2$$

bekannt.

Bei σ^2 handelt es sich aber um einen unbekannten Parameter, der zu schätzen ist.

Theorem 2.3

Unverzerrter Schätzer für σ^2

Sei der Schätzer für σ^2 definiert durch:

$$\hat{\sigma}^2 = \frac{1}{n-2} \sum_{i=1}^n \hat{u}_i^2$$

Unter den Annahmen SLR 1, SLR 3, SLR 4 und SLR 5+b gilt:

$$E[\hat{\sigma}^2|x] = \sigma^2$$

Begründung für Theorem 2.3

Der Beweis ist etwas aufwendiger und wird später nachgeholt.

Warum gilt nicht $\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n \hat{u}_i^2$?

Durch die beiden Gleichungen

$$\sum_{i=1}^n \hat{u}_i = 0$$

$$\sum_{i=1}^n x_i \hat{u}_i = 0$$

werden zwei Freiheitsgrade „verbraucht“.

Daher teilen wir die Summe der quadrierten Residuen durch $n - 2$.

Schätzer für die Varianzen der Schätzer

Mit dem Schätzer

$$\hat{\sigma}^2 = \frac{1}{n-2} \sum_{i=1}^n \hat{u}_i^2$$

ergibt sich

$$\widehat{Var}(\hat{\beta}_0|x) = \hat{\sigma}^2 \frac{1}{n} \frac{\overline{xx}}{\overline{xx} - \bar{x}\bar{x}}$$

$$\widehat{Var}(\hat{\beta}_1|x) = \hat{\sigma}^2 \frac{1}{n} \frac{1}{\overline{xx} - \bar{x}\bar{x}}$$

Mit den Standardfehlern der Schätzer der Regressionskoeffizienten

$$se(\hat{\beta}_0) = \sqrt{\widehat{Var}(\hat{\beta}_0|x)} \text{ und } se(\hat{\beta}_1) = \sqrt{\widehat{Var}(\hat{\beta}_1|x)}$$

können wir später Intervallschätzungen für die Regressionskoeffizienten vornehmen.