

Faculty of Health Sciences

Basal Statistik

Multipel regressionsanalyse i SPSS.

Lene Theil Skovgaard

5. oktober 2020



Multipel lineær regression

- ▶ Regression med to kvantitative kovariater:
 - ▶ Eksempel om ultralyd
 - ▶ Eksempel om fedme hos børn
- ▶ Selektion, modelvalg
- ▶ Modelkontrol
- ▶ Diagnostics

E-mail: ltsk@sund.ku.dk

*: Siden er lidt teknisk



Multipel regression

Ét outcome, mange forklarende variable

Vi har allerede set eksempler på dette:

- ▶ Sammenligning af vægt for mænd og kvinder, korrigeret for højde
- ▶ Vurdering af skuldersmerter efter to behandlinger, korrigeret for baseline smerter
- ▶ Effekt af P-piller på hormon-niveauet, med korrektion for alder

I begge disse situationer havde vi netop to kovariater:

- ▶ Én kvantitativ (højde, baseline smerter, alder)
- ▶ Én kategorisk kovariat (køn, behandling, P-piller)



Problemstillinger ved multipel regression

- ▶ **Prediktion:**
Konstruktion af normalområde til diagnostisk brug
(som i eksemplet om ultralyds scanning, s. 5)
- ▶ **Udredning af årsagssammenhænge,**
til brug for interventioner
(mere som eksemplet s. 41ff)
- ▶ **Videnskabelig indsigt,**
f.eks. eksemplet om P-pillers indvirkning på hormoner,
fra sidste forelæsning



Eksempel med to kvantitative kovariater

107 kvinder er blevet ultralyds scannet få dage inden fødslen, og der ønskes etableret en prediktion af fostervægt/fødselsvægt ud fra BPD (hoved-diameter) og AD (maveomfang).

Data er:

OBS	VAEGT	BPD	AD
1	2350	88	92
2	2450	91	98
3	3300	94	110
.	.	.	.
.	.	.	.
.	.	.	.
105	3550	92	116
106	1173	72	73
107	2900	92	104

Kilde: Secher, N.J., Århus Kommunehospital



Multipel regression

DATA: n personer, dvs. n sæt af sammenhørende observationer:

person	x_1		x_p	y
1	x_{11}		x_{1p}	y_1
2	x_{21}		x_{2p}	y_2
3	x_{31}		x_{3p}	y_3
.	.	.	.	.
n	x_{n1}		x_{np}	y_n

Den **lineære regressionsmodel** med p forklarende variable skrives:

$$y = \beta_0 + \beta_1 x_1 + \cdots + \beta_p x_p + \varepsilon$$

Parametre:

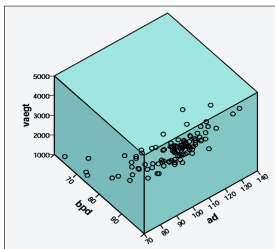
β_0 afskæring, intercept
 $\beta_1, \cdots, \beta_p$ regressionskoefficienter

Grafik

er lidt vanskelig, her et par forsøg på 3-dimensionalt plots:

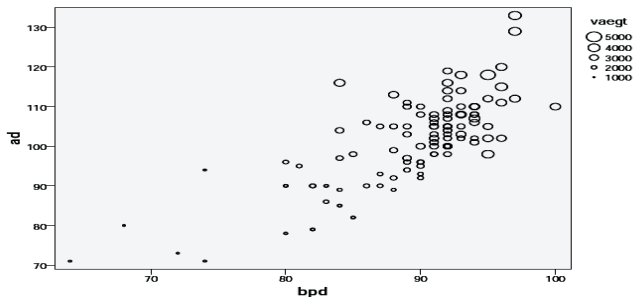
Benyt Graphs/Graphboard Template Choser, marker de 3 involverede variable (vaegt, bpd og ad) og klik på Scatterplot.

For at bytte om på akserne, klikkes på Detailed, hvor man så kan sætte vaegt på Y-aksen, bpd på X-aksen og ad på Z-aksen.



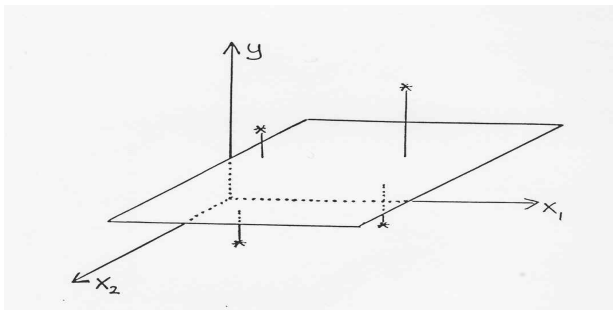
Grafik, fortsat

En anden mulighed er at vælge Bubble Plot under Graphs/Graphboard Template Choser, i stedet for Scatterplot (se s. 7)



*Middelværdistruktur

$$\beta_0 + \beta_1 x_1 + \cdots + \beta_p x_p$$



Mindste kvadraters metode: Minimer størrelsen

$$S(\beta_0, \beta_1, \dots, \beta_p) = \sum (y_i - \beta_0 - \beta_1 x_{i1} - \cdots - \beta_p x_{ip})^2$$

Multipel regression uden transformation

Den *minimale* opsætning:

- ▶ Benyt Analyze/Regression/Linear
- ▶ Sæt vægt over i Dependent
- ▶ Sæt både bpd og ad over i Independent(s)

Som minimum bør man dog også klikke Statistics og afkrydse Confidence Intervals.

Ekstra muligheder kan tilføjes, som det vil fremgå af nogle af de efterfølgende sider (s. 18 og s. 91-92)



Output fra multipel regression

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,898 ^a	,806	,803	306,573

a. Predictors: (Constant), ad, bpd

b. Dependent Variable: vaegt

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	40736853,73	2	20368426,87	216,715	,000 ^b
	Residual	9774647,335	104	93986,994		
	Total	50511501,07	106			

a. Dependent Variable: vaegt

b. Predictors: (Constant), ad, bpd



Output fra multipel regression, fortsat

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	-4628,118	455,990		-10,150	,000
	bpd	37,133	7,615	,322	4,876	,000
	ad	39,763	4,164	,630	9,549	,000

Coefficients^a

Model		95,0% Confidence Interval for B	
		Lower Bound	Upper Bound
1	(Constant)	-5532,363	-3723,873
	bpd	22,032	52,234
	ad	31,506	48,020

a. Dependent Variable: vægt

med fortolkning på næste side



Fortolkning af output s. 11-12

- ▶ Stærk signifikant effekt af begge kovariater ($P = 0.000$)
- ▶ Hvis vi sammenligner to fostre med samme AD, men hvor foster A har en BPD, der er 1mm større end foster B, så vil vi forvente, at A vejer 37.1g mere end B
- ▶ Hvis vi sammenligner to fostre med samme BPD, men hvor foster A har en AD, der er 1mm større end foster B, så vil vi forvente, at A vejer 39.8g mere end B
- ▶ Residualspredningen
(Std Error of the Estimate)= 306.573 , så usikkerheden på prediktionen af fødselsvægt er $\pm 2 \times 306.57\text{g}$, altså *ganske betragtelig* (næsten til lagkage)



Modelkontrol

Hvilke antagelser skal vi checke?

- ▶ Uafhængighed:

Tænk: Er der flere observationer på hvert individ?
Søskende el.lign?

- ▶ Linearitet i begge kovariater
- ▶ Varianshomogenitet (residualer har samme spredning)
- ▶ Normalfordelte residualer

Obs: Intet krav om normalfordeling på kovariaterne!!



Grafisk modelkontrol: residualer

Ud over de *sædvanlige* residualer = observeret - fittet værdi:

$\hat{\varepsilon}_i = y_i - \hat{y}_i$, er der 4 andre typer at vælge imellem:

1. Normerede/Standardiserede: ZRESID
2. Studentized: SRESID
3. Deleted / Leave-one-out: DRESID
4. Studentized Deleted: SDRESID

Alle disse kan fås i nyt datasæt ved at anvende Save-knappen, og krydse det af, man gerne vil bruge (se også s. 85)

Hvilke residualer skal man benytte?

Fordele og ulemper

- ▶ Rart med residualer, der bevarer enhederne (de sædvanlige og type 3)
- ▶ Lettest at finde outliers ud fra de residualer, hvor observationerne udelades en ad gangen (type 3 og 4)
- ▶ Bedst at normere, når observationen selv er med og man ikke kan tegne rådata



Residualplots til modelkontrol

Residualer (af passende type) plottes mod

1. den forklarende variabel x_i
 - for at checke **linearitet**
(se efter *krumninger*, *buer*)
2. de fittede værdier $\hat{y}_i$
 - for at checke **varianshomogenitet**
(se efter *trompeter*)
3. fraktildiagram eller histogram
 - for at checke **normalfordelingsantagelsen**
(se efter *afvigelse fra en ret linie*)

Disse plots ses på s. 20 og 19, og kommenteres s. 22



Residualplots, fortsat

Et **bedre check af varianshomogeniteten** er dog et plot af kvadratroden af den numeriske værdi af normerede residualer, mod de predikterede værdier, eller mod de forklarende variable

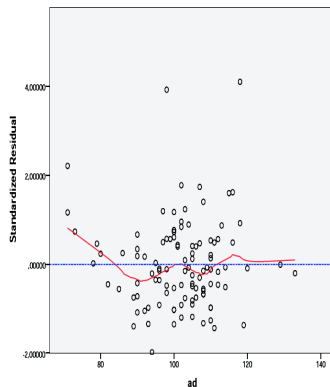
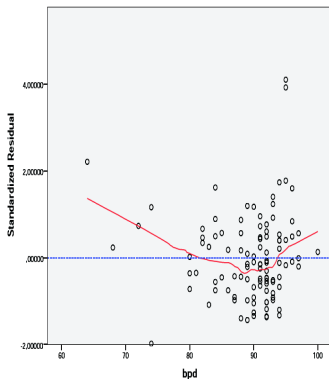
De normerede residualer hedder `SRE_1`, og for at definere kvadratroden af de numeriske residualer går vi ind i Transform/Compute, skriver `sqrtres` i Target Variable og `sqrt(abs(SRE_1))` i definitionsboksen.

Disse *specialplots* ses på side 21

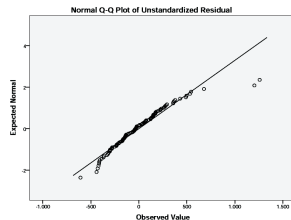
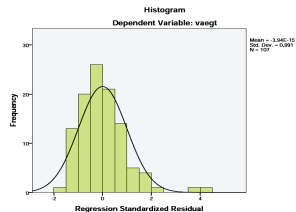
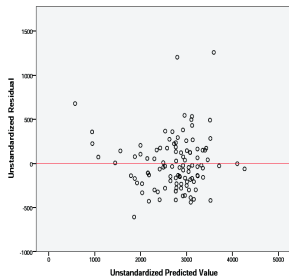


Residualplots (type 1 fra s. 17)

`mod kovariater`, med indlagte udglattede kurver (Loess-kurver):

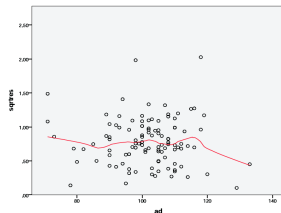
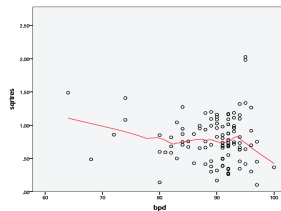
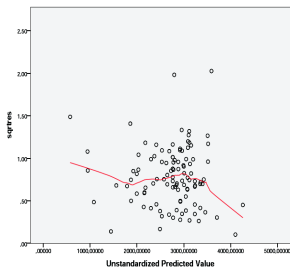


Residualplots (type 2 og 3 fra s. 17)



Specialplots til vurdering af konstant spredning

som beskrevet s. 18 (eller s. 87)



Vurdering af modellen

- ▶ Normalfordelingen halter en anelse, med nogle enkelte ret store positive afvigelse, hvilket kunne tale for at logaritmetransformere vægten.
- ▶ Måske lidt trompetfacon i plot af residualer mod predikterede værdier, men husk på, at observationerne ikke er ligeligt fordelt over x-aksen.

Figuren s. 21 viser ingen oplagte tendenser.

- ▶ Linearitet er ikke helt god, men det skyldes hovedsageligt de tidligst fødte børn.
- ▶ Teoretiske argumenter fra den faglige ekspertise foreslår en samtidig logaritmetransformation af såvel outcome som kovariater



Analyse af logaritmerede data

Vi benytter forskellige logaritmer til outcome og kovariater, for definition, se s. 88:

$\text{logvaegt} = \text{Ln}(\text{vaegt})$

$\text{logbpd} = \text{Ln}(\text{bpd}/90) / \text{Ln}(1.1)$ /* 1 enhed er 10% */

$\text{logad} = \text{Ln}(\text{ad}/100) / \text{Ln}(1.1)$ /* mere om dette s. 25 og 29 */

Analysen med logaritmerede værdier, svarende til s. 10 ses nedenfor og på næste side.

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,926 ^a	,858	,856	,10679

a. Predictors: (Constant), logad, logbpd

b. Dependent Variable: logvaegt



Analyse af logaritmerede data

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	7,876	,011		726,497	,000
	logbpd	,148	,022	,396	6,764	,000
	logad	,140	,014	,585	9,998	,000

Coefficients^a

Model		95,0% Confidence Interval for B	
		Lower Bound	Upper Bound
1	(Constant)	7,855	7,898
	logbpd	,105	,191
	logad	,112	,168

a. Dependent Variable: logvaegt

Se fortolkning s. 25-26



Fortolkning af resultater

- ▶ Stærkt signifikant effekt af begge kovariater ($P < 0.0001$)
- ▶ Hvis vi sammenligner to fostre med samme AD, men hvor foster A har en BPD, der er 10% større end foster B, så vil vi forvente, at A vejer $\exp(0.148) = 1.16$ gange så meget som B, dvs. 16% mere.
- ▶ Hvis vi sammenligner to fostre med samme BPD, men hvor foster A har en AD, der er 10% større end foster B, så vil vi forvente, at A vejer $\exp(0.140) = 1.15$ gange så meget som B, dvs. 15% mere.



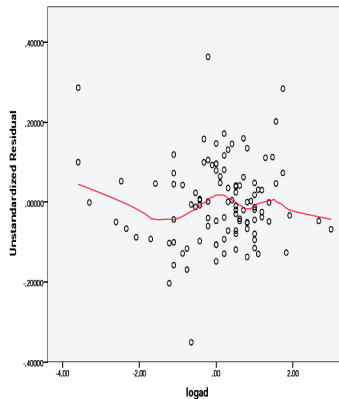
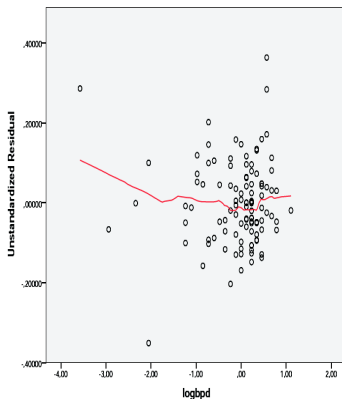
Fortolkning af resultater, fortsat

- ▶ Root MSE=0.10679, så usikkerheden på prediktionen af $\log(\text{fødselsvægt})$ er $\pm 2 \times 0.10679 = 0.21358$, svarende til faktorerne $(\exp(-0.21358), \exp(0.21358)) = (0.808, 1.238)$, altså en variation gående fra -19.2% op til +23.8%, igen en *ganske betragtelig* biologisk variation.
Bemærk asymmetrien i denne kvantificering!
- ▶ Bemærk, at vi i denne model har konstant *relativ* usikkerhed. Variationskoefficienten kan faktisk nu aflæses som Std. Error of the Estimate, fordi vi arbejder med *naturlige* logaritmer.
Vi finder tallet 0.10679 (se s. 23), svarende til CV=10.7%

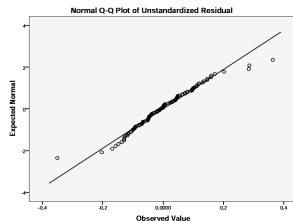
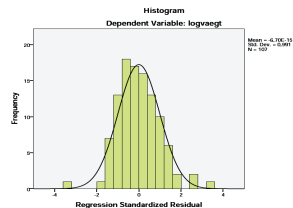
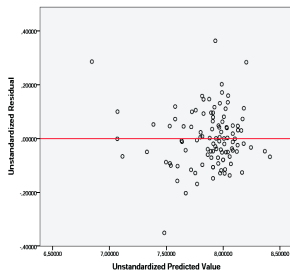


Residualplots for transformerede data

`mod kovariater`, med indlagte udglattede kurver (Loess-kurver):

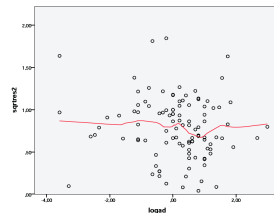
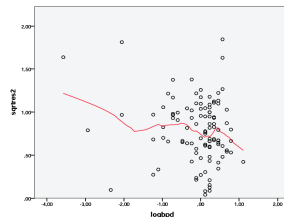
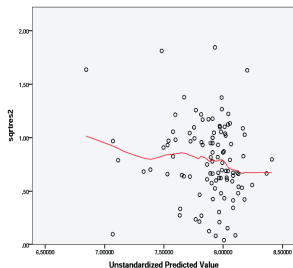


Residualplots for transformede data, fortsat



Specialplots til vurdering af konstant spredning

som beskrevet s. 18



Sammenligning af modeller

nemlig de to *marginal*e modeller (hver med sin kovariat) og *den multiple* regressionsmodel:

Estimaterne for disse modeller

(med tilhørende standard errors (*se*) i parentes):

Intercept	β_1 (logbpd)	β_2 (logad)	s	R^2
7.91	0.318 (0.019)	-	0.149	0.72
7.85	-	0.213 (0.011)	0.128	0.80
7.88	0.148 (0.022)	0.140 (0.014)	0.107	0.86

Bemærk fortolkningen af interceptet: På grund af konstruktionen af logbpd og logad på s. 23, svarer interceptet til den forventede fødselsvægt for et barn med bpd=90 og ad=100.



Fortolkning af koefficienter

f.eks. effekten af bpd:

- ▶ **Marginal model:**

Ændringen i logvægt, når kovariaten $\log bpd$ ændres 1 enhed, dvs. når bpd øges med 10%

- ▶ **Multipel regressionsmodel**

Ændringen i logvægt, når kovariaten $\log bpd$ ændres 1 enhed, men hvor alle andre kovariater (her kun ad) **holdes fast**

I sidstnævnte situation siger vi, at vi har **korregeret** for effekten af de andre kovariater i modellen.

Forskellen kan være markant, fordi kovariaterne typisk er **relaterede**:

- ▶ Når en af dem ændres, ændres de andre også



Goodness-of-fit mål: R-i-anden

$$R^2 = \frac{\text{Sum Sq}(\text{Model})}{\text{Sum Sq}(\text{Total})}$$

“Hvor stor en del af variationen kan forklares af modellen?”
Her finder vi 0.858, dvs. 85.8% (se s. 23)

Dette mål har

- ▶ **Fortolkningsproblemer** når kovariaterne er styret (ganske som for korrelationskoefficienten)
- ▶ R^2 stiger med antallet af kovariater – selv hvis disse er uden betydning!

Derfor ser man ofte i stedet på **Adjusted R^2** :

$$R^2_{\text{adj}} = 1 - \frac{\text{Mean Sq}(\text{Residual})}{\text{Mean Sq}(\text{Total})} \quad (\text{her } 0.856, \text{ se s. } 23)$$



Fitted = predikterede værdier

kan udregnes som (se estimator s. 24)

$$\begin{aligned}\log(\text{vaegt}) &= 7.876 + 0.140 \times \log ad \\ &\quad + 0.148 \times \log bpd \Rightarrow \\ \text{vaegt} &= 0.0028 \times ad^{1.47} \times bpd^{1.55}\end{aligned}$$

I praksis vil man dog benytte Save-knappen, og afkrydse Predicted Values, (se evt. s. 85), således at de predikterede værdier bliver lagt som variabelen PRE_1 i datasættet.



* Prediktioner til brug for illustrationer

- ▶ Tilføj de fiktive personer, man gerne vil prediktere outcome for
- ▶ Angiv deres kovariatværdier, med tilhørende missing value for outcome
- ▶ Predikter outcome i nyt datasæt
- ▶ Tilbagetransformer til original skala
- ▶ Tegn

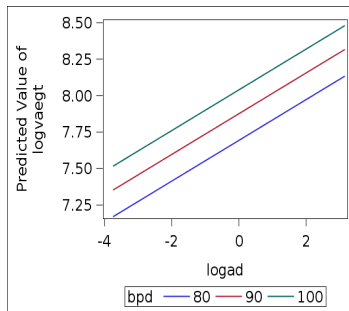
De fiktive personer har ingen indflydelse på estimation i modellen, da de mangler outcome-værdier

Se flere detaljer s. 89-90,
der i princippet kunne danne figurerne s. 35

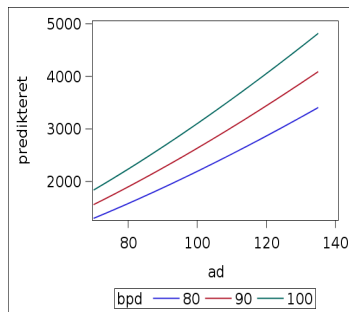


Predikterede værdier

På logaritmisk skala:



Tilbagetransformeret
til oprindelig skala:



Disse figurer er dog ikke lavet i SPSS....

Regression diagnostics

Understøttes konklusionerne af *hele* materialet?

Eller er der observationer med meget stor indflydelse på resultaterne?

Leverage = *potentiel* indflydelse (hat-matrix)

Hvis der kun er én kovariat er det simpelt:

$$h_{ii} = \frac{1}{n} + \frac{(x_i - \bar{x})^2}{S_{xx}}$$

Observationer med *ekstreme* x -værdier *kan* have stor indflydelse på resultaterne, men de har det *ikke nødvendigvis*!

- ▶ ikke hvis de ligger 'pænt' i forhold til regressionslinien, dvs. har et *lille residual*, så derfor ...-->



Regression diagnostics, fortsat

- ▶ Udelad den i'te person og find nye estimater for regressionskoefficienterne
- ▶ Udregn **Cook's afstand**, et samlet mål for ændringen i parameterestimaterne
- ▶ Spalt Cooks afstand ud i koordinater og angiv: Hvor mange *se*'er ændres f.eks. $\hat{\beta}_1$, når den i'te person udelades?

Se mere s. 91

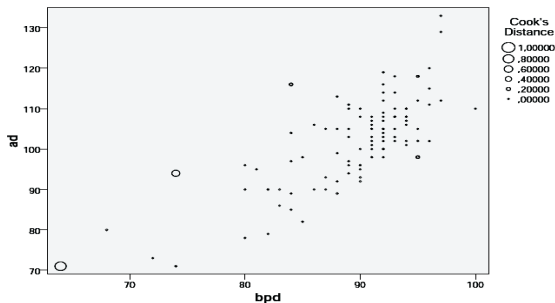
Hvad gør vi ved indflydelsesrige observationer?

- ▶ udelader dem?
- ▶ anfører et mål for deres indflydelse?



Cooks afstand som mål for indflydelse

Bubble plot, se s. 7-8



De mest indflydelsesrige observationer er dem med en usædvanlig kombination af kovariater.

Outliers

Observationer, der *ikke passer* ind i sammenhængen

- ▶ de er ikke nødvendigvis indflydelsesrige
- ▶ de har ikke nødvendigvis et stort residual

Hvad gør vi ved outliers?

- ▶ ser nærmere på dem,
de er tit ganske interessante

Hvornår kan vi *udelade* dem?

- ▶ hvis de ligger meget *yderligt*, dvs. har høj leverage
 - ▶ husk at afgrænse konklusionerne tilsvarende!
- ▶ hvis man kan finde årsagen
 - ▶ og da skal **alle sådanne** udelades!



Interaktion mellem kvantitative kovariater

er lidt svært.....

Hvis modellen ikke beskrives ved en plan, hvad gør man så?

Hvis man blot inkluderer et produktled mellem x_1 og x_2 , svarer det til, at antage, at

effekten af x_1 ændrer sig lineært med værdien af x_2 :

$$\begin{aligned}\mu &= \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \gamma x_1 x_2 \\ &= \beta_0 + (\beta_1 + \gamma x_2) x_1 + \beta_2 x_2, \quad \text{dvs.}\end{aligned}$$

Effekten af x_1 er $\beta_1 + \gamma x_2$

Alternativt må man opdele den ene variabel i grupper.....



Nyt eksempel: Fedme i skolealderen

Spørgsmål:

Hvordan hænger fedmegraden i skolealderen sammen med højde og vægt i 1-årsalderen?

For 197 børn har vi sammenhørende registreringer af

- ▶ Højde og vægt i 1-års alderen
- ▶ Fedmescore i skolealderen, dvs. en score, der udtrykker barnets vægt i forhold til alder og højde

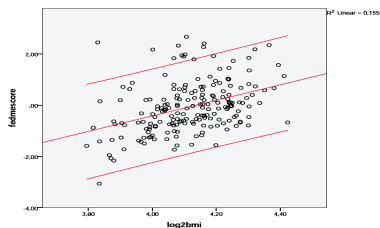
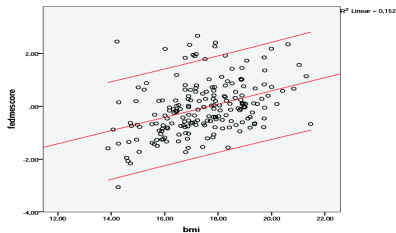
Det er en normeret størrelse,
der typisk vil ligge mellem -2 og 2

Da fedmescoren kan ses som et slags bmi-mål, er det rimeligt at forestille sig, at det hænger sammen med bmi i 1-års alderen.



Fedme vs. bmi

vs. body mass index, evt. logaritmeret:



Det ser nogenlunde pænt lineært ud, omend der er et par observationer, der ser “for store” ud.....

Måske skal vi bruge en anden kombination af højde og vægt end lige body mass index? (som ikke plejer at virke så godt for børn...)

Fedme i skolealderen, transformation

Hvis vi logaritmerer bmi (højre figur s. 42), får vi

$$\log(\text{bmi}) = \log(\text{vægt}) - 2 \times \log(\text{højde})$$

så en oplagt mulighed ville være at forsøge en multipel regression, med $\log(\text{vægt})$ og $\log(\text{højde})$ som kovariater, dvs. med multiplikative effekter af højde og vægt i 1-årsalderen.

Vi bruger igen 1.1-logaritmen:

$$\log_{\text{højde}} = \ln(\text{højde}/0.75)/\ln(1.1)$$

$$\log_{\text{vægt}} = \ln(\text{vægt}/10)/\ln(1.1)$$

og har samtidig normeret, således at interceptet kommer til at svare til et 1-årigt barn på 75 cm, der vejer 10 kg.



Fedme i skolealderen, analyse

Den *minimale* opsætning:

- Benyt Analyze/Regression/Linear
- Sæt fedmescore over i Dependent, og såvel logvaegt som loghoejde i Independent(s)
- Afkryds Confidence Intervals under Statistics

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,485 ^a	,235	,227	,87927

a. Predictors: (Constant), loghoejde, logvaegt



Fedme i skolealderen, resultater

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	,005	,065		,070	,944
	logvaegt	,457	,067	,672	6,803	,000
	loghoejde	-,460	,157	-,290	-2,935	,004

Coefficients^a

Model		95,0% Confidence Interval for B	
		Lower Bound	Upper Bound
1	(Constant)	-,125	,134
	logvaegt	,324	,589
	loghoejde	-,769	-,151

a. Dependent Variable: fedmescore

med fortolkning på næste side

Fortolkning af resultater

- ▶ Stærkt signifikant effekt af begge kovariater ($P = 0.000$ hhv. $P = 0.004$)
- ▶ Hvis vi sammenligner to børn **med samme højde**, men hvor barn A vejer 10% mere end barn B ved 1-års alderen, så vil vi forvente, at A's fedmescore er 0.457 større end B's.
- ▶ Hvis vi sammenligner to børn **med samme vægt**, men hvor barn A er 10% højere end barn B ved 1-års alderen, så vil vi forvente, at A's fedmescore er 0.460 **lavere** end B's.



Fortolkning af resultater, fortsat

Hvis det var bmi, der var den relevante kovariat, skulle vi have haft ca.

$$\text{koeff til vægt} = -2\text{koeff til højde}$$

men **det har vi ikke**

Vi ser på outputtet s. 45, at de to koefficienter snarere er lige store, blot med modsat fortegn, altså at $\text{koeff til vægt} = -\text{koeff til højde}$, således at **samme procentvise stigning i højde og vægt** giver uforandret fedmescore i skolealderen.



Sammenligning af modeller

nemlig de to marginale modeller (med hver sin kovariat)
og *den multiple regressionsmodel*:

Estimaterne for disse modeller

(med tilhørende standard errors (*s.e.*) i parentes):

Intercept	β_1 (loghøjde)	β_2 (logvægt)	s	R^2
-0.100	0.363 (0.111)	-	0.976	0.052
-0.051	-	0.305 (0.044)	0.896	0.201
0.005	-0.460 (0.157)	0.457 (0.067)	0.879	0.235

Bemærk:

- ▶ Koefficienterne ændres ved overgang til multipel regression
specielt den for højde, som ligefrem skifter fortegn!
- ▶ Usikkerhederne på estimerne bliver større, selvom residualspredningen s bliver mindre.



VIGTIGT

Hvordan kan man forstå to *så forskellige* estimater
for loghoejde som 0.363 og -0.460?

Fordi:

Den biologiske fortolkning af parameterestimerne
er helt forskellig

Det videnskabelige spørgsmål, der besvares,
er *ikke* det samme spørgsmål!



Fortolkning af estimater

Koefficienten β_1 til loghøjde:

- ▶ **Simpel/univariat regressionsmodel:**

Forventet øgning i fedmescore er 0.36 for 1 enheds øgning af kovariaten loghøjde, dvs. for en 10% forøgelse af 1-års højden.

- ▶ **Multipel regressionsmodel:**

Forventet *fald* i fedmescore er -0.46 for en 10% forøgelse af 1-års højden *for to individer, hvor alle andre kovariater* (her kun logvægt) *er identiske* ("holdes fast").

Det kaldes, at vi har **korrigeret** ("adjusted") for effekten af de andre kovariater.

Forskellen er her meget markant, fordi kovariaterne er stærkt relaterede: Når en af dem ændres, ændres den anden typisk også



Fortolkning af højde i marginal model

Når højden er den eneste kovariat, er den et udtryk for, hvor stort barnet er, så det videnskabelige spørgsmål, der belyses, er,

- Er børn, der er store i 1-årsalderen, i gennemsnit federe i skolealderen?

Den positive koefficient til højden betyder, at fedmegraden i skolealderen vokser signifikant med størrelsen i 1-års alderen.

Biologisk mekanisme: Overernæring i barnealderen bruges til at vokse, så barnet bliver stort af sin alder, men det kan øjensynligt øge risikoen for fedme senere. Heldigvis er det ikke en stærkt deterministisk sammenhæng (residualspredningen er relativt stor).



Fortolkning af højde i multipel regressionsmodel

Når vægten i 1-års alderen fastholdes, ved vi noget mere om det højeste af de to 1-årige børn:

Det højeste barn må også være det slankeste barn!

Så det nye videnskabelige spørgsmål, der belyses, er,

- Er børn, der er slanke i 1-årsalderen, i gennemsnit slankere i skolealderen?

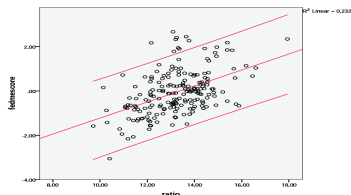
Den negative koefficient til højden betyder, at slanke småbørn generelt er slankere i skolealderen.



Fortolkning af højde i multipel regressionsmodel, II

Koefficienterne til de logaritme-transformerede højde og vægt er *lige store og med modsat fortegn*. Det betyder:

1. at vægt/højde (altså *ikke* $BMI = \text{vægt}/\text{højde}^2$) i 1-års alderen er prædiktivt for fedme i skolealderen
2. at to småbørn med samme vægt/højde-ratio har samme forventede fedmegrad i skolealderen



Eksempel: Lungefunktion og cystisk fibrose

Et *lille* studie, af 25 patienter, hvor outcome er pemax (et udtryk for lungefunktion) og hvor vi har hele 9 potentielle kovariater:

Table 12.11 Data for 25 patients with cystic fibrosis (O'Neill *et al.*, 1983)

Sub	Age	Sex	Height	Weight	BMP	FEV ₁	RV	FRC	TLC	PEmax
1	7	0	109	13.1	68	32	258	183	137	95
2	7	1	112	12.9	65	19	449	245	134	85
3	8	0	124	14.1	64	22	441	268	147	100
4	8	1	125	16.2	67	41	234	146	124	85
5	8	0	127	21.5	93	52	202	131	104	95
6	9	0	130	17.5	68	44	308	155	118	80
7	11	1	139	30.7	89	28	305	179	119	65
8	12	1	150	28.4	69	18	369	198	103	110
9	12	0	146	25.1	67	24	312	194	128	70
10	13	1	155	31.5	68	23	413	225	136	95
11	13	0	156	39.9	89	39	206	142	95	110
12	14	1	153	42.1	90	26	253	191	121	90
13	14	0	160	45.6	93	45	174	139	108	100
14	15	1	158	51.2	93	45	158	124	90	80
15	16	1	160	35.9	66	31	302	133	101	134
16	17	1	153	34.8	70	29	204	118	120	134
17	17	0	174	44.7	70	49	187	104	103	165
18	17	1	176	60.1	92	29	188	129	130	120
19	17	0	171	42.6	69	38	172	130	103	130
20	19	1	156	37.2	72	21	216	119	81	85
21	19	0	174	54.6	86	37	184	118	101	85
22	20	0	178	64.0	86	34	225	148	135	160
23	23	0	180	73.8	97	57	171	108	98	165
24	23	0	175	51.1	71	33	224	131	113	95
25	23	0	179	71.5	95	52	225	127	101	195

Ex. O'Neill et.al. (*Am Rev Respir Dis*, 1983)



Univariate analyser

også kaldet *marginale* analyser,
hvor man ser på *en enkelt kovariat ad gangen*:

Table 12.12 Results of separately regressing PEmax on each explanatory variable

Explanatory variable	Regression coefficient	Standard error	<i>t</i>	P
Age	4.055	1.088	3.73	0.0011
Sex	-19.045	13.176	-1.45	0.16
Height	0.932	0.260	3.59	0.0016
Weight	1.187	0.301	3.94	0.0006
BMP	0.639	0.565	1.13	0.27
FEV ₁	1.354	0.555	2.44	0.023
RV	-0.123	0.077	-1.59	0.12
FRC	-0.319	0.145	-2.20	0.038
TLC	-0.358	0.404	-0.89	0.38

Her er der 5 signifikante variable

Er det så disse variable, der skal med i modellen?



Valg af model

kommer helt og holdent an på, hvad spørgsmålet er!

De univariate analyser kan være stærkt misvisende pga korrelationer mellem de enkelte variable, f.eks. de 5 signifikante variable: Age, Height, Weight, FEV₁ og FRC:

Correlations

		age	height	weight	fev1
age	Pearson Correlation	1	,926**	,906**	,294
	Sig. (2-tailed)		,000	,000	,153
	N	25	25	25	25
height	Pearson Correlation	,926**	1	,921**	,317
	Sig. (2-tailed)	,000		,000	,123
	N	25	25	25	25
weight	Pearson Correlation	,906**	,921**	1	,449*
	Sig. (2-tailed)	,000	,000		,024
	N	25	25	25	25
fev1	Pearson Correlation	,294	,317	,449*	1
	Sig. (2-tailed)	,153	,123	,024	
	N	25	25	25	25

** . Correlation is significant at the 0.01 level (2-tailed).

* . Correlation is significant at the 0.05 level (2-tailed).

Bemærk især korrelationerne mellem alder, højde og vægt.



Videnskabelig variabelselektion

Endnu en gang:

Gennemtænk præcis hvilket videnskabeligt spørgsmål, man ønsker besvaret – det præcise spørgsmål bestemmer hvilke variable, der skal inkluderes i modellen.

Det er svært

– men den eneste måde at opnå egentlig videnskabelig indsigt!

(og så bliver det i øvrigt lettere bagefter at skrive en god artikel og lettere at svare på reviewernes kommentarer ...)



Automatisk variabelseleksion

Undgå så vidt muligt dette!

Der findes forskellige fremgangsmåder (se s. 92):

- ▶ **Forlæns selektion:**

Medtag hver gang den mest signifikante

Slutmodel: Weight BMP FEV1

- ▶ **Baglæns elimination:**

Start med alle, udelad hver gang den mindst signifikante

Slutmodel: Weight BMP FEV1

- ▶ **Forskellige hybrid-metoder**

Det ser da meget stabilt ud!?



men men men...

Hvis nu observation nr. 25 tilfældigvis ikke havde været med?

Så ville forlæns selektion have *taget højde ind som den første*,
og baglæns elimination ville have *smidt højde ud som den første*!

Tommelfingerregel (vedrører kun stabiliteten)

Antallet af observationer skal være mindst 10 gange så stort som
antallet af *undersøgte* parametre!

Advarsel ved variabelselektion

- ▶ **Massesignifikans!**
- ▶ *Enhver* variabelselektion baseret på signifikansvurderinger (dvs. også når I selv fjerner enkelte variable pga. statistisk insignifikans) vil give følgende effekt:
 - ▶ Signifikanserne overvurderes!
 - ▶ Regressionsparametrene er “for store”, dvs. for langt væk fra 0.
- ▶ **Automatisk** variabelselektion:
Hvad kan vi sige om 'vinderne'?
 - ▶ Var de hele tiden signifikante, eller blev de det lige pludselig?
 - ▶ I sidstnævnte tilfælde kunne de jo være blevet smidt ud, mens de var insignifikante. . .



Model med alle 9 kovariater

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,798 ^a	,637	,420	25,471

a. Predictors: (Constant), tlc, sex, bmp, age, rv, fev1, height, frc, weight

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	17101,390	9	1900,154	2,929	,032 ^b
	Residual	9731,250	15	648,750		
	Total	26832,640	24			

a. Dependent Variable: pemax

b. Predictors: (Constant), tlc, sex, bmp, age, rv, fev1, height, frc, weight



Output, fortsat

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	176,058	225,891		,779	,448
	age	-2,542	4,802	-,385	-,529	,604
	sex	-3,737	15,460	-,057	-,242	,812
	height	-,446	,903	-,287	-,494	,628
	weight	2,993	2,008	1,602	1,490	,157
	bmp	-1,745	1,155	-,627	-1,510	,152
	fev1	1,081	1,081	,362	1,000	,333
	rv	,197	,196	,507	1,004	,331
	frc	-,308	,492	-,403	-,626	,540
	tlc	,189	,500	,096	,377	,711



Model med alle 9 kovariater, II

Bemærk på outputtet s. 61-62:

- ▶ Alle kovariater er enkeltvis non-signifikante
- ▶ *Overall F-test* giver signifikans ($P=0.032$)

Og hvad kan vi overhovedet bruge modellen til....?



Baglæns elimination

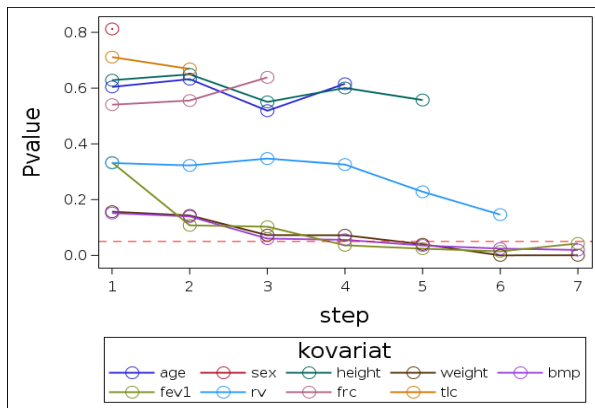
Tabel over successive p -værdier

	[,1]	[,2]	[,3]	[,4]	[,5]	[,6]	[,7]	[,8]	[,9]
age	0.604	0.632	0.519	0.616	-	-	-	-	-
sex	0.812	-	-	-	-	-	-	-	-
height	0.628	0.649	0.550	0.600	0.557	-	-	-	-
weight	0.157	0.143	0.072	0.072	0.040	0.000	0.000	0.000	0.001
bmp	0.152	0.140	0.060	0.056	0.035	0.024	0.019	0.098	-
fev1	0.333	0.108	0.103	0.036	0.024	0.014	0.043	-	-
rv	0.331	0.323	0.347	0.326	0.228	0.146	-	-	-
frc	0.540	0.555	0.638	-	-	-	-	-	-
tlc	0.711	0.669	-	-	-	-	-	-	-

Altman stopper ved skridt nr. 7, se figur næste side



P-værdier ved baglæns elimination



Ingen kode til denne figur

Krydsvalidering = Cross validation

Hvis man har en tilstrækkelig stor mængde data, er det en god fremgangsmåde at:

- ▶ Foretage modelfittet på en del af data, f.eks. to trediedele
- ▶ Afprøve modellen på resten bagefter, og måske blive lidt chokeret

Hvis modellen predikterer dårligt på den sidste trediedel, predikterer den nok også dårligt i nye situationer.



Sammenligning af modeller

Hvad sker der ved udeladelse af en forklarende variabel?

- ▶ Fittet bliver dårligere, dvs. residualkvadratsummen bliver større.
- ▶ Antallet af frihedsgrader (for residualkvadratsummen) stiger.
- ▶ Estimatet s^2 for residualvariansen σ^2 kan både stige og falde (fordi vi dividerer med frihedsgraderne)
- ▶ %-delen af variation, som forklares af modellen, R^2 , falder. Dette kompenseres der for i den *justerede determinationskoefficient* R_{adj}^2

Som **kriterium for, om modellen er god** kan vi altså bruge s^2 eller R_{adj}^2



Sammenligning af modeller

nemlig de to marginale modeller for `height` og `weight`, samt den multiple regressionsmodel med disse:

β_0	$\beta_1(\text{height})$	$\beta_2(\text{weight})$	s	R^2	Adj. $\bar{R}^2$
-33.276	0.932(0.260)	-	27.34	0.36	0.33
63.546	-	1.187(0.301)	26.38	0.40	0.38
47.355	0.147(0.655)	1.024(0.787)	26.94	0.40	0.35

- ▶ Hver af de to forklarende variable har betydning, vurderet ud fra de marginale modeller.
- ▶ I den multiple regressionsmodel ser *ingen af dem* ud til at have nogen betydning.
- ▶ Mindst en af dem har en betydning, men det er svært at sige hvilken. (Det ser dog mest ud til at være vægten.)



Sammenligning af modeller, II

Flere kommentarer til sammenligningen på forrige side:

- ▶ Størrelsen R^2 stiger næsten ikke fra den marginale model med kun `weight` som kovariat, til den multiple regressionsmodel (fra 0.4035 til 0.4049, dvs. uændret med 2 decimaler).
- ▶ Den justerede R^2 favoriserer derfor den simpleste af disse to modeller, hvor usikkerheden på parameterestimatet også er betydeligt lavere.
- ▶ Intercepterne i de 3 modeller kan ikke sammenlignes, da de refererer til helt forskellige individer, som i øvrigt alle er meget uinteressante: hhv. `height=0`, `weight=0` og `height=weight=0`

Modelkontrol: giver ikke frygtelig meget mening med den ratio af observationer/kovariater:



Sammenligning af kovariat-effekter

er vanskelig, bare at formulere:

Er effekten af vægt større end effekten af alder? Øh....

- ▶ Sammenligner vi P-værdier?
Det har kun noget med *evidensen* for en effekt at gøre, ikke med selve størrelsen af effekten.
- ▶ Man kan udregne såkaldte **standardiserede koefficienter**, som angiver effekten af 1 SD svarende til denne kovariat, i enheder af SD i outcome-variablen.

Er dette fornuftigt?

Næppe: Den biologiske effekt er nok ligeglad med, hvordan resten af personerne ser ud...

En sådan standardiseret koefficient vil afhænge af det aktuelle sample - ligesom en korrelation.

Kollinearitet: Kovariaterne er lineært relaterede

Det vil de altid til en vis grad være, undtagen i designede forsøg (f.eks. landbrugsforsøg)

Symptomer på kollinearitet:

- ▶ Visse af kovariaterne er stærkt korrelerede
- ▶ Nogle parameterestimerer har meget store standard errors
- ▶ Alle kovariater i den multiple regressionsanalyse er insignifikante, men R^2 er alligevel stor
- ▶ Der sker store forskydninger i estimererne, når en kovariat udelades af modellen
- ▶ Der sker store forskydninger i estimererne, når en observation udelades af modellen
- ▶ *Resultaterne er anderledes end forventet*



Ekstra udskrifter fra regressionsanalysen

- ▶ **Variance Inflation Factor** (`vif`):
Den faktor, som variansen af regressionskoefficienten er blevet ganget op med, på grund af kollineariteten mellem kovariaterne.
- ▶ **Tolerance** (`tol`):
Blot den reciprokke af VIF, altså $1/\text{VIF}$.
- ▶ **Standardiserede koefficienter** (`stb`):
som forklaret s. 70.

Se s. 92 for at få disse frem i SPSS.



Kollinearitet i praksis

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients			Collinearity Statistics
		B	Std. Error	Beta	t	Sig.	Tolerance
1	(Constant)	176,058	225,891		,779	,448	
	age	-2,542	4,802	-,385	-,529	,604	,046
	sex	-3,737	15,460	-,057	-,242	,812	,441
	height	-,446	,903	-,287	-,494	,628	,072
	weight	2,993	2,008	1,602	1,490	,157	,021
	bmp	-1,745	1,155	-,627	-1,510	,152	,141
	fev1	1,081	1,081	,362	1,000	,333	,185
	rv	,197	,196	,507	1,004	,331	,095
	frc	-,308	,492	-,403	-,626	,540	,058
	tlc	,189	,500	,096	,377	,711	,376

Coefficients^a

Model		Collinearity Statistics
		VIF
1	(Constant)	
	age	21,830
	sex	2,269
	height	13,955
	weight	47,781
	bmp	7,116
	fev1	5,420
	rv	10,538
	frc	17,143
	tlc	2,660

a. Dependent Variable: pemax

Tommelfingerregel: vif større end 5-10 stykker (dvs. to1 mindre end 0.2-0.1) er *kriminelt*....



Bemærk

Hvad gør man så, når der er kollinearitet?

1. Gennemtænk grundigt, hvad *den enkelte variabel* står for afhængigt af hvilke af de andre mulige variable, der fastholdes (= er med i modellen)
 - Overvej også, om *responsvariablen* ændrer fortolkning
2. Lav analyser af fokusvariablen med og uden justering for forskellige grupper af de andre variable og prøv at forstå forskellene i resultaterne
 - Præsenter gerne begge/alle analyser i artiklen med fortolkning af forskellene
3. Spar eventuelt på antallet af variable for grupper af variable, der hænger sammen: Drejer det sig om ét fælles aspekt af interesse, så kan man måske nøjes med én af variablene og begrunde, hvorfor man vælger netop den
4. **Fortolk med stor forsigtighed**



Vigtige pointer

- ▶ Man må *ikke* nøjes med at præsentere univariate analyser for alle variablene!

Som f.eks. s. 55

Problemet med fortolkningen forsvinder nemlig *ikke* af, at man tillægger hver enkelt variabel al forklaringsevnen en ad gangen.

- ▶ Man må *ikke* tro, at det er ligemeget hvilke effekter, man justerer for - for det videnskabelige spørgsmål ændrer sig jo.



Bemærk

Den statistiske model bestemmes af
det videnskabelige spørgsmål.

Eneste undtagelse fra denne regel er
fordelingsantagelserne,
der bestemmes af data.



Bemærk

Nogle kilder til forkerte modelvalg

Forsimple virkeligheden **for meget**, eksempelvis

- ▶ tro, at frasen *alt andet lige* giver mening (og ikke bare er en bekvem undskyldning for ikke at tænke det ordentligt igennem)
- ▶ tro, at det giver mening at tale om sammenhængen mellem exposure og respons, som om der kun kan eksistere ét eneste videnskabeligt spørgsmål, der involverer denne exposure og dette respons
- ▶ tro, at *confounder*er er givet ud fra exposure og respons, så man kan vælge sine kovariater uden at overveje konsekvenserne for fortolkningen



Et par falske påstande

Påstand: Når man skal vurdere effekten af en forklarende variabel (“exposure”) på et bestemt outcome, så skal man så vidt muligt justere for alle confoundere.

Sandhed: Nej! Man skal **kun** justere for de variable, som man gerne ville have kunnet holde fast – og man skal kunne gennemskue, hvilken konsekvens justering for hver eneste af de inkluderede confoundere har for fortolkningen af den estimerede effekt af exposure på outcome!



Et par falske påstande, II

Påstand: Man må ikke inkludere mediator variable, dvs. variable, der er en del af virkningsmekanismen.

Sandhed: Mediator variable er ligesom alle andre variable: Hvis man inkluderer dem, ændrer man det videnskabelige spørgsmål, der besvares ved analysen! Når man inkluderer en eller flere mediator variable, undersøger man størrelsen af den del af effekten, der *ikke* går via effekten på disse mediator variable, altså styrken af eventuelle *andre* virkningsmekanismer.



Eksempel på fejlslutning

Kan et øget fedtindtag være gavnligt for hjertet?

f.eks reducere risikoen for hjerteinfarkt, eller sænke kolesterol i blodet?

Tja.....:

Øget motion giver øget fødeindtag og dermed formentlig øget fedtindtag, så hvis man har fedtindtag som *eneste* kovariat, så kan det se gavnligt ud at spise meget fedt....

Bemærk

Påstand: Signifikansen for den enkelte variabel bliver altid svagere, når de andre tages med.

Sandhed: Ofte, men ikke altid. Nogle gange bliver signifikanserne væsentligt stærkere.

Vi så dette ved forelæsningsen om kovariansanalyse, hvor effekten af P-piller først kom frem, da vi justerede for alderen.



APPENDIX

med vejledninger svarende til nogle af slides

- ▶ Figurer: s. 83, 85-87, 89-91
- ▶ Regressionsanalyse: s. 84, 88, 91-92
- ▶ Modelkontrol: s. 85-87
- ▶ Prediktion: s. 89-90
- ▶ Diagnostics: s. 91-92



Tredimensionalt plot

Slide 7-8

Benyt Graphs/Graphboard Template Choser, marker de 3 involverede variable (vaegt, bpd og ad) og klik på Scatterplot eller Bubble Plot.

For at bytte om på akserne, klikkes på Detailed, hvor man så kan sætte vaegt på Y-aksen, bpd på X-aksen og ad på Z-aksen.

Man kan også evt. vælge en variabel til at angive farven (Color).

Regressionsanalyse

Slide 10

Den *minimale* opsætning:

- ▶ Benyt Analyze/Regression/Linear
- ▶ Sæt vægt over i Dependent
- ▶ Sæt både bpd og ad over i Independent(s)

Som minimum bør man dog også klikke Statistics og afkrydse Confidence Intervals.

Adskillige ekstra options kan tilføjes, som det vil fremgå af nogle af de efterfølgende sider (s. 85, 91-92)



Modelkontrol

Slide 15-21

Den automatiske modelkontrol i SPSS er ikke særlig god.....

Brug i stedet nedenstående:

I regressionsopsætningen fra s. 84 klikkes på Save, hvorefter man typisk vil vælge

- ▶ Under Predicted Values: vælg Unstandardized, som så kommer til at ligge i datasættet under navnet PRE_1
- ▶ Under Residuals er der 5 muligheder, f.eks.
 - ▶ Unstandardized: De *sædvanlige*, altså observeret minus forventet, kaldet RES_1 i datasættet
 - ▶ Studentized deleted: Normerede Press-residualer (altså hvor den pågældende observation er udeladt ved fit af modellen), kaldet SDR_1 i datasættet



Modelkontrol, II

Slide 15-21

Ud fra det udbyggede datasæt dannet s. 85 kan vi nu lave diverse figurer:

- ▶ Benyt Graph/Chart Builder, vælg Scatter, klik Unstandardized Predicted Values over på X-aksen og Unstandardized Residuals over på Y-aksen.
- ▶ Vælg Analyze/Descriptive Statistics/Q-Q Plots og sæt Unstandardized Residuals over i Variables.
- ▶ Benyt Graph/Chart Builder, vælg Histogram og klik det ønskede residual (f.eks. RES_1) ned på X-aksen. Klik så på Distribution Curve og afkryds Normal, klik Apply og Close



Check af varianshomogenitet

Slide 21

De standardiserede residualer hedder `ZRE_1`, og for at definere kvadratroden af de numeriske residualer går vi ind i Transform/Compute, skriver `sqrtres` i Target Variable og `sqrt(abs(ZRE_1))` i definitionsboksen.

Herefter fremstilles figurerne som sædvanligt i Graph/Chart Builder, og ved at dobbeltklikke på grafen efterfølgende, kan man lægge udglattede kurver på ved at klikke på Add Fit Line at Total og derefter i Properties-boksen afkrydse Loess/Apply.

Transformation af variable

Slide 23

Her skal vi definere logaritmetransformerede variable ved at benytte Transform/Compute, hvor vi først i Target Variable skriver logvaegt og i Numeric Expression skriver $\text{Ln}(\text{vaegt})$

Derefter gør vi det igen, men denne gang sætter vi logad som Target Variable, og i Numeric Expression skriver vi $\text{Ln}(\text{ad}/100)/\text{Ln}(1.1)$

Og endelig en sidste gang, hvor vi sætter logbpd som Target Variable, og i Numeric Expression skriver $\text{Ln}(\text{bpd}/90)/\text{Ln}(1.1)$



Prediktioner til brug for illustrationer

Slide 34 og 35

- ▶ Tilføj fiktive personer til datasættet, med manglende værdier af logvægt, men med de værdier af logbpd og logad (dvs. af bpd og ad), som man gerne vil lave prediktioner for (her 3 værdier af bpd og adskillige af ad)
- ▶ Lav en kolonne, kaldet f.eks. ny, som skal have værdien 0 for alle de oprindelige observationer, men 1 for de nye fiktive
- ▶ Estimer nu i samme model som før i det udbyggede datasæt og gem de predikterede værdier (nu hedder de sikkert PRE_2, fordi vi har gemt nogen tidligere)
- ▶ Transformer de predikterede værdier tilbage til oprindelig skala ved at benytte Transform/Compute Variable, skriv predikteret i Target Variable og `exp(PRE_2)` i Numeric Expression



Prediktioner til brug for illustrationer, II

Slide 34 og 35

Selve figurerne dannes således:

- ▶ Begræns data til de fiktive personer ved at gå ind i Data/Selet Cases, klik i If og skrive `ny=1`
- ▶ **Venstre figur:**
Gå ind i Graph/Chart Builder, vælg Lines (det til højre, der kan opdeles efter gruppe), og i den fremkomne boks trækkes PRE_2 over på Y-aksen, logad over på X-aksen og bpd over i Set Color
- ▶ **Højre figur:**
Som venstre, blot med predikteret på Y-aksen i stedet for PRE_2, og ad på X-aksen i stedet for logad



Diagnostics

Slide 37-38

Vi anvender her igen Save-knappen i regressionsopsætningen (se s. 85), og krydser det af, vi gerne vil bruge, nemlig Cook's og under Influence Statistics: Standardized DfBetas (og evt. også DfBetas(s))

Herefter benyttes Graph/Chart builder/Scatter som adskillige gange tidligere.



Variabelselektion og kollinearitet

Slide 58, 64

I regressionsopsætningen (se s. 84) klikkes på Method, der nu kan skiftes fra Enter ("*den kontrollerede og anbefalede metode*") til Backward eller Forward

Slide 71-73

I regressionsopsætningen Analyze/Regression/Linear (se s. 84) klikkes på Statistics, hvor man kan vælge Collinearity Diagnostics

