

Faculty of Health Sciences

Basal statistik

Korrelerede målinger, SPSS

Lene Theil Skovgaard

9. november 2020



Korrelerede målinger

- ▶ Tilfældige effekter
- ▶ Varianskomponentmodeller
- ▶ Modeller for longitudinelle målinger
- ▶ Korrelationsstrukturer
- ▶ Baseline problematik

E-mail: ltsk@sund.ku.dk

*: Siden er lidt teknisk



En gennemgående antagelse hidtil

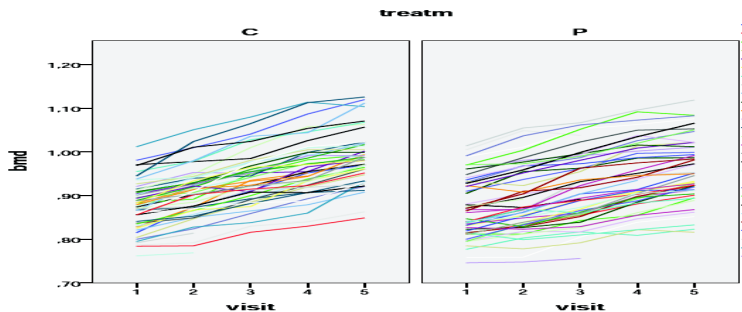
Uafhængighed

- ▶ Kun en enkelt observation for hvert individ (unit) (bortset fra den parrede situation, med to målemetoder eller før/efter undersøgelser)
- ▶ Ingen tvillinger/søskende ...

men i dag har vi flere end to for hvert individ, og vi må forvente, at observationer fra samme individ ser *mere ens* ud end observationer fra forskellige individer, de er **korrelerede**.

Eksempel: Calcium tilskud

Spaghettiplot: Individuelle profiler (se fremgangsmåde s. 99-100)



Desværre måtte specialmarkering af forløbene med manglende værdier opgives.

Eksempel: Calcium tilskud

til styrkelse af knogledannelse hos unge piger

I alt 112 11-årige piger randomiseres til at få 'tilskud' i form af enten calcium eller placebo.

Observationer: Y_{git} (gruppe g , pige=individ i , visit=tid t)

Outcome: BMD=bone mineral density, i $\frac{g}{cm^2}$, måles 5 gange over 2 år (hvert halve år).

Det videnskabelige spørgsmål:

Kan et calciumtilskud øge knoglevæksten hos præ-pubertale piger?
og i givet fald - hvor meget?



Datastruktur ved gentagne målinger

Hvert individ bidrager med ligeså mange linier, som vedkommende har observationer

- ▶ 5 linier for de fleste piger
- ▶ En variabel, der angiver nummeret på (tiden for) målingen

Obs	treatm	girl	visit	bmd
1	C	101	1	0.815
2	C	101	2	0.875
3	C	101	3	0.911
4	C	101	4	0.952
5	C	101	5	0.970
6	P	102	1	0.813
7	P	102	2	0.833
8	P	102	3	0.855
9	P	102	4	0.881
10	P	102	5	0.901
11	P	103	1	0.812
.	.	.	.	.

Bemærk: I SAS-versionen hedder treatm blot grp



Terminologi for korrelerede målinger

- ▶ Multivariate responser (berøres ikke her):
Forskellige outcomes (responser) på det samme individ, f.eks. en række hormonmålinger, der skal vurderes samlet.
- ▶ Cluster design:
Samme outcome (respons) målt på f.eks. alle individer i en række familier.
- ▶ Repeated measurements / Gentagne målinger:
Samme outcome (respons) målt i forskellige situationer (eller på forskellige steder) på samme individ.
- ▶ Longitudinelle målinger:
Samme outcome (respons) målt gentagne gange over tid for samme individ.



Mixed models for korrelerede målinger

To typer af generaliseringer:

- ▶ **Varianskomponentmodeller**

Generelle lineære modeller, hvor nogle af effekterne, typisk intercept og evt. hældning (dvs. effekt af tid) er gjort *tilfældige* (dvs. varierer mellem individer)

- ▶ **Kovariansstrukturer**

Generelle lineære modeller, hvor korrelationen mellem målinger (og evt. forskelle i varians) specificeres direkte.

Begge disse kaldes under et for **Mixed Models**

Mix af systematiske og tilfældige (*random*) effekter

Hvis korrelationen ignoreres, får man **forkerte spredninger**,
(for små eller for store,...)



Tilfældige effekter

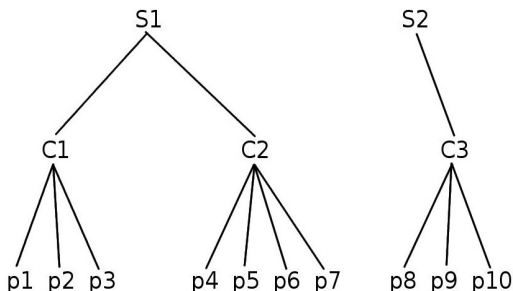
Variationskilder (varianskomponenter)

- ▶ geografisk/miljømæssig variation
 - ▶ mellem regioner, hospitaler, skoler eller lande
- ▶ biologisk variation
 - ▶ variation mellem individer, familier eller dyr
- ▶ variation mellem målinger på samme individ (within-individual)
 - ▶ variation mellem arme, tænder, injektionssteder, (dage)
- ▶ variation, der skyldes ukontrollable forsøgsomstændigheder
 - ▶ tidspunkt på dagen, temperatur, observatør
- ▶ målefejl/målevariation



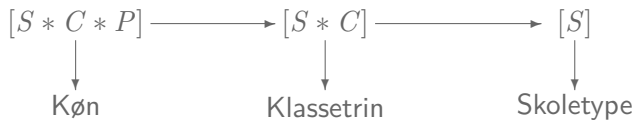
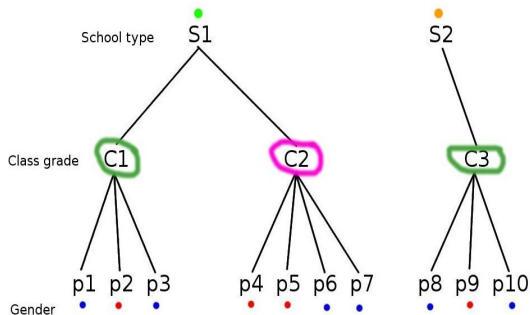
Cluster design (hierarkisk design)

f.eks. skole (School), skole klasse (Class) og elev (Pupil)



$$[I] = [S * C * P] \longrightarrow [S * C] \longrightarrow S$$

Cluster design, med kovariater



Et klassisk set-up

Gentagne målinger på samme individ, ofte over tid.

Vi skelner imellem forskellige typer af kovariater:

- ▶ **Within**-individuals kovariat (**level 1**):
En kovariat, der varierer *indenfor* person,
f.eks. **tid**, dosis, blodtryk
Disse effekter estimeres svarende til en **parret sammenligning**.
- ▶ **Between**-individuals kovariat (**level 2**):
En kovariat, der *ikke* varierer indenfor person, men kun
mellem personer, f.eks. **behandling**, køn, genotype
Disse effekter estimeres svarende til en **uparret sammenligning**.

Individer kunne også være *clustre*, såsom
familier, hospitaler, klasser etc.



Hvis man ignorerer korrelationen

mellem målinger på samme individ opstår der fejl

- ▶ lav efficiens (type 2 fejl) ved vurdering af level 1 kovariater (within-individual effekter)
- ▶ for små standard errors (type 1 fejl) for estimer af level 2 effekter (between-individual effekter)
- ▶ muligvis bias i middelværdistruktur (i tilfælde af missing values, eller ubalanceret design)

Det er altså vigtigt at medtage alle relevante variationskilder!



Formål med “mixed effects” designs

- ▶ At opnå **større præcision** ved vurdering af effekt af tid, dosis osv., ikke mindst interaktionen mellem tid og evt. behandling (udvikler de to grupper sig forskelligt?)
- ▶ At opnå **viden** om den relative størrelse af de forskellige varianskomponenter, for derved i fremtidige undersøgelser at kunne bruge ressourcerne der, hvor de er bedst anbragt, eksempelvis:
 - ▶ Hvor ofte skal man måle outcome?
 - ▶ Skal man have data på mange børn i hver klasse, eller hellere mange klasser på hver skole?



Simpelt eksempel: Hævelse ved vaccination

Eksperiment:

6 kaniner, hver især vaccineret 6 gange, forskellige steder på ryggen

Outcome y_{rs} : hævelse i cm^2 , med notationen

$r = 1, \dots, R=6$ angiver kaninen (**rabbit**),

$s = 1, \dots, S=6$ angiver stedet (**spot**)

Formål med undersøgelsen:

- ▶ Giv et estimat for “*overall mean*”
- ▶ Hvor mange gange kan kaninerne *genbruges*, altså hvor mange stik kan det svare sig at udføre pr. kanin?

Vi har i alt 36 målinger af hævelse, **men** vi må forvente, at hævelse kan være specifikt for den enkelte kanin, således at nogle har tendens til stor hævelse, andre til mindre hævelse.



Data - omstrukturering, se s. 101

Oprindeligt datasæt (*bredt*):

```
kanin a b c d e f
1 7.9 6.1 7.5 6.9 6.7 7.3
2 8.7 8.2 8.1 8.5 9.9 8.3
3 7.4 7.7 6 6.8 7.3 7.3
4 7.4 7.1 6.4 7.7 6.4 5.8
5 7.1 8.1 6.2 8.5 6.4 6.4
6 8.2 5.9 7.5 8.5 7.3 7.7
```

skal omdannes til *langt* (kaldet *rabbit*):

```
rabbit sted swelling
1 a 7.9
2 a 8.7
3 a 7.4
4 a 7.4
5 a 7.1
6 a 8.2
1 b 6.1
. . .
```

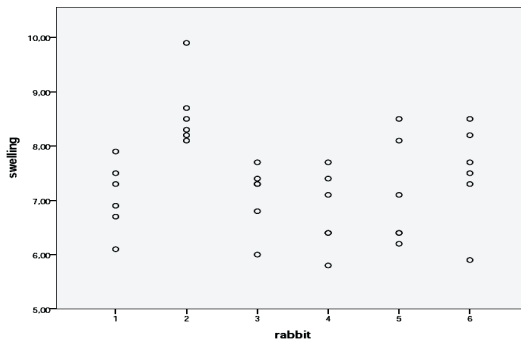


Scatter plot

36 observationer, 2 variable:

Outcome: Hævelse = swelling

X-akse: Arbitrær nummerering af kaniner



Naiv kvantificering af hævelse

Fra Analyze/Descriptive Statistics/Explore får vi (bl.a.):

Case Processing Summary

	Valid		Cases Missing		Total	
	N	Percent	N	Percent	N	Percent
swelling	36	100.0%	0	0.0%	36	100.0%

Descriptives

		Statistic	Std. Error
swelling	Mean	7,3667	,15523
	95% Confidence Interval for Mean	Lower Bound	7,0515
		Upper Bound	7,6818
	5% Trimmed Mean	7,3401	
	Median	7,3500	
	Variance	,867	
	Std. Deviation	,93136	

Hvad er der galt med denne kvantificering?

Tænk, hvis alle målinger på samme kanin var *helt ens*?
Så har vi jo i virkeligheden kun 6 observationer.

Men de er jo ikke *helt ens*, hvad så?

Analyse (med hovedet under armen)

- ▶ Hver kanin har *et niveau* (en middelværdi)
- ▶ Herudover er der *variation mellem indstikssteder*

I computer sprog:

Kaninen er en **faktor**, analysen er en ensidet variansanalyse (Analyze/General Linear Model/Univariate):

Tests of Between-Subjects Effects

Dependent Variable: swelling

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	12,833 ^a	5	2,567	4,393	,004
Intercept	1953,640	1	1953,640	3344,002	,000
rabbit	12,833	5	2,567	4,393	,004
Error	17,527	30	,584		
Total	1984,000	36			
Corrected Total	30,360	35			

a. R Squared = ,423 (Adjusted R Squared = ,326)



Output, fortsat

Parameter Estimates

Dependent Variable: swelling

Parameter	B	Std. Error	t	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Intercept	7,517	,312	24,089	,000	6,879	8,154
[rabbit=1]	-,450	,441	-1,020	,316	-1,351	,451
[rabbit=2]	1,100	,441	2,493	,018	,199	2,001
[rabbit=3]	-,433	,441	-,982	,334	-1,335	,468
[rabbit=4]	-,717	,441	-1,624	,115	-1,618	,185
[rabbit=5]	-,400	,441	-,906	,372	-1,301	,501
[rabbit=6]	0 ^a	.	.	.	.	.

a. This parameter is set to zero because it is redundant.

Kaninerne har signifikant forskellige niveauer ($P=0.0040$)

Men: Er det overhovedet interessant information?

Det er i hvert fald ikke svar på spørgsmålet!

- ▶ Vi er jo ikke interesserede i *disse specifikke* 6 kaniner, men snarere kaniner i al almindelighed, *som art betragtet*!
- ▶ Vi antager, at disse 6 kaniner er *tilfældigt udvalgt*



Varianskomponentmodel

Vi vil i stedet se på en model, hvor forskellen på kaniner modelleres som en ekstra variationskilde:

$$y_{rs} = \mu + a_r + \varepsilon_{rs}$$

hvor a_r 'erne og ε_{rs} 'erne antages at være *uafhængige*, normalfordelte med

$$\text{Var}(a_r) = \omega_B^2, \quad \text{Var}(\varepsilon_{rs}) = \sigma_W^2$$

Variationen mellem kaniner er nu en **tilfældig effekt (random factor)**,

ω_B^2 og σ_W^2 er **varianskomponenter**, og modellen kaldes også en **two-level model**



Hvad indebærer denne varianskomponentmodel

Alle observationer af hævelse har **samme middelværdi** og **samme varians** (summen af varianskomponenterne):

$$y_{rs} \sim N(\mu, \omega_B^2 + \sigma_W^2)$$

Men: Målinger foretaget på den samme kanin er korrelerede, med **intra-class korrelationen**

$$\text{Corr}(y_{r1}, y_{r2}) = \rho = \frac{\omega_B^2}{\omega_B^2 + \sigma_W^2}$$

Målinger på samme kanin har en tendens til at se **mere ens** ud end målinger på forskellige kaniner.

Alle målinger på samme kanin ser **lige ens ud**. Denne korrelationsstruktur kaldes **compound symmetry (CS)** eller **exchangeability**.



Estimation i varianskomponentmodel

Her skal man sørge for at definere

- ▶ rabbit som *random factor*
- ▶ En *tom* middelværdi
(de har alle den samme, nemlig interceptet, μ)

Benyt Analyze/Mixed Models/Linear, sæt rabbit over i Subjects og klik Continue.

Nu sættes swelling i Dependent Variable og rabbit i Factor(s). Under Fixed sættes ingenting, men under Random sættes rabbit over i Model, hvorefter man klikker Add/Continue.

Under Statistics, afkrydser vi Parameter estimates.



Output fra varianskomponentmodel

fra modellen s. 23

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Intercept	7,366667	,267014	5	27,589	,000	6,680286	8,053047

a. Dependent Variable: swelling.

Estimates of Covariance Parameters^a

Parameter	Estimate	Std. Error
Residual	,584222	,150846
rabbit Variance	,330407	,271716

a. Dependent Variable: swelling.

Sammenlignes med $CI=(7.052, 7.682)$ fra tidligere (s. 18):
 At ignorere korrelationen fører her til en **type 1** fejl,
 nemlig *for smalt CI*.



Fortolkning af varianskomponenter

Værdierne er taget fra output s. 24:

Variation	Varianskomponent	Estimat	Andel af variationen
Between	ω_B^2	0.3304	36%
Within	σ_W^2	0.5842	64%
Total	$\omega_B^2 + \sigma_W^2$	0.9146	100%

Typiske forskelle (95% prædiktionsintervaller):

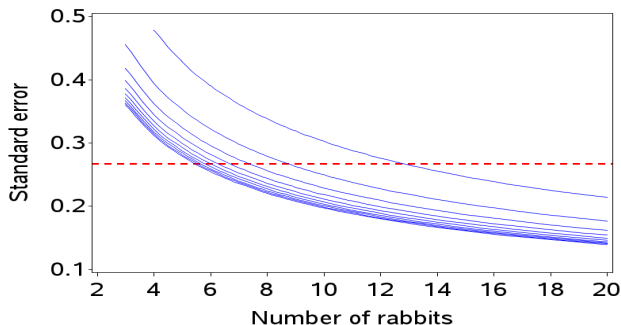
- ▶ for to steder på den **samme** kanin
 $\pm 2 \times \sqrt{2 \times 0.5842} = \pm 2.16 \text{ cm}^2$
- ▶ for to steder på **forskellige** kaniner
 $\pm 2 \times \sqrt{2 \times 0.9146} = \pm 2.70 \text{ cm}^2$



Standard error på estimat for hævelse

For R =antal kaniner, fra 3 til 20:

For S =antal injektionssteder, fra 1 til 10:



$$\text{Var}(\bar{y}) = \frac{\omega_B^2}{R} + \frac{\sigma_W^2}{RS}$$

*Den “effektive” sample size

Hvis vi kun havde **en enkelt** observation for hver af k kaniner, hvor mange kaniner skulle vi så bruge til at få samme præcision?

$$k = \frac{R \times S}{1 + \rho(S - 1)}$$

$$\text{Her har vi } \rho = \frac{\omega_B^2}{\omega_B^2 + \sigma_W^2} = \frac{0.3304}{0.3304 + 0.5842} = 0.361 \Rightarrow k = 12.8$$

Vi har altså, effektivt set, kun hvad der svarer til to uafhængige observationer fra hver kanin!



Longitudinelle målinger

Eksempel: Aspirin optagelse for raske og syge
(Matthews et.al.,1990)

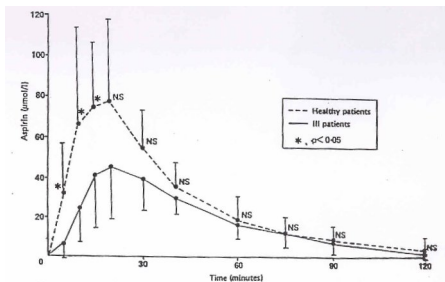


FIG 2—Mean and standard deviation of aspirin concentrations in nine healthy and nine ill patients over time. ("Usual" method of display)

T-tests for hvert tidspunkt:

- ▶ massesignifikansproblem
- ▶ testene er ikke uafhængige
- ▶ fortolkning kan være vanskelig

Pas på med gennemsnit - og gennemsnitskurver

specielt når der er *missing values*

- ▶ De giver ingen fornemmelse for variationen mellem personer
Individuelle forløb bør altid tegnes
(men ikke nødvendigvis publiceres)
- ▶ De kan skjule vigtige strukturer, hvis tidsforløbene adskiller sig
kvalitativt fra hinanden,
(hvis de ikke er rimeligt parallelle):

Alternativ:

Regn videre på individuelle karakteristika



Individuelle tidsprofiler

Spaghettiplot

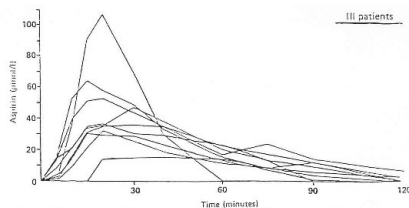
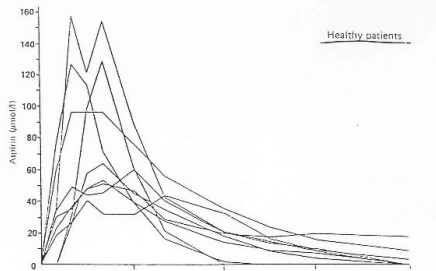


FIG 4—Individual plots of aspirin concentrations against time in healthy patients and ill patients

Har de samme facon allesammen?
Er gennemsnittet repræsentativt?

Analyse af udvalgte karakteristika

Eksempelvis:

- ▶ Endpoint (sidste måling), hældning
dur helt klart ikke i dette eksempel
- ▶ Maksimal værdi, toppunkt, og dettes placering
- ▶ AUC, Arealet under kurven

Sådanne karakteristika sammenlignes ved hjælp af traditionelle metoder, typisk et T-test.

Husk (som altid) konfidensintervaller



Fordele og ulemper ved de simple analyser

Fordele:

- ▶ simple (at forstå og forklare til andre)
- ▶ for det meste simple at udføre

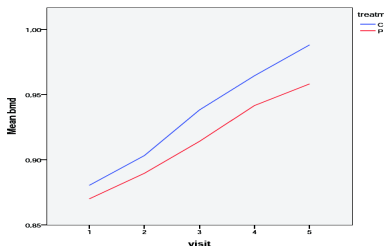
Ulemper:

- ▶ Informationstab
- ▶ Kræver et passende antal observationer pr. individ
- ▶ Kræver et fornuftigt karakteristika og ofte målinger på ens tidspunkter
- ▶ Umuligt at inddrage tidsafhængige kovariater
- ▶ Kan ikke tage hensyn til baseline



Eksempel: Calcium tilskud

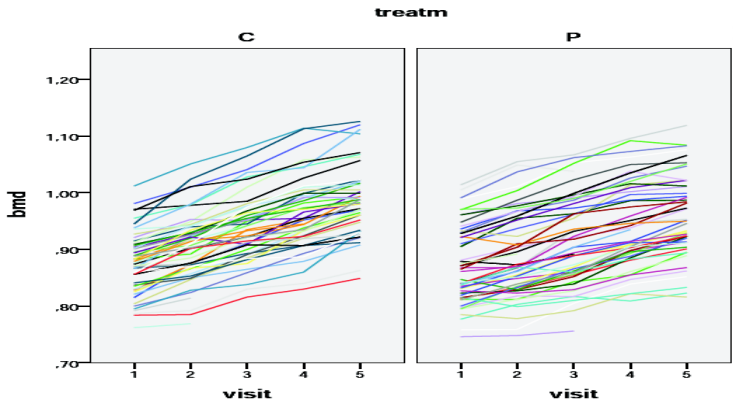
Gennemsnit for de 2 grupper, (se fremgangsmåde s. 104)



men er gennemsnitskurver rimeligt her?

Ja, nogenlunde....pga de ret ensartede forløb, se næste side,
men der er missing values... og baseline issues

Individuelle profiler - igen (som s.4)



Simple analysemuligheder

Udvælg et (eller flere) karakteristika til

sammenligning af grupperne:

- ▶ **Gennemsnit for hvert tidspunkt:** Ingen evidens for forskel
husk korrektion for massesignifikans.
- ▶ **Ændringer for successive tidspunkter:**
husk korrektion for massesignifikans.
Ændring fra start til slut er større i Calcium-gruppen:
0.0190 (0.0062, 0.0318), $P=0.004$
- ▶ Individuelle hældninger er større i Calcium-gruppen:
Estimat for forskel: **0.0039 (0.0019), $P=0.050$**

men disse analyser er **suboptimale**

– hvis andet er muligt...



Struktur for Calcium-eksemplet

	Unit/Enhed	Variation	Kovariater
Level 2	piger	<i>mellem</i> piger ω_B^2	treatm
Level 1	enkeltobservationer	<i>indenfor</i> piger σ_W^2	visit treatm*visit

Vi er specielt interesserede i *within*-kovariaten treatm*visit, fordi den udtrykker en *forskel i mønsteret* over tid,



Modelspecifikation i praksis

Mixed model, med

Systematiske effekter: De to faktorer: `treatm` og `visit`, samt en interaktion mellem disse,
dvs. *ingen bindinger på tidsudviklingen*

Tilfældige effekter: `girl` som *random factor*

Dette specificeres i SPSS således:

Benyt Analyze/Mixed Models/Linear, sæt `girl` over i Subjects og klik Continue

Sæt nu `bmd` i Dependent Variable og såvel `treatm`, `visit` og `girl` i Factor(s).

fortsættes næste side...



Modelspecifikation i praksis, fortsat

Under Fixed markeres både `treatm` og `visit`, der vælges `Factorial` og klikkes `Add`.

Under Random skal vi benytte `Build nested terms` og markerer først `girl` og trykker pil ned. Derefter klikkes `Within`, man markerer `treatm`, trykker pil ned og klikker `Add/Continue`, således at effekten fremstår som `girl(treatm)`.

Under Statistics, afkrydser vi `Parameter estimates`.



Bemærkning til modelspecifikation

- ▶ Pigerne er **nestet** i grupperne,
Vurdering af gruppeforskelle er uparret

Nestingen skrives som den tilfældige effekt `girl(treatm)`

- ▶ Alle visits forekommer (i princippet) for den enkelte pige,
Vurdering af tidseffekter er parrede



Systematisk vs. tilfældig effekt?

Systematisk = Fixed:

- ▶ Alle faktorens niveauer observeres (typisk kun et par stykker, f.eks. et antal **behandlinger**, eller nogle **tidspunkter**)
- ▶ Der kan kun drages konklusioner om **netop disse** behandlinger (ikke om andre *ikke-benyttede* behandlingstyper)
- ▶ Vi er interesserede i forskellen på de enkelte behandlinger
- ▶ Der skal være et rimeligt antal observationer for hvert niveau af faktoren (ikke for få i nogen behandlingsgrupper)



Systematisk vs. tilfældig effekt?, fortsat

Tilfældig = Random:

- ▶ En repræsentativ stikprøve (sample) af faktorniveauer er observeret
f.eks. et antal patienter, skoleklasser etc.
- ▶ Vi er ikke interesserede i netop disse individer, eller disse skoleklasser, *men*
- ▶ vi ønsker at drage konklusioner *i al almindelighed*, dvs.
for **andre** patienter, skoleklasser eller kaniner,
- ▶ Det er **nødvendigt** at lade faktoren være tilfældig, hvis man interesserer sig for kovariater hørende til dette level, f.eks. behandling, klassesettrin...
eller selve niveauet (af hævelsen)



Output fra mixed model fra s. 37-38

(idet der dog af pladshensyn ikke er konfidensgrænser på estimerterne på næste side....)

Covariance Parameters

Estimates of Covariance Parameters^a

Parameter	Estimate	Std. Error
Residual	,000235	1,700832E-5
girl(treatm) Variance	,004439	,000608

a. Dependent Variable: bmd.

Type III Tests of Fixed Effects^a

Source	Numerator df	Denominator df	F	Sig.
Intercept	1	109,913	21261,517	,000
treatm	1	109,913	2,629	,108
visit	4	381,551	619,421	,000
treatm * visit	4	381,551	5,297	,000

a. Dependent Variable: bmd.

Analysen viser en hel klar **interaktion** treatm*visit, dvs. at der *ikke* er parallelle tidsforløb i de to grupper.

... hvis modellen altså er rimelig

Output, fortsat

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.
Intercept	,957569	,009131	122,183	104,869	,000
[visit=1]	-,087499	,003100	382,386	-28,224	,000
[visit=2]	-,067483	,003103	381,193	-21,750	,000
[visit=3]	-,043422	,003117	381,016	-13,931	,000
[visit=4]	-,016185	,003148	380,886	-5,142	,000
[visit=5]	0 ^b	0	.	.	.
[treatm=C] * [visit=1]	,010384	,012922	118,295	,804	,423
[treatm=C] * [visit=2]	,016957	,012964	119,780	1,308	,193
[treatm=C] * [visit=3]	,023290	,012996	120,922	1,792	,076
[treatm=C] * [visit=4]	,022718	,013021	121,856	1,745	,084
[treatm=C] * [visit=5]	,029506	,013037	122,447	2,263	,025
[treatm=P] * [visit=1]	0 ^b	0	.	.	.
[treatm=P] * [visit=2]	0 ^b	0	.	.	.
[treatm=P] * [visit=3]	0 ^b	0	.	.	.
[treatm=P] * [visit=4]	0 ^b	0	.	.	.
[treatm=P] * [visit=5]	0 ^b	0	.	.	.
[treatm=C]	0 ^b	0	.	.	.
[treatm=P]	0 ^b	0	.	.	.

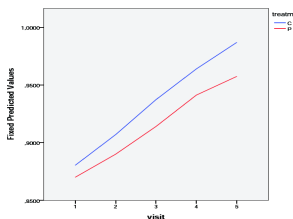
Se bemærkninger til output på s. 46



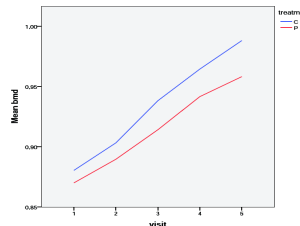
Figur af predikterede kurver

svarer ikke helt til gennemsnittene fra s. 33

Prediktioner:



Gennemsnit:

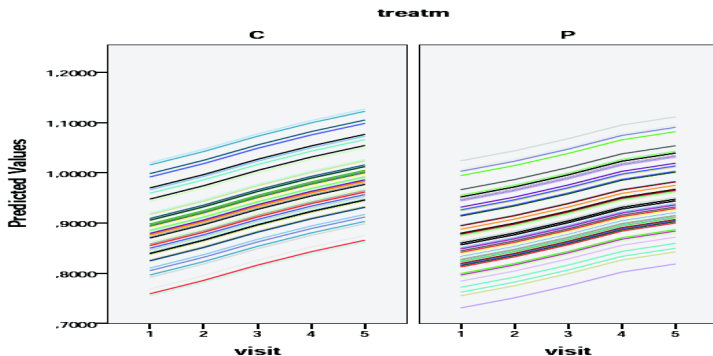


Dette skyldes de enkelte missing values....

Vedr. venstre figur: se s. 107

Predikterede individuelle forløb

som inkluderer de **tilfældige niveauer** for pigerne, se s. 108



Foreløbig konklusion om calcium-eksemplet

- ▶ De predikterede forløb (2×5 estimerede middelværdier) er vist s. 44, venstre plot
- ▶ Vi fandt (s. 42) en signifikant interaktion `treatm*visit`, dvs. grupperne udvikler sig forskelligt over tid ($P = 0.000$)
- ▶ På s. 43 ses, at forskellen på grupperne ved sluttidspunktet er 0.02951 (0.01304) i C-gruppens favør (aflæses ud for `[treatm=C]*[visit=5]`) samt at dette er signifikant ($P=0.025$)
- ▶ Forskellen stiger over tid, hvilket også er forventeligt ved kontinuerlig behandling
- ▶ Intra-individ korrelationen er 0.95 (se s. 53)



Modelkontrol

To typer residualer bør checkes:

De sædvanlige Observeret minus predikteret (*gruppe*-)middelværdi
(kun systematiske effekter fratrækkes),
gerne i *skaleret* version, hvis det er muligt

De betingede Observeret minus predikteret *individue*l prediktion
(både systematiske og tilfældige effekter fratrækkes),
Conditional (betinget med individet)

Vi ser på *de sædvanlige* tegninger:

- ▶ Residualer plottet mod predikterede værdier
- ▶ Histogram og fraktildiagram

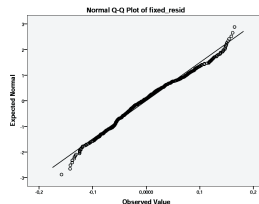
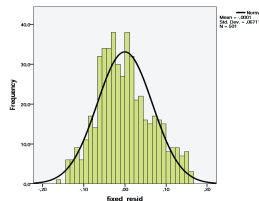
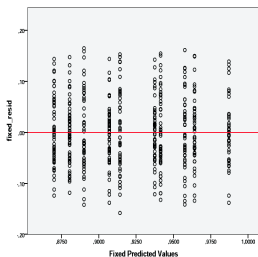
Se figurerne s. 48 og 49 (se konstruktion af disse s. 109)



Modelkontrol, sædvanlige (store) residualer

Afvigelse fra **gruppe**-middelværdier (se s. 47):

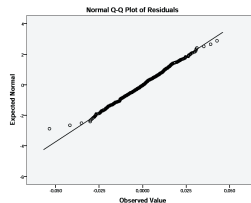
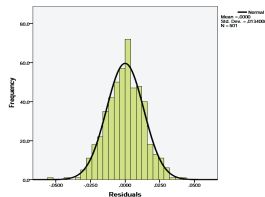
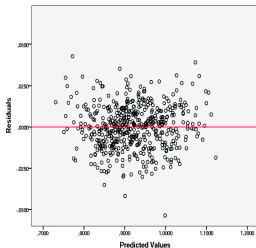
Check af varianshomogenitet
og normalfordelte residualer



Modelkontrol, betingede (små) residualer

Afvigelse fra **individuelle** prediktioner (se s. 47):

Check af varianshomogenitet
og normalfordelte residualer



Effekt af korrelerede observationer

Når korrelationen **ignoreres**, sker der følgende:

- ▶ Level 1 kovariater (tidsrelaterede effekter, “*within*”):
Lav styrke (**type 2 fejl**):
Vi opdager ikke de effekter, der er der.
Her drejer det sig om $\text{treatm} \times \text{visit}$, samt (hvis denne ikke er med i modellen) om selve visit
- ▶ Level 2 kovariater (behandlinger, gruppe, “*between*”):
For små standard errors, og dermed for små P-værdier (**type 1 fejl**):
Vi kommer til at finde effekter, der slet ikke er der.
Her er dette kun relevant i en model *uden interaktion*, hvor vi evt. vil vurdere en generel forskel på grupperne (treatm).



Fejlagtig analyse

Tosidet ANOVA, **korrelation ignoreret**

Tests of Between-Subjects Effects

Dependent Variable: bmd

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	,694 ^a	9	,077	16,820	,000
Intercept	425,874	1	425,874	92897,855	,000
treatm	,051	1	,051	11,046	,001
visit	,644	4	,161	35,103	,000
treatm * visit	,006	4	,002	,354	,841
Error	2,251	491	,005		
Total	428,763	501			
Corrected Total	2,945	500			

a. R Squared = ,236 (Adjusted R Squared = ,222)

Type 2 fejl: Her findes fejlagtigt ingen vekselvirkning, da vi har “glemt” parringen

Vi kan altså *ikke* se, at de to grupper udvikler sig forskelligt....

Synonymer for modellen s. 37-38

- ▶ Two-level model
- ▶ Varianskomponentmodel (med 2 varianskomponenter)
- ▶ Model med tilfældig person effekt
- ▶ Model med tilfældige intercepter (niveauer)
- ▶ Model med “compound symmetry” kovariansstruktur (eller “exchangeability” kovariansstruktur)

Korrelation = Kovarians, normeret med spredninger

se mere s. 110



Compound symmetry = Exchangeability

Varianskomponentmodellen antager, at alle målinger har **samme varians** ($\omega_B^2 + \sigma_W^2$) og at alle par af observationer *på samme individ* er **lige stærkt korrelerede**:

$$\text{Corr}(Y_{git_1}, Y_{git_2}) = \rho = \frac{\omega_B^2}{\omega_B^2 + \sigma_W^2}$$

kaldet **intra-class korrelationen**

Korrelationen ρ estimeres ud fra output s. 42

$$\hat{\rho} = \frac{0.004439}{0.004439 + 0.000235} \approx 0.95$$

Observationerne er **ombyttelige=exchangeable**,
altså **CS: Compound Symmetry**



Korrelationsstruktur

Her har vi 5 tidspunkter, og derfor en 5*5 korrelations-matrix, hvor entry (i,j) angiver korrelationen mellem visit i og visit j .

Compound Symmetry har strukturen:

$$\begin{pmatrix} 1 & \rho & \rho & \rho & \rho \\ \rho & 1 & \rho & \rho & \rho \\ \rho & \rho & 1 & \rho & \rho \\ \rho & \rho & \rho & 1 & \rho \\ \rho & \rho & \rho & \rho & 1 \end{pmatrix}$$

som siger, at **alle tids-par er lige korrelerede**,
se kode s. 110

Det betyder, at der slet ikke tages hensyn til, at observationerne er taget over tid, altså i en bestemt rækkefølge!!

Er det en fornuftig antagelse?



Valg af varians- og korrelationsstruktur?

Den mest generelle: **Ustruktureret**:

Benyt Analyze/Mixed Models/Linear, sæt girl over i Subjects, visit over i Repeated og skift Repeated Covariance Type til Unstructured.

Gå herefter ind i Random og sørg for, at der ikke står noget her (fjern evt. girl fra denne). Nu sættes (eller bibeholdes) bmd i Dependent Variable og såvel treatm, visit og girl i Factor(s). Under Fixed markeres både treatm og visit, der vælges Factorial og klikkes Add.

Under Statistics, afkrydser vi Parameter estimates og evt. Covariances of residuals.



Ustruktureret kovarians

Denne struktur er *helt uden bånd*, både på korrelation og på de 5 spredninger, i alt 15 parametre, dog *antaget ens i de to grupper*

Spredningsestimater (måske svagt stigende):

Visit 1	Visit 2	Visit 3	Visit 4	Visit 5
0.0629	0.0688	0.0705	0.0730	0.0700

Korrelations strukturen bliver så (ikke hentet fra SPSS):

Row	Col1	Col2	Col3	Col4	Col5
1	1.0000	0.9699	0.9414	0.9250	0.8987
2	0.9699	1.0000	0.9727	0.9585	0.9399
3	0.9414	0.9727	1.0000	0.9809	0.9592
4	0.9250	0.9585	0.9809	1.0000	0.9755
5	0.8987	0.9399	0.9592	0.9755	1.0000



Skal man vælge ustruktureret kovarians?

Fordele:

- ▶ Vi *tvinger* ikke en forkert kovariansstruktur ned over vores observationer
- ▶ Vi får indsigt i mønsteret i spredninger og korrelationer (hvis der er tilstrækkelig information, dvs. mange personer og ikke alt for mange tidspunkter)

Ulemper:

- ▶ Vi bruger en masse parametre på beskrivelsen af modellen kovariansen. Resultatet kan derfor blive ustabil, og kan derfor ikke anvendes for små datasæt
- ▶ Den kan kun bruges for balancerede data (alle individer skal være målt til de samme tidspunkter)

så måske hellere noget “midt imellem”?



Mulige korrelationsstrukturer

Observationer, der er foretaget **tæt på hinanden** i tid vil formentlig være **stærkere korrelerede** end observationer, der tidsmæssigt ligger længere fra hinanden.

Der findes (forfærdeligt) mange muligheder

- ▶ **Autoregressiv** struktur (se s. 111)
- ▶ **Autoregressiv** struktur, *samtidig* med den tilfældige effekt (niveau) for hvert individ
- ▶ Flere tilfældige effekter, f.eks. tillige en tilfældig hældning **Random regression** (kommer lidt senere)
- ▶ Et væld af andre, prøv f.eks. at google *"sas mixed repeated type"*



Autoregressiv kovariansstruktur, AR(1)

Hvis tiderne er ækvidistante, er det følgende struktur

$$(\omega_B^2 + \sigma_W^2) \begin{pmatrix} 1 & \rho & \rho^2 & \rho^3 & \rho^4 \\ \rho & 1 & \rho & \rho^2 & \rho^3 \\ \rho^2 & \rho & 1 & \rho & \rho^2 \\ \rho^3 & \rho^2 & \rho & 1 & \rho \\ \rho^4 & \rho^3 & \rho^2 & \rho & 1 \end{pmatrix}$$

dvs. korrelationen falder som potenser af afstanden mellem observationerne.

I SPSS specificeres dette som s. 55, idet man dog under Repeated Covariance Type ændrer til AR(1).

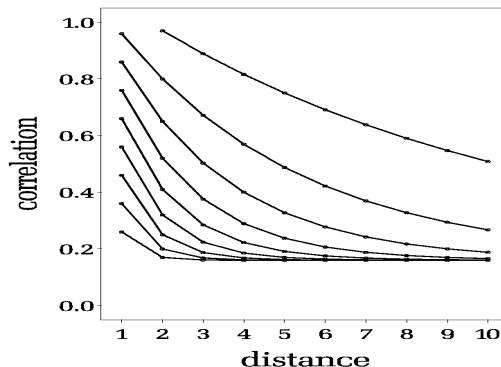
Hvis tiderne ikke er ækvidistante, svarer det til korrelationen $\text{Corr}(Y_{git_1}, Y_{git_2}) = \rho^{|t_1 - t_2|}$, se s. 111.



Autoregressiv korrelation

med overlejret tilfældigt niveau

– som funktion af afstanden mellem målingerne
for $\rho = 0.1, \dots, 0.9$



Test af interaktionen $\text{treatm}^*\text{visit}$?

– for forskellige valg af kovariansstruktur

Kovarians struktur	Teststørrelse $\sim$ fordeling	P-værdi
Uafhængighed	$0.35 \sim F(4,491)$	0.84
Compound symmetry	$5.30 \sim F(4,382)$	0.0004
Autoregressiv (evt. med CS oveni)	$2.86 \sim F(4,383)$	0.023
Ustruktureret	$2.72 \sim F(4,107)$	0.034

Altså ikke samme konklusion!

Kovariansstrukturen kan være vigtig, så hvordan vælger man?

Det vender vi tilbage til....s. 76



Tidspunkt for de 5 visits

- ▶ Der var (naturligvis) ikke præcis et halvt år mellem alle successive målinger.
- ▶ Pigerne var heller ikke præcis lige gamle ved start

Hvad gør vi så?

Hvad er den fornuftige tidsskala?

- ▶ Alder?
- ▶ Tid siden randomisering?

Vi antager, at datoen for den første måling også er dato for randomisering af den pågældende pige, så dette er tid 0.

Mere om håndtering af baseline senere



Tid siden randomisering

Denne er ikke forsøgt udregnet i SPSS, men overført fra SAS.

- ▶ Nu er tiden helt individuel for den enkelte pige
- ▶ De starter dog alle med tid 0 ved første besøg
- ▶ Der er ikke mere noget, der hedder visit1, visit2 osv., det svarer i hvert fald ikke til et bestemt tidspunkt
- ▶ Tiden er også omregnet til years, år siden randomisering

Bemærk vedr. output på næste side:

- ▶ Det faldende antal målinger over de to år
missing values/dropout
- ▶ Variationen i prøvetagningen til de enkelte *visits*



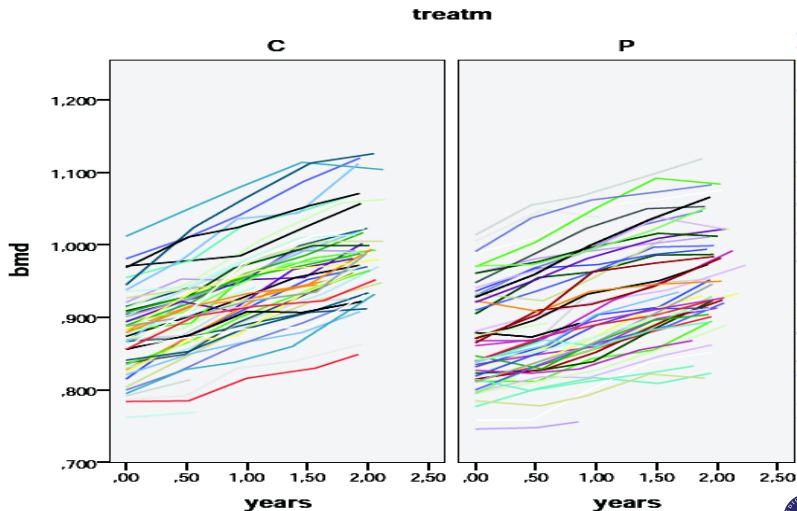
Overblik over nye tider

Descriptive Statistics

treatm	visit		N	Minimum	Maximum	Mean
C	1	bmd	55	,762	1,028	,88045
		years	55	,00	,00	,0000
		Valid N (listwise)	55			
	2	bmd	52	,769	1,051	,90327
		years	52	,42	,72	,5203
		Valid N (listwise)	52			
	3	bmd	48	,816	1,080	,93825
		years	48	,87	1,19	,9630
		Valid N (listwise)	48			
	4	bmd	46	,830	1,114	,96450
		years	46	1,34	1,78	1,4977
		Valid N (listwise)	46			
	5	bmd	44	,849	1,126	,98816
		years	44	1,84	2,15	1,9763
		Valid N (listwise)	44			
P	1	bmd	57	,746	1,014	,87007
		years	57	,00	,00	,0000
		Valid N (listwise)	57			
	2	bmd	53	,748	1,055	,88968
		years	53	,45	,60	,5182
		Valid N (listwise)	53			
	3	bmd	51	,756	1,067	,91416
		years	51	,85	1,13	,9585
		Valid N (listwise)	51			
	4	bmd	48	,809	1,096	,94165
		years	48	1,34	1,71	1,5017
		Valid N (listwise)	48			
	5	bmd	47	,816	1,119	,95823
		years	47	1,80	2,23	1,9759
		Valid N (listwise)	47			



De nye (faktiske) individuelle forløb



Kovariansstrukturer i tilfælde af uens tidspunkter

Når alle personer ikke er målt til de samme tidspunkter, er visse middelværdistrukturer og korrelationsstrukturer

ikke længere mulige

- ▶ Ustruktureret middelværdi, dvs. en parameter for hvert tidspunkt (*visit*)
- ▶ Ustruktureret kovarians/korrelation

Men man kan stadig benytte

- ▶ CS-strukturen
- ▶ random regression, kommer nu
- ▶ Erstatte AR(1)-strukturen med den tilsvarende for irregulære tidspunkter, se s. 111



Ny ide: Individuelle vækstrater?

Vi antager, at tidsudviklingen er lineær, men **ikke helt lige stejl** for alle piger, dvs. vi indfører **individuelle hældninger**:

Lad Y_{git} være den observerede BMD for den i 'te pige (i den g 'te gruppe) til tid t . Vi specificerer modellen:

$$y_{git} = a_{gi} + b_{gi}t + \varepsilon_{git}, \quad \varepsilon_{git} \sim N(0, \sigma_W^2)$$

altså en **individuel linie** for hver eneste pige, med intercept a_{gi} og hældning b_{gi}

Bemærk, at interceptet her svarer til time=0 (years=0), altså randomiseringstidspunktet (tidspunkt for baseline måling).



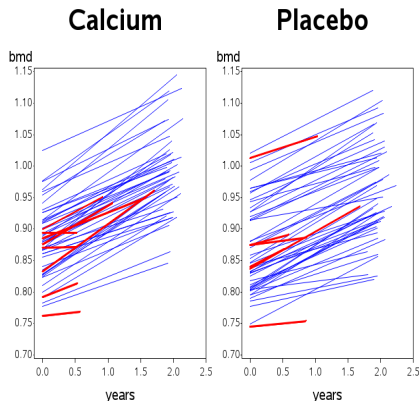
Stokastisk regression/Random regression

– en generalisering af ideen om tilfældigt niveau

Vi lader godt nok
hver pige have

- ▶ sit eget niveau a_{gi}
- ▶ sin egen hældning b_{gi}

men...



Random regression, II

... vi *binder* disse individuelle 'parametre' (a_{gi} og b_{gi}) sammen med en antagelse om Normalfordelingen (den todimensionale)

$$\begin{pmatrix} a_{gi} \\ b_{gi} \end{pmatrix} \sim N_2 \left(\begin{pmatrix} \alpha_g \\ \beta_g \end{pmatrix}, G \right)$$

hvor G er en 2×2 -matrix, der beskriver **populationsvariationen** af linierne, dvs.

- ▶ variationen mellem niveauerne
- ▶ variationen mellem hældningerne
- ▶ korrelationen mellem niveau og hældning



Random regression i praksis

Her specificeres:

- ▶ To *vilkårlige linier* under Fixed, se s. 71 (baseline-korrektion senere)
- ▶ En 2×2 -kovarians for de personspecifikke intercepter og hældninger (G) under Random, se s. 71
- ▶ De estimerede middelværdier gemmes via Save-knappen

Benyt Analyze/Mixed Models/Linear, sæt girl over i Subjects og klik Continue

Sæt nu bmd i Dependent Variable og såvel treatm som girl i Factor(s), men years i Covariate(s).

fortsættes næste side....



Random regression i praksis, fortsat

- ▶ Under Fixed markeres både `treatm` og `years`, der vælges Factorial og klikkes Add.
- ▶ Under Random sættes `girl` fra Subjects over i Combinations, og `years` sættes over i Model.

Desuden sættes flueben ved Include intercept, og, **meget vigtigt:** Covariance Type skiftes til Unstructured.

- ▶ Under Statistics, afkrydser vi Parameter estimates.



Dele af output fra random regression s. 70-71

G-matricen:

Estimates of Covariance Parameters^a

Parameter		Estimate	Std. Error
Residual		,000124	1,035965E-5
Intercept + years [subject = girl]	UN (1,1)	,004178	,000574
	UN (2,1)	,000103	,000103
	UN (2,2)	,000179	3,442941E-5

a. Dependent Variable: bmd.

Ikke megen afhængighed mellem interceper og hældninger,

$$\rho = \frac{0.000103}{\sqrt{0.004178 \cdot 0.000179}} = 0.12, \text{ men det er en tilfældighed...}$$

Information Criteria^a

-2 Restricted Log Likelihood	-2351,368
Akaike's Information Criterion (AIC)	-2343,368



Dele af output,II

Korrelationsmatricen for de 5 visits for en specifik pige, fordi de nu har individuelle tidspunkter, og dermed også individuelle korrelationer (der er ikke lige langt mellem observationerne)

Dette er hentet fra SAS:

Variansestimater (svagt stigende):

Visit 1	Visit 2	Visit 3	Visit 4	Visit 5
0.004303	0.004457	0.004677	0.005014	0.005429

```
Estimated V Correlation Matrix for girl(grp) 101 C
Row      Col1      Col2      Col3      Col4      Col5
  1      1.0000      0.9663      0.9540      0.9329      0.9072
  2      0.9663      1.0000      0.9688      0.9571      0.9396
  3      0.9540      0.9688      1.0000      0.9700      0.9598
  4      0.9329      0.9571      0.9700      1.0000      0.9727
  5      0.9072      0.9396      0.9598      0.9727      1.0000
```



Dele af output, III

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.
Intercept	,869385	,008643	109,941	100,583	,000
[treatm=C]	,011565	,012334	109,938	,938	,350
[treatm=P]	0 ^b	0	.	.	.
years	,045320	,002150	96,061	21,080	,000
[treatm=C] * years	,008887	,003074	96,589	2,891	,005
[treatm=P] * years	0 ^b	0	.	.	.

I denne model kvantificerer vi effekten af calciumtilskud (forskellen på de to hældninger) til **0.0089 (0.0031) g pr cm³ pr år**.

Ingen signifikant forskel ved baseline, P=0.35 (ses under [treatm=C]). Mere om dette fra s. 82

Dele af output, IIIa

Hvis vi erstatter kovariaten years med years2=years-2, flyttes interceptet til 2 år, dvs. slut-tidspunktet:

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.
Intercept	,960025	,009828	107,571	97,685	,000
[treatm=C]	,029339	,014026	107,653	2,092	,039
[treatm=P]	0 ^b	0	.	.	.
years2	,045320	,002150	96,174	21,080	,000
[treatm=C] * years2	,008887	,003074	96,703	2,891	,005
[treatm=P] * years2	0 ^b	0	.	.	.

Estimeret **fordel efter 2 år**: 0.0293 (0.0140) g pr cm³, P=0.039.

Sammenligning af modelfit

under antagelse om linearitet

$-2 \log L$ (og AIC) skal være lille, dvs. et stort negativt tal

Kovarians struktur	$-2 \log L$	Kov.par.	AIC	Differens i hældning	P
Uafhængighed	-1252.4	1	-1250.4	0.0105 (0.0086)	0.22
Compound Symmetry	-2253.7	2	-2249.7	0.0089 (0.0020)	< 0.0001
Autoregressiv	-2373.6	2	-2369.6	0.0094 (0.0033)	0.0037
Random Regression	-2351.4	4	-2343.4	0.0089 (0.0031)	0.0048

Random regression ser rimelig ud (AIC lille),
 måske er den autoregressive struktur dog *en anelse bedre*



Random regression = Stokastisk regression

Fordele:

- ▶ Bruger al forhåndenværende information
- ▶ Optimal procedure **hvis modellen holder**
- ▶ Let at inkludere kovariater
- ▶ Kan tage hensyn til baseline (kommer lige om lidt)

Ulemper:

- ▶ Sværere at forstå og kommunikere videre
- ▶ Biased i tilfælde af informative manglende værdier
f.eks. hvis piger med lav stigning konsekvent udgår af studiet
(den såkaldte “healthy worker” effekt)

Hvorfor ikke bare bruge individuelle linier?

(besværligere, suboptimalt, somme tider umuligt, se s. 32)



Sammenligning af de to fremgangsmåder:

Random regression eller individuelle regressioner

Forskel på hældninger C vs. P:

- ▶ Random regression: 0.0089 (0.0031), $P=0.0048$
P: 0.0453, C: 0.0542
- ▶ Individuelle regressioner: 0.0077 (0.0039), $P=0.049$
P: 0.0413, C: 0.0490

Hvorfor denne forskel?

- ▶ De korte forløb er mere usikkert bestemt, og det tages der hensyn til i Random Regression, men (selvfølgelig) ikke, når man tager gennemsnit af individuelle hældninger.
- ▶ De korte forløb har lidt lavere hældninger i C-gruppen, se s. 80.
Dette kunne være indikation af **selektivt bortfald...uha!**



Individuelle hældninger i SPSS

At udregne en hældning for hver pige er nemt at forstå og beskrive, men:

- ▶ kan kun lade sig gøre for piger med mindst to observationer
- ▶ er faktisk ret besværligt..., se s. 113.

Man ender med et nyt datasæt, indeholden (bl.a.) de enkelte hældninger, under navnet *years*, fordi dette er kovariaten.

og dette kan man nu benytte til at regne videre på, f.eks. som på næste side.

Hældninger, opdelt efter behandling og dropout-status

Her er benyttet Split File efter såvel treatm som dropout, hvor dropout=1 betyder, at pigen ikke gennemfører hele forløbet.

Descriptive Statistics

treatm	dropout		N	Mean
C	0	years	44	.0547
		Valid N (listwise)	44	
	1	years	7	.0417
		Valid N (listwise)	7	
P	0	years	47	.0459
		Valid N (listwise)	47	
	1	years	6	.0335
		Valid N (listwise)	6	

OBS-OBS-OBS:

SPSS laver fejl!! Pige nr. 305 er udeladt, fordi bmd ikke ændrer sig!? Hvad er der galt med en vandret linie??



Manglende observationer - missing values

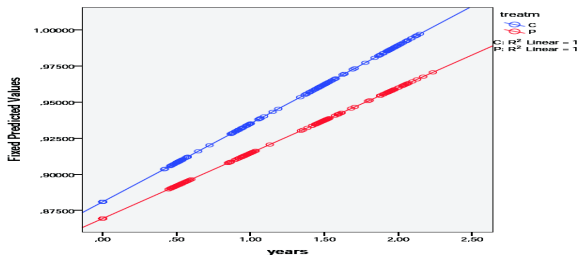
MCAR Missing completely at random
mangler alene pga tilfældigheder

MAR Missing at random
Missingness kan afhænge af kovariater (x)
og evt. også af tidligere outcome-værdier (y)

NI Non-ignorable: (Informative missing)
Missingness afhænger af
de uobserverede outcome værdier!!

Predikterede forløb fra random regression

Kode s. 114



Som vi tidligere har set, er der en vis forskel allerede fra starten (ved **baseline**), selv om denne er *insignifikant*, ($P=0.35$, se s. 74).

Bør vi justere for baseline?

Justering for baseline?

Baseline måling:

Observation foretaget inden (eller ved) behandlingsstart.

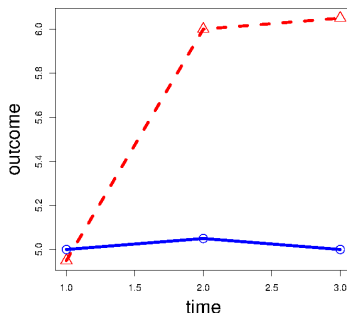
Vi diskuterede håndtering af sådanne i forbindelse med “Ancova”:

- ▶ Der skal *ikke* (nødvendigtvis) justeres i tilfælde af observationelle studier
- ▶ Justering **bør foretages**, hvis der er tale om **randomiserede undersøgelser**, fordi vi vil sammenligne to personer, som starter med at være ens, men som modtager forskellige behandlinger – og det er et *tilfælde*, hvis grupperne ikke starter på samme niveau



Hypotetisk naiv sammenligning af to grupper, I

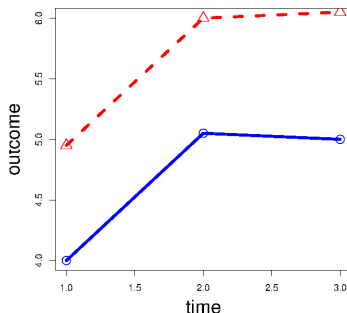
uden hensyntagen til “principielt ens” baselines:



- **Konklusion:** Interaktion mellem tid og behandling
- **Sandhed:** Konstant forskel på de to behandlinger

Hypotetisk naiv sammenligning af to grupper, II

uden hensyntagen til “principielt ens” baselines



- ▶ **Konklusion:** Konstant forskel på behandlingerne
- ▶ **Sandhed:** Ingen behandlingseffekt

Hvordan korrigeres for baseline?

- ▶ Baseline inddrages som kovariat:
 - ▶ **ikke helt godt**, fordi det svarer til at antage, at korrelationen mellem baseline og hver af de efterfølgende observationer er lige stærk
- ▶ Baseline fratrækkes:
ikke altid så godt pga *Regression to the mean*,
men fungerer fint for langsomt varierende outcome
- ▶ Middelværdierne i de to grupper sættes lig hinanden ved starttidspunktet
 - ▶ **det fungerer nemt** i “random regression”
 - ▶ Ofte kan man simpelthen omdefinere sin behandlingsvariabel, så den angiver “kontrolgruppe” for alle ved baseline.



Med baseline som kovariat (ikke helt rimeligt)

Nu er det kun de sidste 4 tidspunkter, der er outcome
Se fremgangsmåde s. 115

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.
Intercept	,019922	,024504	105,546	,813	,418
[treatm=C]	,017424	,006249	99,847	2,788	,006
[treatm=P]	0 ^b	0	.	.	.
years2	,046239	,002277	93,766	20,309	,000
[treatm=C] * years2	,007284	,003261	93,940	2,234	,028
[treatm=P] * years2	0 ^b	0	.	.	.
baseline	1,081138	,027735	101,197	38,981	,000

Bemærk kovariaten years2=years-2, der lægger interceptet hen ved 2 år.

Estimeret fordel efter 2 år: 0.0174 (0.0063) g pr cm³



Analyse af differenser (ikke helt rimeligt)

Igen er det kun de sidste 4 tidspunkter, der er outcome (se s. 116)

- ▶ Baseline fratrækkes alle efterfølgende værdier
- ▶ Baseline selv benyttes *ikke* i analysen

a. Dependent Variable: bmd_afv.

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.
Intercept	.090496	.004347	98,941	20,817	.000
[treatm=C]	.018048	.006209	99,597	2,907	.005
[treatm=P]	0 ^b	0	.	.	.
years2	.046232	.002278	93,649	20,291	.000
[treatm=C] * years2	.007330	.003263	93,817	2,246	.027
[treatm=P] * years2	0 ^b	0	.	.	.

Hvis baseline benyttes som kovariat her, fås de samme resultater som s. 87



Fælles middelværdi ved baseline

Hvis der er tale om en **lineær tidsudvikling**, som her years:

Udelad `treatm` af modelsætningen, men behold interaktionen, se s. 117

- ▶ Når years er 0 (ved baseline), har bmd middelværdi svarende til interceptet, for begge grupper.
- ▶ men der er to forskellige hældninger

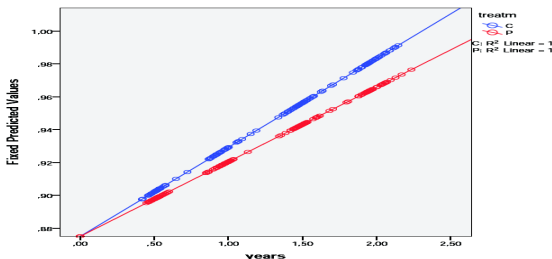
Parameter	Estimate	Std. Error	df	t	Sig.
Intercept	,875064	,006163	110,927	141,995	,000
years	,045384	,002149	96,289	21,120	,000
[<code>treatm=C</code>] * years	,008758	,003071	96,857	2,852	,005
[<code>treatm=P</code>] * years	0 ^b	0	.	.	.

Estimeret **fordel efter 2 år** fås ved simpel opgangning af forskel på hældningerne: 0.0175 (0.0062) g pr cm³



Predikterede forløb, med ens baselines

Se tilsvarende fremgangsmåde for tidligere model, s. 114



Vi ser her, at gruppernes middelværdi tvinges til at starte på samme niveau, så nu er der **justeret for baseline**

Fælles middelværdi ved baseline, II

Hvis tidsudviklingen blot er “*forskellige middelværdier*”, altså hvis tiden er en factor (såsom visit), må vi bruge et *trick*, hvor vi omdefinerer gruppen, fra `treatm` to `adj_treatm` (`adj=adjusted`), så **alle piger tilhører P-gruppen ved første måling**.

Denne definition af `adj_treat` er noget besværlig, fordi den skal foretages i to trin:

$$\text{adj_treatm} = \begin{cases} P & \text{for visit}=1 \\ \text{treatm} & \text{for visit} > 1 \end{cases}$$

Se denne definition i praksis s. 118-119, selve analysen s. 120 og output s. 92



Model fra s. 37, med ens baselines

a. Dependent Variable: bmd.

Estimates of Fixed Effects^a

Parameter	Estimate	Std. Error	df	t	Sig.
Intercept	,962412	,006850	149,778	140,492	,000
[visit=1]	-,087242	,003084	390,589	-28,292	,000
[visit=2]	-,067483	,003103	381,530	-21,750	,000
[visit=3]	-,043422	,003117	381,353	-13,931	,000
[visit=4]	-,016185	,003148	381,222	-5,142	,000
[visit=5]	0 ^b	0	.	.	.
[visit=2] * [adj_treat=C]	,007095	,004175	400,033	1,699	,090
[visit=3] * [adj_treat=C]	,013428	,004272	399,487	3,143	,002
[visit=4] * [adj_treat=C]	,012856	,004349	398,966	2,956	,003
[visit=5] * [adj_treat=C]	,019644	,004397	398,628	4,468	,000
[visit=1] * [adj_treat=P]	0 ^b	0	.	.	.
[visit=2] * [adj_treat=P]	0 ^b	0	.	.	.
[visit=3] * [adj_treat=P]	0 ^b	0	.	.	.
[visit=4] * [adj_treat=P]	0 ^b	0	.	.	.
[visit=5] * [adj_treat=P]	0 ^b	0	.	.	.

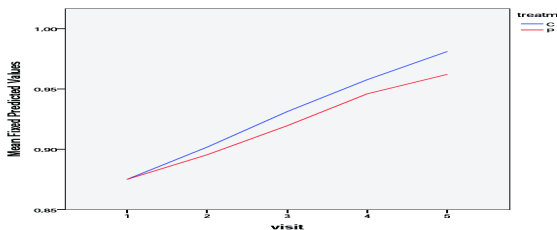
Bemærk, at der her kun er forskel på grupperne fra visit=3 og frem.



Predikterede forløb, med ens baselines

f.eks. her med CS-struktur, dvs. fra opsætningen s. 120.

Vedrørende figuren, se tilsvarende s. 114



Vi ser her, at gruppernes middelværdi tvinges til at starte på samme niveau, så nu er der **justeret for baseline**

Forskellige vurderinger af forskelle i tidsudvikling

Her i form af **Estimerede forskelle efter 2 år**:

Uden hensyntagen til baseline :

simpel model, <i>s.43</i>	0.0295	(0.0130)
random regression, <i>s.75</i>	0.0293	(0.0140)

Med hensyntagen til baseline :

5 visits, ens baseline, <i>s.93</i>	0.0196	(0.0044)
RR, ens baseline, <i>s.89</i>	0.0175	(0.0062)
baseline som kovariat, <i>s.87</i>	0.0174	(0.0063)
....		
Differenser, <i>s.88</i>	0.0181	(0.0062)
Sammenligning af rå tilvækster, <i>s.35</i>	0.0190	(0.0065)



Effekt af at korrigere for baseline

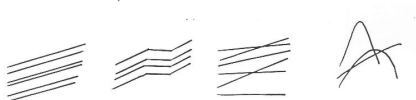
- ▶ Noget af (men ikke hele) forskellen mellem grupperne efter 2 år kan "*bortforklares*" ved, at grupperne har forskelligt udgangspunkt (baseline værdi):
- ▶ Samtidig forøges præcisionen af den estimerede forskel (standard error bliver mindre)

Forskellen bliver overbevisende signifikant

ligesom da vi så på de rå differenser

Variationskilder

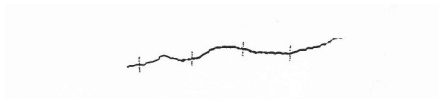
1. Tilfældige/stokastiske/random effects:



2. Seriel korrelation ('korrelationsmønster')



3. Målefejl



Flere anvendelser af Mixed Models

- ▶ Udnyttelse af alle observationer i en parret sammenligning med manglende værdier
- ▶ Cross-over studier, hvor forskellige behandlinger afprøves på samme individ, i forskellig rækkefølge
- ▶ Håndtering af flere forløb for hvert individ, f.eks. før og efter en behandling, eller ved forskellig træning
 - ▶ Her skal der tages hensyn til korrelationen mellem samtlige målinger på samme person
 - ▶ men vi må forvente højere korrelation mellem målinger foretaget i samme situation
- ▶ Adskillelse af **individuelle effekter** og **populations-effekter**

Vi kører et Mixed-kursus i nov/dec, der dog **ikke holdes i SPSS**



APPENDIX

Programbidder svarende til diverse slides:

- ▶ Plots: s. 99-100, 104, 114
- ▶ Estimation i varianskomponentmodel: s. 103, 105, 111
- ▶ Prediktion og modelkontrol: s. 107-109, 114
- ▶ Alternative kovariansstrukturer: s. 110-111
- ▶ Random regression: s. 112, 114-115, 117
- ▶ Baseline håndtering: s. 115-120



Spaghettiplot

Slide 4

Desværre måtte specialmarkering af forløbene med manglende værdier opgives.

Men selve spaghettiplottene laves ved at gå i Graph/Chart Builder/Line og vælge det til højre.

Herefter klikkes Groups/Point ID og sættes flueben i Columns panel variable, hvorved der kommer et felt i øverste venstre hjørne med navnet Panel?.

Sæt nu `treatm` i `Panel?`, `bmd` på Y-aksen, `visit` på X-aksen og `girl` i `Set color` (kræver, at `girl` bliver gjort `Nominal`)

Akserne bliver tossede, se s. 100



Ændring af akser

Slide 4, 33, 44-45

For at ændre akserne på en figur, dobbeltklikker man på den og går ind under **Y** (hvis man skal ændre på Y-aksen). Her sættes minimum og maksimum manuelt i stedet for “Auto”.

Man kan også gøre det, mens man er i gang med at danne selve figuren, nemlig ved at gå ind i **Element Properties**, hvor man klikker på den ønskede akse (her **Y-Axis1**), fjerner fluebenet i **Automatic** og skriver den ønskede værdi i stedet for.

Omstrukturering til langt datasæt

Slide 16

Dette er vanskeligt at beskrive i SPSS....

I dette tilfælde skal data først transponeres
(Data/Restructure/Transpose all data), og bagefter
benyttes igen Data/Restructure med
Restructure selected variables into cases

Følg herefter anvisningerne, herunder at sætte

- ▶ kanin som betegnelse for Case group identification
- ▶ swelling som betegnelse for Target Variable
- ▶ sted som betegnelse for Index Variable



ANOVA, sammenligning af kaniner

Slide 19-20

Ensidet variansanalyse (ANOVA) udføres i
Analyze/General Linear Model/Univariate, hvor vi sætter
swelling i Dependent Variable og rabbit i
Fixed Factor(s).

Her kan vi vælge at se parameterestimer under Options,
hvor vi afkrydser Parameter estimates.



Estimation i varianskomponentmodel

Slide 23

Benyt Analyze/Mixed Models/Linear sæt rabbit over i Subjects og klik Continue

Nu sættes swelling i Dependent Variable og rabbit i Factor(s). Under Fixed sættes ingenting, men under Random sættes rabbit over i Model, hvorefter man klikker Add/Continue.

Under Statistics, afkrydser vi Parameter estimates.



Figur af gennemsnitskurver

Slide 33

Gå ind i Graph/Chart Builder/Line og vælg det til højre.
Sæt nu bmd på Y-aksen, visit på X-aksen og treatm i
Set color.

Ovre i boksen Element Properties under Statistics vælges
Mean.

Desværre bliver akserne tossede, da de er afpassede efter selve
observationerne, men dette kan afhjælpes som beskrevet s. 100.

Mixed model for Calcium eksempel

Slide 37

Benyt Analyze/Mixed Models/Linear sæt girl over i Subjects og klik Continue

Nu sættes bmd i Dependent Variable og såvel treatm, visit og girl i Factor(s). Under Fixed markeres både treatm og visit, der vælges Factorial og klikkes Add. Under Random skal vi benytte Build nested terms og markerer først girl og trykker pil ned. Derefter klikkes Within, man markerer treatm, trykker pil ned og klikker Add/Continue.

Under Statistics, afkrydser vi Parameter estimates.



*Teknisk detalje

Slide 39

Kenwardrogers - eller Sattertwaite?

Det ser ud som om SPSS benytter en udmærket approksimation, men hvilken vides ikke pt.

Resultaterne bliver ikke altid *helt* de samme som SAS-analyserne giver, men afvigelserne er små.



Predikterede forløb

Slide 44(-45)

Følg opskriften fra s. 105, men benyt også Save-knappen, hvor man under Fixed Predicted Values sætter flueben ved Predicted Values (som får navnet FIXPRED_1).

Under Predicted Values & Residuals sættes flueben ved Predicted Values og Residuals (som får navnene hhv PRED_1 og RESID_1).

Brug herefter sædvanligt Lineplot for variablen FIXPRED_1 vs. visit, med forskellige farver for treatm.



Predikterede individuelle værdier

Slide (44-)45

Spaghettiplottene laves ved at gå i Graph/Chart Builder/Line og vælge det til højre.

Herefter klikkes Groups/Point ID og sættes flueben i Columns panel variable, hvorved der kommer et felt i øverste venstre hjørne med navnet Panel?.

Sæt nu `treatm` i Panel?, `PRED_1` på Y-aksen, `visit` på X-aksen og `girl` i Set color.

Desværre bliver akserne tossede, da de er afpassede efter selve observationerne..... se s. 100



Modelkontrol i mixed model

Slide 48-49

Efter brug af Save-knappen (se s. 107) mangler vi stadig nogle residualer, nemlig observationer minus predikterede middelværdier. Disse må vi selv danne ved hjælp af Transform/Compute, hvor vi i Target Variable sætter det nye variabelnavn, her `fixed_resid`, og i feltet Numeric Expression skriver `bmd-Fixed Predicted Values`.

Herefter benyttes Graph/Chart Builder/ til de forskellige figurer.

For at lægge en vandret linie i 0 oveni en figur, dobbeltklikker man på figuren, klikker på Options-ikonet og vælger Reference line from Equation, hvorefter man i Custom Equation skriver $y=0$ og herefter Apply.

Evt kan man vælge liniens farve under Lines.



Specifikation af CS-struktur (model fra s. 37)

Slide 53-54

Benyt Analyze/Mixed Models/Linear, sæt girl over i Subjects, visit over i Repeated og skift Repeated Covariance Type til Compound Symmetry.

Gå herefter ind i Random og sørg for, at der ikke står noget her (fjern evt. girl fra denne). Nu sættes (eller bibeholdes) bmd i Dependent Variable og såvel treatm, visit og girl i Factor(s). Under Fixed markeres både treatm og visit, der vælges Factorial og klikkes Add.

Under Statistics, afkrydser vi Parameter estimates og evt. Covariances of residuals.



Specifikation af kovariansstrukturer

Slide 55, 58, 59

Typen er her **Un**structured.

Fremgangsmåden er som på s. 110, men i Repeated skal Repeated Covariance Type sættes til Unstructured.

Typen er her **AR1**

Fremgangsmåden er som på s. 110, men i Repeated skal Repeated Covariance Type sættes til AR(1).

Man kan godt have en random effect af girl *samtidig med* en AR(1)-korrelationsstruktur.

I visse versioner af SPSS kan man også vælge strukturen SP_POWER, som er den autoregressive struktur for irregulære tidspunkter.



Random regression, med ny tid

Slide 70-71

så nulpunktet svarer til randomisering

Benyt Analyze/Mixed Models/Linear sæt girl over i Subjects og klik Continue

Nu sættes bmd i Dependent Variable og såvel treatm girl i Factor(s), men years i Covariate(s).

Under Fixed markeres både treatm og years, der vælges Factorial og klikkes Add.

Under Random sættes girl fra Subjects over i Combinations, og years sættes over i Model.

Desuden sættes flueben ved Include intercept, og (**meget vigtigt**): Covariance Type skiftes til Unstructured.

Under Statistics, afkrydser vi Parameter estimates.



Person-specifikke hældninger

Slide 79

Benyt Split File, vælg Compare groups og sæt såvel treatm, dropout og girl over i Groups Based on, så alle analyser bliver lavet for hver pige for sig, men så også de to andre variable kommer med i det nye datasæt, vi skal definere nedenfor.

Gå derefter i Analyze/Regression/Linear, udfyld med bmd som Dependent og years som Independent(s)

For at danne et datasæt (f.eks. kaldet slopes) med de individuelle hældninger, benyttes nu Save-knappen, hvor der sættes flueben i Create Coefficient Statistics og Write a new data file

I datasættes slopes kan man nu benytte Data/Select Cases, idet hældningsestimerne findes i de linier, hvor
`ROWTYPE_='EST'`



Predikterede forløb fra random regression

Slide 82

Følg opskriften fra s. 112, men benyt også Save-knappen, hvor man under Fixed Predicted Values sætter flueben ved Predicted Values (som får navnet FIXPRED_?), hvor spørgsmålstegnet er et tal, svarende til, hvor mange gange tidligere, man har gemt predikterede værdier.

Gå herefter i Graph/Chart Builder/Line, med FIXPRED_? på Y-aksen, years på X-aksen, og treatm i Colors



Med baseline som kovariat

Slide 87

Variablen `baseline` er defineret som værdien af `bmd` ved første besøg (`visit=1`).

Når denne benyttes som kovariat, bruges **kun de 4 follow-up visits** som outcome, dvs. man skal i `Data/Select cases` udvælge `If visit>1`.

Herefter inkluderes `baseline` blot under `Covariate(s)` og `Fixed` i modellen fra s. 112

Hvis man benytter kovariaten `years2=years-2` i stedet for `years`, kan forskellen ved 2 år aflæses som forskellen i intercept.



Analyse af differenser

Slide 88

Variablen `bmd_afv` defineres i `Data/transform/Compute` som `bmd-baseline`.

Når man bruger differenser som outcome, benyttes **kun de 4 follow-up visits** som outcome, dvs. man skal i `Data/Select cases` udvælge `If visit>1`.

Herefter følges opskriften fra s. 112, blot med `bmd_afv` som outcome.

Hvis man benytter kovariaten `years2=years-2` i stedet for `years`, kan forskellen ved 2 år aflæses som forskellen i intercept.



Random regression, med ens baseline

Slide 89-90

Følg fremgangsmåden som angivet s. 112, men under Fixed sættes kun `years` og `treatm*years` over i Model.

`treatm` skal ikke stå der som solo-effekt, idet grupperne skal have samme middelværdi ved starttidspunktet.



Omdefinering af gruppe ved baseline

Slide 91-93

Først skal vi definere en ny variabel `adj_treatm`, der skal være identisk med `treatm`, bortset fra ved baseline (`visit=1`), hvor alle skal tilhøre gruppen P.

Dette gøres ved at gå ind i Transform/Compute Variable, skrive `adj_treatm` i Target variable, hvorefter man under Type & Label vælger, at der skal være tale om en String variabel.

Herefter defineres der ad to omgange, som det beskrives på næste side.



Omdefinering af gruppe ved baseline, II

Slide 91-93

1. Først klikkes på If (nederst til venstre), hvorefter man under Include if case satisfies condition skriver `visit=1` og derefter Continue. Nu sættes "P" ind i feltet String expression.
2. Herefter går man igen ind i Transform/Compute Variable, skriver igen `adj_treatm` i Target variable, men under Include if case satisfies condition skriver man nu `visit>1` og derefter Continue. Nu sættes `treatm` ind i feltet String expression.

Modellen sættes op på næste side.



Modellen s. 37, baseline justeret, III

Slide 91-93

Benyt Analyze/Mixed Models/Linear sæt girl over i Subjects og klik Continue

Nu sættes bmd i Dependent Variable og såvel treatm, adj_treatm, visit og girl i Factor(s). Under Fixed sættes visit og adj_treatm*visit over i Model, *men ikke* selve adj_treatm.

Under Random skal vi benytte Build nested terms og markerer først girl og trykker pil ned. Derefter klikkes Within, man markerer treatm, trykker pil ned og klikker Add/Continue.

Under Statistics, afkrydser vi Parameter estimates.

