

Basal Statistik

Begreber. Parrede sammenligninger, i SPSS

Lene Theil Skovgaard

3. februar 2020

Indhold

- ▶ Planlægning af undersøgelse, protokol
- ▶ Grafik, Basale begreber
- ▶ Parrede sammenligninger
- ▶ Limits of agreement
- ▶ Appendix med uddybende SPSS-vejledning

Home pages:

<http://publicifsv.sund.ku.dk/~sr/BasicStatistics>

E-mail: ltsk@sund.ku.dk

*: Siden er lidt teknisk



Ide, Problemstilling

- ▶ Har “folk” et tilstrækkeligt højt niveau af vitamin D?
- ▶ Og hvis ikke, kan vi så gøre noget ved det?
 - eller i hvert fald forstå hvorfor...

Vi ser her på:

studie af kvinder fra 4 lande:

Danmark, Polen, Finland og Irland

Udvælgelse af personer?

- ▶ Hvem? Inklusionskriterier kontra repræsentativitet.
- ▶ Hvor mange? Dimensionering.
- ▶ Design?



Planlægning af undersøgelse

- ▶ Formulering af de(t) centrale spørgsmål
 - ▶ Er folk generelt oppe på det anbefalede niveau på 25 nmol/l?
 - ▶ Er der forskel på landene?
 - ▶ I givet fald, hvorfor?
- ▶ Hvilke oplysninger skal registreres?

Formodede forklarende variable = kovariater

 - ▶ spisevaner
 - ▶ sol eksponering
 - ▶ fedme
 - ▶ rygning
 - ▶ alkohol



Skriv en protokol

Dette er en vigtig del af processen!!

- ▶ Man får tænkt sig om på forhånd
- ▶ Der bliver udarbejdet information til brug for kolleger mv.
- ▶ Det tjener som “ekstra hukommelse” - man glemmer en del hvis dataindsamling eller andet trækker ud
- ▶ Det er en nødvendig del af dokumentation i forbindelse med f.eks. etisk komite, ansøgning om midler, anmeldelse af trial etc.
- ▶ I forbindelse med den statistiske analyse dokumenterer det, hvad der var den oprindelige strategi og hvad der bør betegnes som tilfældige fund

5 / 100



Eksempel på data

Obs	country	category	vitd	age	bmi	sunexp	vitdintake
1	Ireland	Woman	37.6	71.153	26.391	Sometimes in sun	5.430
2	Ireland	Woman	53.0	70.233	20.540	Sometimes in sun	9.257
3	Ireland	Woman	66.7	70.301	23.500	Sometimes in sun	30.040
4	Ireland	Woman	62.7	70.203	20.800	Avoid sun	3.005
5	Ireland	Woman	89.1	69.932	21.800	Avoid sun	4.068
6	Ireland	Woman	24.3	70.652	36.000	Prefer sun	3.443
38	Ireland	Woman	26.2	71.518	26.950	Sometimes in sun	2.158
39	Ireland	Woman	43.7	70.326	25.723	Prefer sun	3.205
40	Ireland	Woman	35.2	70.638	21.107	Sometimes in sun	7.753
41	Ireland	Woman	17.0	72.049	30.978	Prefer sun	2.906

Det oprindelige datasæt i tekstformat, samt vejledning til indlæsning fremgår af appendix bagest i disse slides (s. 81-93).

6 / 100



Datastruktur, terminologi

- ▶ Rækkerne kaldes **observationer** (typisk 1 pr. person)
- ▶ Søjlerne kaldes **variable** (en bestemt type oplysning). De kan være
 - ▶ **Kvantitative variable (Numeriske variable)**, dvs. tal, som man kan regne på
 - ▶ Vitamin D koncentration (vitd, i nmol/l)
 - ▶ Alder (age)
 - ▶ Body mass index (bmi)
 - ▶ **Kategoriske variable (Class-variable, factors)**, som kun kan antage nogle få bestemte værdier, her repræsenteret ved **tekst (string)**
 - ▶ Personens hjemland (country)
 - ▶ Personens solvaner (sunexp)

7 / 100



Anbefalet rækkefølge af aktiviteter

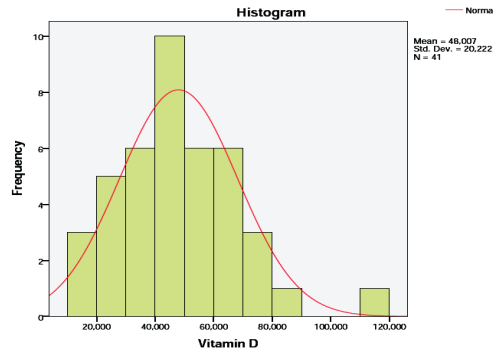
1. **Tænk** (forhåbentlig allerede på protokolstadiet)
2. **Tegn**
 - ▶ Histogram
 - ▶ Boxplot (typisk for at sammenligne grupper)
 - ▶ Scatter plot
3. **Regn**
 - ▶ Tabeller
 - ▶ Summary statistics
4. **Lav analyser**
 - ▶ Model
 - ▶ Estimation
 - ▶ Test

8 / 100



Histogram for Irske kvinder

Benyt Graphs/Chart Builder og vælg det enkleste Histogram.
Følg derefter vejledningen s. 84



God til vurdering af fordelingen

9 / 100



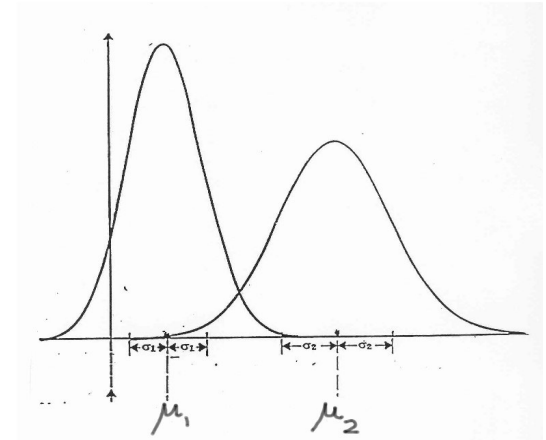
Normalfordelinger

som **ikke er nær så vigtige**, som nogle af jer sikkert tror!

Middelværdi = mean,
ofte benævnt μ , δ el.lign.

Spredning, ofte benævnt σ
(eller s , når den udregnes):

$$N(\mu, \sigma^2)$$

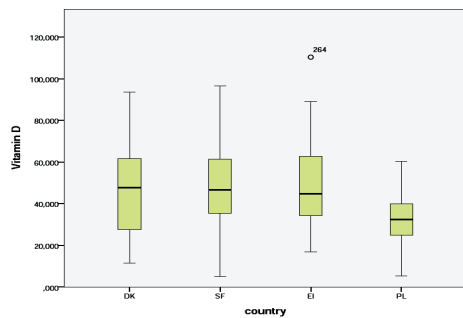


10 / 100



Box-plot for alle kvinder

Vi benytter Analyze/Descriptive Statistics/Explore, og
følger derefter vejledningen s. 85



God til
sammenligninger

- ▶ Box: 25% - 75% fraktil
- ▶ Streg: Median
- ▶ $\diamond$: Gennemsnit
- ▶ Whiskers: definitionsafhængig

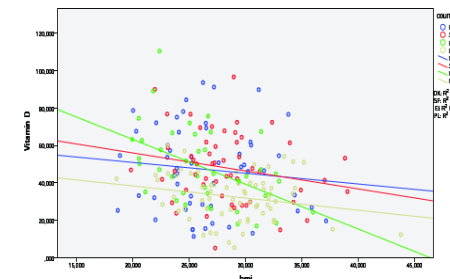
...noget med **variansanalyse**

11 / 100



Scatter-plot af Vitamin D niveau mod BMI

Benyt Graphs/Chart Builder og vælg Scatter (det, der er
nummer to fra venstre), sæt bmi over på X-aksen, vitd på
Y-aksen, og country over i Set Color. Se mere s. 86



Er der en afhængighed af BMI? Måske lineær?

... noget med **regressionsanalyse**

12 / 100



Regn: Summary statistics, I

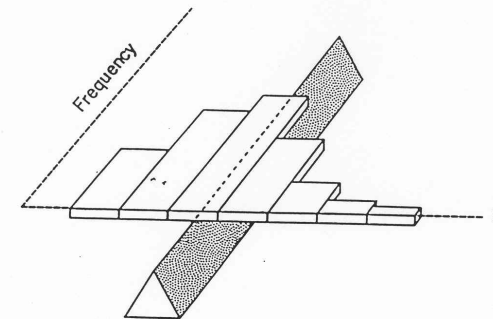
Observationer $y_1, \dots, y_n$

- Location, centrum
 - **Gennemsnit:** $\bar{y} = \frac{1}{n}(y_1 + \dots + y_n)$
 - **Median:** midterste observation, efter størrelsesorden
- I symmetriske fordelinger vil gennemsnit og median være ens (pånær tilfældigheder, naturligvis)
- I skæve fordelinger vil de *ikke* være ens:
Typisk er der hale mod de høje værdier, så *gennemsnittet er større end medianen*.

13 / 100



Gennemsnit=tyngdepunkt



- kan opfattes som ligevægtspunkt
- påvirkes kraftigt af yderlige observationer

Eksempel:

Indlæggelsestider:

5,5,5,7,10,16,106 dage

Gennemsnit: $154/7=22$ dage.

Repræsentativt for hvad?

På den anden side, hvis omkostninger er proportionale med indlæggelsestiden, så er det måske gennemsnittet, der er interessant for hospitalsledelsen.

14 / 100



Regn: Summary statistics, II

Observationer $y_1, \dots, y_n$

Variation

- Varians: $s^2 = \frac{1}{n-1} \sum (y_i - \bar{y})^2$
Spredning = Standardafvigelse = Standard Deviation =
 $\sqrt{\text{varians}} = s = \text{SD}$
- Fraktiler
Medianen er 50% fraktilen, men der er en fraktil svarende til alle procenter, se næste side

15 / 100



Fraktiler for vitamin D

Sorter data med den mindste først, tæl:

5% fraktil: 5% er mindre end dette, 95% er større

25% fraktil: 25% er mindre, 75% er større
kaldes også **nedre kvartil** eller **Q1**

50% fraktil: 50% er mindre, 50% er større
Midterste observation, kaldes også **median**

75% fraktil: 75% er mindre, 25% er større
kaldes også **øvre kvartil** eller **Q3**

$2\frac{1}{2}\%$ og $97\frac{1}{2}\%$ er vigtige,
fordi 95% af observationerne ligger imellem disse

16 / 100



Summary statistics for vitamin D

For at få opdelt udregningerne i de enkelte lande, går vi først ind i Data/Split File, vælger Compare groups og sætter country over i Groups Based on.

Herefter benyttes enten

- Analyze/Descriptive Statistics/Descriptives (s. 18):
Desværre er der ikke her mulighed for at vælge median og kvartiler, se mere s. 87
- Analyze/Descriptive Statistics/Frequencies (s. 19):
Denne giver desværre et ret uoverskueligt output, se mere s. 87

17 / 100



Summary statistics for vitamin D, II

udført med Analyze/Descriptive Statistics/Descriptives:

		Descriptive Statistics				
country		N	Minimum	Maximum	Mean	Std. Deviation
DK	Vitamin D	53	11,400	93,600	47,16604	22,782922
	Valid N (listwise)	53				
SF	Vitamin D	54	5,200	96,600	47,99074	18,724711
	Valid N (listwise)	54				
EI	Vitamin D	41	17,000	110,400	48,00732	20,222121
	Valid N (listwise)	41				
PL	Vitamin D	65	5,400	60,300	32,56154	12,464483
	Valid N (listwise)	65				

18 / 100



Summary statistics for vitamin D, III

udført med Analyze/Descriptive Statistics/Frequencies:

Statistics			
Vitamin D			
DK	N	Valid	53
		Missing	0
	Mean		47,16604
	Std. Deviation		22,782922
	Percentiles	25	27,25000
		50	47,80000
		75	64,45000
SF	N	Valid	54
		Missing	0
	Mean		47,99074
	Std. Deviation		18,724711
	Percentiles	25	34,72500
		50	46,60000
		75	61,47500

19 / 100



Summary statistics for vitamin D, IV

Bemærk:

- Median og gennemsnit er nogenlunde ens, svarende til rimelig symmetri i Boxplottene s. 11
- Dette kan *ikke* bruges til at påstå, at der er tale om en Normalfordeling
- Polen ligger lavere end de øvrige lande, undtagen måske for Q1. Bemærk, at også spredningen er lavere for Polen.

20 / 100



Traditionel fortolkning af spredningen s

Hovedparten af observationerne ligger inden for $\bar{y} \pm ca. 2 \times s$ dvs. sandsynligheden for at en tilfældig udtrukket person fra populationen har en værdi i dette interval er *stor...*

For Vitamin D blandt irske kvinder finder vi
reference område = normalområde

$$48.0 \pm 2 \times 20.2 = (7.6, 88.4)$$

Hvis data er normalfordelt, vil dette interval indeholde ca. 95% af fremtidige observationer. Hvis ikke, tja....

Note: "Ca. 2-tallet" er i virkeligheden $(1 + \frac{1}{n})t_{97.5\%}(n-1)$, og usikkerheden på grænserne (st.err.) er ca. $\sqrt{\frac{3}{n}}s \approx 5.46$

21 / 100



Er folk oppe på det anbefalede niveau på 25 nmol/l?

Vi har lige fundet normalområdet til (7.6, 88.4), hvilket fortæller os, at en del personer må forventes at have værdier under 25 - hvis vi har at gøre med en normalfordeling.

Med definitionen $lav = (vitd < 25)$ kan vi også bare lave en lille tabel (det lærer I senere, se dog s. 90):

lav				
	Frequency	Percent	Valid Percent	Cumulative Percent
Valid .00	36	87.8	87.8	87.8
1.00	5	12.2	12.2	100.0
Total	41	100.0	100.0	

som altså viser, at

mere end 12% af de irske kvinder har et for lavt niveau.

22 / 100



Normalområde / Referenceområde

Område, der omslutter de centrale 95% af observationerne:

- ▶ nedre grænse: $2\frac{1}{2}\%$ fraktil
- ▶ øvre grænse: $97\frac{1}{2}\%$ fraktil

(For irske kvinder fås 18 hhv 89.1 nmol/l, se s. 91)

Hvis fordelingen kan beskrives ved en normalfordeling $N(\mu, \sigma^2)$, kan de sande fraktiler udtrykkes som

$$2\frac{1}{2}\% \text{ fraktil: } \mu - 1.96\sigma \approx \bar{y} - 1.96s \approx \bar{y} - 2s$$

$$97\frac{1}{2}\% \text{ fraktil: } \mu + 1.96\sigma \approx \bar{y} + 1.96s \approx \bar{y} + 2s$$

og normalområdet udregnes derfor som

$$\bar{y} \pm ca. 2 \times s \approx (\bar{y} - 2 \times s, \bar{y} + 2 \times s)$$

23 / 100



Praktisk konstruktion af referenceområde

▶ Store datasæt:

Brug fraktiler

▶ Mellemstore datasæt:

Brug en rimelig fordelingsantagelse, typisk normalfordelingen, evt. efter transformation

Her er normalfordelingsantagelsen vigtig!

Mellemstort...: Med 100 observationer er usikkerheden på grænserne ca. 20%

▶ Små datasæt:

Lad være med det!!

24 / 100



Hvad er en rimelig fordelingsantagelse?

Her: passer normalfordelingen nogenlunde?

- ▶ **Gode argumenter**
 - ▶ Tegn histogram, er det symmetrisk?
 - ▶ Fraktildiagram, er det lineært? (kræver tilvænning)
- ▶ **Svagere indicier**
 - ▶ Er gennemsnit og median tæt på hinanden?
 - ▶ Er fraktilerne (f.eks. 25% og 75%) symmetriske omkring medianen?

Bemærk:

Et stort antal observationer sikrer **ikke**,
at der er tale om en normalfordeling.
– og et lille materiale **kan** sagtens være et sample fra en
Normalfordeling – vi kan bare ikke afgøre det....

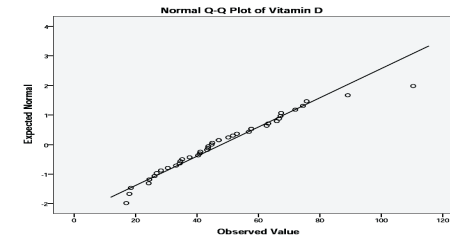
25 / 100



Fraktildiagram, “qqplot”

til check af Normalfordelingen:

Sammenlign observationerne (**X**-aksen) med **teoretiske** fraktiler, baseret på en fittet normalfordeling, her normeret (**Y**-aksen).
Bemærk, at akserne vender omvendt af dem i SAS (og R)



Denne figur kan laves på forskellig vis, se s. 92

Punkterne bør ligge nogenlunde **på en ret linie**, hvis der er tale om en normalfordeling.

26 / 100



Hvorfor normalfordelingen?

- ▶ Det er ofte en **rimelig approksimation**
 - ▶ Evt. efter transformation med logaritme, kvadratrods, invers,...
- ▶ **Central grænseværdisætning:**
 - ▶ **Sum (eller gennemsnit)** af et **stort antal** variable får en fordeling, der efterhånden kommer til at *ligne* en normalfordeling (sum af normalfordelinger er igen en normalfordeling).
- ▶ **Rimelig let at arbejde med**, fordi standard programmel er udviklet for normalfordelingen.

men som regel er antagelsen ikke specielt vigtig!
Undtagelsen er konstruktion af referenceområder

27 / 100



Eksempler på pæne normalfordelinger

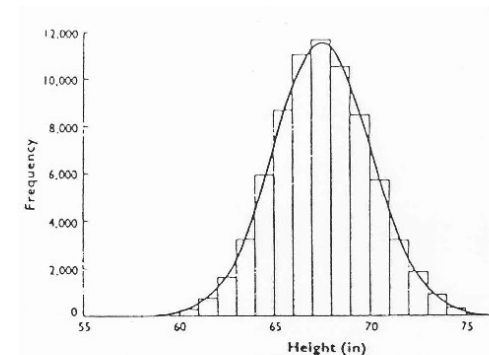


Fig. 2.10. A distribution of heights of young adult males, with an approximating normal distribution (Martin, 1949, Table 17 (Grad

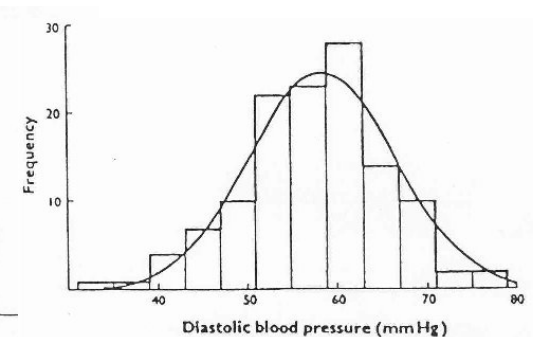


Fig. 2.11. A distribution of diastolic blood pressures of schoolboys with an approximating normal distribution (Rose, 1962, Table 1).

28 / 100



Typisk afvigelse fra normalfordelingen

som regel når der er tale om ret lave værdier, f.eks. hormonmålinger (eller immunoglobulin):

- ▶ Histogrammet er skævt, med en hale mod de høje værdier
- ▶ Gennemsnittet er en del større end medianen

Løsning: Transformer med en **logaritme**

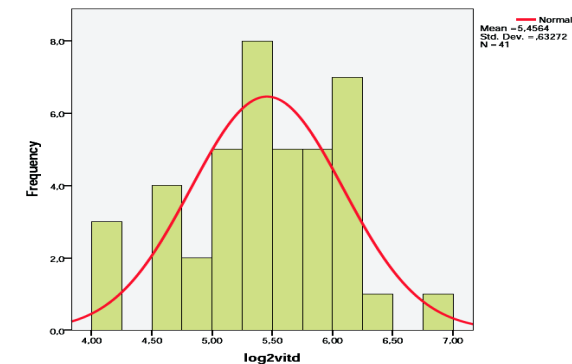
- ▶ ligegyldig hvilken: naturlig, 10-tals, 2-tals... (se s. 93)
- ▶ bare man transformerer tilbage med den samme anti-logaritme, *når man har regnet færdigt...* dvs. $\exp(\text{noget})$, 10^{noget} eller 2^{noget}

29 / 100



Histogram for logaritmerede værdier af vitamin D

her 2-tals logaritmen



igen med fittet overlejret normalfordeling

Her har vi - måske - *lidt* bedre symmetri

30 / 100



Referenceområde, baseret på logaritmer

fremgangsmåde som s. 17

Descriptive Statistics

	N	Minimum	Maximum	Mean	Std. Deviation
log2vitd	41	4,09	6,79	5,4564	,63272
Valid N (listwise)	41				

For logaritmen til Vitamin D for irske kvinder finder vi $5.456 \pm 2 \times 0.633 = (4.19, 6.72)$

Dette interval skal **tilbagetransformeres** med **anti-logaritmen**: $(2^{4.19}, 2^{6.72}) = (18.3, 105.6)$ nmol/l

Sammenlign med (7.6, 88.4) uden transformation, eller de empiriske fraktiler (18.0, 89.1)

31 / 100



Skæve fordelinger: Immunoglobulin

Summary statistics for 298 personer:

Statistics

Igm		
N	Valid	298
	Missing	0
Mean		,803
Std. Deviation		,4695
Minimum		,1
Maximum		4,5
Percentiles	25	,500
	50	,700
	75	1,000

32 / 100



Immunoglobulin summary statistics

Bemærk:

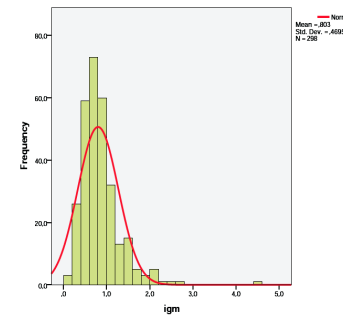
- Gennemsnittet er noget højere end medianen, så vi har nok en hale med høje værdier (jvf. s. 13)
- Maximum (og Q3) viser, at det specielt er de øverste 25% af fordelingen, der er trukket op mod de høje værdier.
- Spredningen er stor i forhold til gennemsnittet - den kan faktisk overhovedet ikke fortolkes!

33 / 100



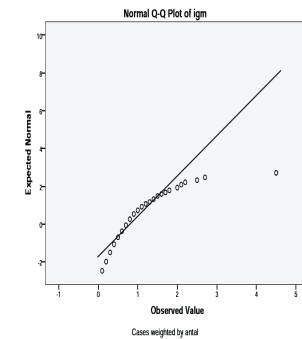
Immunoglobulin, figurer

Histogram



Tydeligt ikke-normalfordelt
(bemærk observationer langt ude til højre)

Fraktildiagram

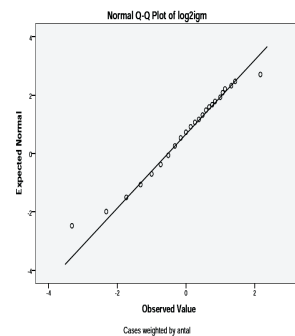
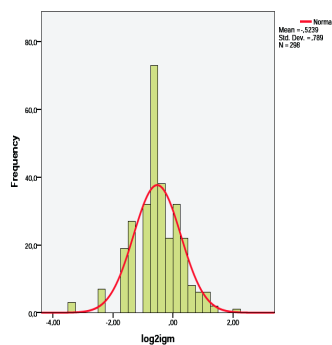


Tydeligt ikke-lineært
(omvendt hængekøjeform)

34 / 100



Immunoglobulin, log2-transformeret



Næsten lineært

Væsentlig bedre
normalfordelingstilpasning

35 / 100



Referenceområde for immunoglobulin

Urimelige værdier er i *rød kursiv*.

Ufortolkelige værdier er bare i *kursiv*

Data	gennemsnit (median)	spredning	Referenceområde
utransformeret	<i>0.803</i>	<i>0.469</i>	<i>(-0.135, 1.741)</i>
log2-transformeret	-0.524	0.789	(-2.102, 1.054)
tilbagetransformeret	(0.695)	-	(0.233, 2.076)
empiriske fraktiler	(0.700)	-	(0.2, 2.0)

Sådan foregår tilbagetransformationen:

Referenceområde for logaritmer: (-2.102, 1.054)

Tilbagetransformeret: $(2^{-2.102}, 2^{1.054}) = (0.23, 2.08)$

Lad være med at tilbagetransformere spredningerne!

36 / 100



Vigtigheden af normalfordelingen

afhænger af *formålet* med undersøgelsen

- ▶ **vigtig**
 - ▶ ved beskrivelser
 - ▶ *specielt* ved konstruktion af **referenceområder**
- ▶ **ikke så vigtig**
 - ▶ ved sammenligninger, vurdering af effekter hvor det kun er **residualerne**, der antages normalfordelte, og hvor **antallet af observationer** kan redde situationen
- ▶ **ikke på nogen måde påkrævet for kovariater!**
 - som I ikke ved så meget om endnu...

37 / 100



Parrede sammenligninger

Vi skal se på en situation, hvor vi ønsker at sammenligne to fordelinger/situationer, men hvor observationer fra den ene situation

“er **parret** med” observationer fra den anden fordeling.

Eksempler:

- ▶ Målinger på samme person før og efter en behandling
- ▶ Sammenligning af to grupper/behandlinger, hvor individerne er individuelt matchet på f.eks. køn, alder, bopæl etc.
- ▶ To målemetoder, der benyttes på samme person/dyr/blodprøve

38 / 100



Formålet med undersøgelsen

kan være flere forskellige:

- ▶ Vurdering af effekten af en behandling (således at man har målinger både før og efter behandling). Sædvanligvis vil man dog her også have en kontrolgruppe, hvis det er muligt (5. uges emne)
- ▶ Sammenligning af to behandlinger, hvor man ved hjælp af matchning eller cross-over har sørget for parrede observationer for behandlingerne
- ▶ Vurdering af, om to målemetoder/apparaturer måler det samme, eller rettere: kvantificering af, hvor stor diskrepans, der ses imellem dem

39 / 100



Sammenligning af målemetoder

To metoder til bestemmelse af slagvolumen:

- ▶ **MF**: bestemt ved Doppler ekkokardiografi
- ▶ **SV**: bestemt ved cross-sectional ekkokardiografi

Ubrugelig tabel:

person	MF	SV
1	47	43
2	66	70
3	68	72
4	69	81
5	70	60
.	.	.
.	.	.
.	.	.
17	104	94
18	105	98
19	112	108
20	120	131
21	132	131
gennemsnit	86.05	85.81
SD	20.32	21.19
SEM	4.43	4.62

Måler de to målemetoder “det samme”?

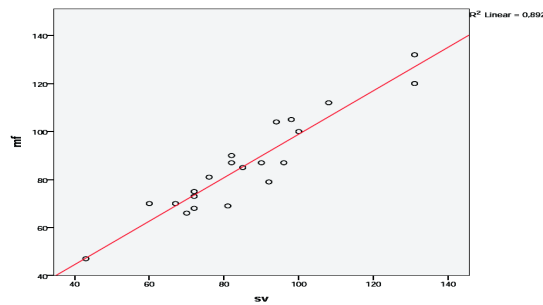
Indlæsning og transformation, se s. 94

40 / 100



Scatter plot af MF vs. SV

med indlagt **identitetslinie**



Er der en pæn lineær sammenhæng mellem de to målemetoder?

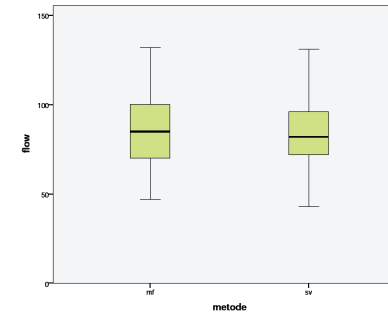
Er de måske endda rimeligt **ens**? (se s. 95)

41 / 100

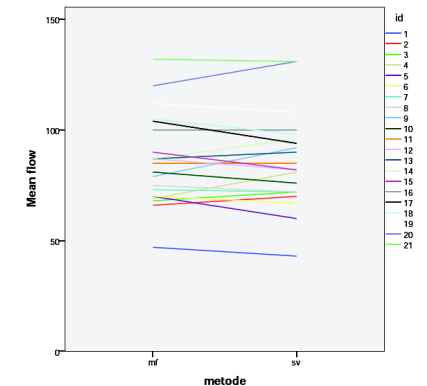


Man skal kunne se parringen!

Forkert tegning



Rigtig tegning



Se s. 96 og 97 (kræver omstrukturering af data til *langt* format)

42 / 100



Analyse af parrede data

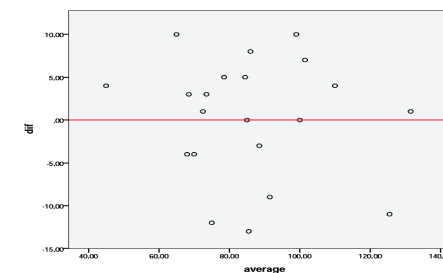
- ▶ Personen er **sin egen kontrol**
Det giver **stor styrke** til at opdage evt. forskelle.
- ▶ Se på **individuelle differenser**
– men på hvilken skala?
 - ▶ Er differensernes størrelse nogenlunde uafhængig af niveauet?
 - ▶ Eller er der snarere tale om **relative** (procentuelle) forskelle:
I så fald skal der tages differenser på en **logaritmisk** skala.
- ▶ Undersøg om differenserne har middelværdi 0:
parret T-test

43 / 100



Bland-Altman plot

Scatter plot af differenser $dif = mf - sv$ mod gennemsnit
 $average = (mf + sv) / 2$ for den enkelte person (se s. 98):



Ligger differenserne omkring 0?

Er der ca. den samme fordeling for alle gennemsnit?

44 / 100



Statistisk model for parrede data

X_i : flowmålingen MF for den i 'te person

Y_i : flowmålingen SV for den i 'te person

Differenser $D_i = X_i - Y_i$ ($i = 1, \dots, 21$)

uafhængige, normalfordelte

med **middelværdi** δ og **spredning** σ

Bemærk:

- ▶ Kun antagelser om differenser er nødvendige
 - fordi det er et **parret design**
- ▶ Intet krav om fordeling af *selve flowmålingerne*!
 - *kun* af differenserne.

45 / 100



Inferens, Statistisk analyse

Med udgangspunkt i indsamlede **data**, hvad kan vi så sige om den sandsynlighedsmekanisme (**model**, f.eks. de ukendte parametre), der har frembragt disse data?

▶ Estimation:

Når vi ser disse 21 differenser, hvad kan vi så sige om de to ukendte parametre, δ og σ ?

▶ Test:

Ser der ud til at være systematisk forskel på de to metoder, dvs. er $\delta = 0$?

▶ Prædiktion:

Hvor store forskelle kan vi forvente i praksis?

46 / 100



Gangen i en statistisk analyse

- ▶ **Modelkontrol**: Er forudsætningerne opfyldt?
Burde komme først, men kommer af praktiske grunde efter estimationen.
- ▶ **Estimation**:
Hvilke parameterværdier passer bedst med observationerne?
Og hvor sikkert er de bestemt?
- ▶ **Modelreduktion** (test af hypoteser):
Er simplere beskrivelser tilladelige?
Passer en simplere model næsten lige så godt?

47 / 100



Antagelser for den parrede sammenligning

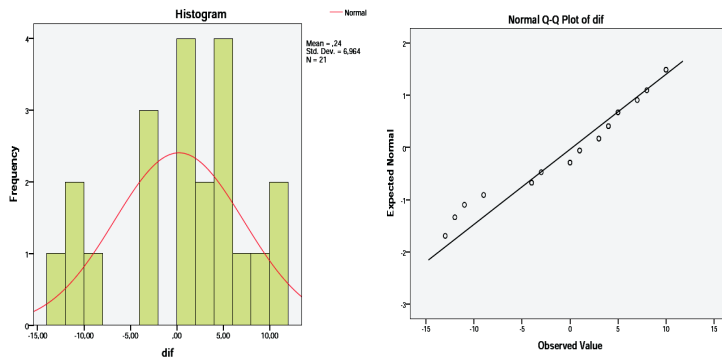
Differenserne $D_i = X_i - Y_i$, $i = 1, \dots, 21$:

- ▶ er **uafhængige**:
personerne har ikke noget med hinanden at gøre
- ▶ har **samme spredning** (varians):
vurderes ved det såkaldte Bland-Altman plot af differenser mod gennemsnit (se s. 44)
- ▶ er **normalfordelte**:
vurderes grafisk eller numerisk
 - ▶ histogram og fraktildiagram, hmm....kun 21 observationer
 - ▶ formelt test? nix...
 - ▶ **somme tider vigtig**, andre gange ikke

48 / 100



Fordeling af differenser



Passer normalfordelingen nogenlunde?

49 / 100

*Estimation

Differenserne $D_i = MF_i - SV_i \sim N(\delta, \sigma^2)$ er 21 uafhængige observationer fra en normalfordeling

Maximum likelihood princippet giver her, at

parametrene δ og σ estimeres ved henholdsvis gennemsnittet $\bar{D}$ og den tidligere omtalte spredning s (s. 15).

Vi markerer sædvanligvis estimator ved at sætte en $\hat{}$ (hat) over, altså $\hat{\delta} = \bar{D}$ (udtales *delta-hat*)

Her finder vi $\hat{\delta} = 0.238$, men

Estimator skal angives med tilhørende usikkerheder!

– gerne i form af et konfidensinterval

50 / 100

Usikkerhed på et estimat

Hvad betyder det?

F.eks. (som her) usikkerhed på et gennemsnit, som estimat for (skøn over) en ukendt middelværdi.

Vi kan tænke på gentagelser af undersøgelsen:

- ▶ Hver gang får vi et nyt estimat (for middelværdien)
- ▶ Vi kan studere fordelingen af sådanne estimater
- ▶ Spredningen i denne fordeling angiver usikkerheden. Den kaldes som regel **standard error of the estimate**

Hvis det er en middelværdi, kaldes den **standard error of the mean**

51 / 100

Altså:

Spredning på gennemsnittet kaldes

Standard error (of the mean), SEM

Hermed angives usikkerheden på gennemsnittet, og der gælder

$$SEM = \frac{SD}{\sqrt{n}}$$

SEM bliver således mindre, når n bliver større

Den bruges til at konstruere **konfidensintervaller**, som kommer nu...

52 / 100

Konfidensinterval = Sikkerhedsinterval

Interval, der “fanger” den ukendte parameter
(her middelværdien δ)
med stor (typisk 95%) sandsynlighed.

Hvor kan vi tro på at den faktiske middelværdi δ ligger?

(Approximativt) 95% konfidensinterval for middelværdi

$$\bar{D} \pm 2 \times \text{SEM}$$

Eksakt for normalfordelingen, bortset fra “ca. 2”,
 (“mere præcist 2-tal” er $t_{97.5\%}(20) = 2.086$)

Vi siger, at intervallet har **dækningsgrad 95%**,
på engelsk *coverage*

53 / 100



Konfidensinterval for δ = forskel MF-SV

95% konfidensinterval: $\bar{D} \pm \text{'ca. 2'} \times \text{SEM}$,
eller med et “mere præcist 2-tal”: $t_{97.5\%}(20) = 2.086$

I stedet for håndregning, bruger vi et T-test for hypotesen om
middelværdi 0 (**kommer om lidt**, se evt. s. 99)

Test Value = 0					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference
					Lower Upper
dif	.157	20	.877	.23810	-2.9317 3.4078

Konfidensintervallet ses at være (-2.93, 3.41)

54 / 100



Fortolkning af konfidensinterval (= sikkerhedsinterval)

Konfidensinterval for middelværdien af forskellen δ
mellem MF og SV blev estimeret til

$$(-2.93, 3.41)$$

Det betyder:

- ▶ Der kan ikke påvises nogen systematisk forskel (**bias**) mellem de to typer målinger
- ▶ Vi kan dog heller ikke afvise, at der *kan* være forskel
- ▶ En evt. bias vil med stor sikkerhed (her 95%) være mindre end. ca. 3 – 3.5 (til hver side)

55 / 100



Test af hypotese (ofte kaldet **nul**hypotese)

Kan vi nøjes med en simple model?

Kunne en eller flere parametre i en model være en kendt værdi
(**ofte 0**, deraf navnet)?

Modelreduktion: Model $\rightarrow$ (nul)hypotese (H_0).

Kan den forenklede model **tænkes at være den rigtige?**

Eksempelvis:

- ▶ Er der **systematisk forskel på de to målemetoder?** ($\delta = 0$)
- ▶ Har mænd og kvinder samme middelværdi af blodtrykket?
($\mu_1 = \mu_2$ eller $\mu_1 - \mu_2 = 0$)
- ▶ Er blodtrykket uafhængig af alderen? (hældning $\beta = 0$)
- ▶ Er der samme sandsynlighed for farveblindhed hos piger og drenge?
($p_1 = p_2$ eller $p_1 - p_2 = 0$)

56 / 100



Test af hypotese, II

Ofte ønsker vi at forkaste hypotesen, fordi vi så har fundet en effekt, f.eks. en forskel på to grupper.

Andre gange ønsker man at vise, at der ingen forskel er, **og så skal man bruge konfidensintervaller i stedet for!**

Teststørrelse: En størrelse, der måler **diskrepans mellem observation og hypotese**

- ▶ **Stor diskrepans:** Forkast hypotesen, fordi den passer dårligt sammen med data. **Men hvor stor?**
- ▶ Undersøg om teststørrelsen (diskrepansen) er **værre / mere ekstrem** end hvad der kan forventes ved tilfældighedernes spil.

57 / 100



Teststørrelsens fordeling

Teststørrelsen måler diskrepansen mellem observationerne og hypotesen.

Selv når hypotesen H_0 er *fuldstændig sand*, vil vi aldrig opnå fuldstændig overensstemmelse mellem model og observationer (f.eks. ikke nøjagtigt samme gennemsnit for MF og SV).

- ▶ Hvor store vil afvigelserne typisk være, når H_0 er sand?
- ▶ Hvilke værdier af teststørrelsen vil vi typisk få, og med hvilken hyppighed (sandsynlighed)?
- ▶ Det kaldes *fordelingen* af teststørrelsen, og den kan beregnes, så vi ved, hvad der er *normalt* og hvad der er **unormalt/ekstremt**

58 / 100



Test af “ingen bias” mellem MF og SV

dvs. test af nulhypotesen $H_0 : \delta = 0$

Vi benytter et T-test på differenserne, og dette fremkommer som brøken

$$\frac{\text{estimat} - \text{hypoteseværdi}}{\text{standard error for estimat}}$$

og det viser sig, at denne (under H_0) er T-fordelt (Student-fordelt) med et antal

frihedsgrader, som er antallet af observationer minus 1

- ▶ Lille (numerisk) t : God tilpasning
- ▶ Stor (numerisk) t : Dårlig tilpasning

59 / 100



T-test for MF vs. SV, fortsat

Her finder vi teststørrelsen:

$$t = \frac{\hat{\delta} - 0}{\text{SEM}} = \frac{0.24 - 0}{1.52} = 0.158 \sim t(20)$$

Der er 20 **frihedsgrader**, fordi der er 21 observationer og kun 1 fælles middelværdi.

Passer denne værdi (0.158) godt med en **t-fordeling med 20 frihedsgrader?**

Ja, den ligger **ret centralt** i fordelingen, og vi kan derfor ikke se noget galt med vores hypotese (se næste side).

60 / 100

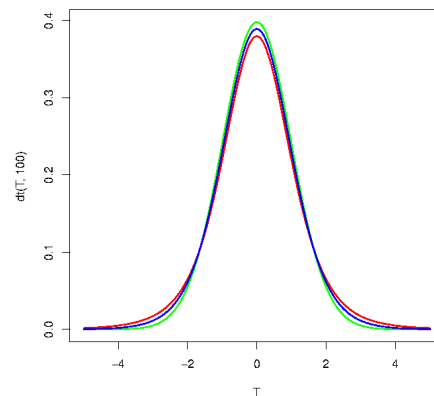


*Teknisk note

t-fordelingen (Student fordelingen)

har en parameter *df*, der kaldes
antallet af frihedsgrader
(her: 5, 10, 100).

- Mange frihedsgrader:
Fordelingen ligner
normalfordeling
- Få frihedsgrader:
Tungere haler.



61 / 100



Parret T-test i praksis

Vi vil teste ens middelværdi for MV og SV, med differenser dif
Der er flere alternativer til at gøre dette:

1. Analyze/Compare Means/Paired-Samples T Test
Her markeres mf og sv samtidigt, og føres over til Pair1
(der kan laves flere samtidige tests, hvis der er behov for det)
2. Analyze/Compare Means/One Sample T Test
hvor testet i virkeligheden blot er et test af middelværdi 0 for
differenserne, så det er blot dif, der skal anvendes.

62 / 100



Typisk output fra parret T-test

Her fra det første alternativ fra forrige side:

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	mf	86,05	21	20,321	4,434
	sv	85,81	21	21,186	4,623

		Paired Differences			95% Confidence Interval of the Difference		t	df	Sig. (2-tailed)
		Mean	Std. Deviation	Std. Error Mean	Lower	Upper			
Pair 1	mf - sv	,238	6,964	1,520	-2,932	3,408	,157	20	,877

Estimeret differens: 0.2381 (1.5196)

P-værdi: P=0.88

63 / 100



Fortolkning af P-værdi

P-værdien er sandsynlighed for “**dette eller værre**”, altså
større diskrepans end den observerede,
under nulhypotesen (dvs. hvis nulhypotesen er sand).

Hvis der kun er en ganske lille sandsynlighed for at få noget, der er
værre end det vi har, så må det være ret slemt, og vi må forkaste.

Her finder vi $P = 0.88$, altså stor sandsynlighed for at få noget, der
er værre end det vi har, så nulhypotesen ser rimelig ud: vi kan ikke
forkaste.

64 / 100



Signifikans

Hvis P-værdien er under 0.05, siger man at testet er **signifikant** på 5% niveau. Man **forkaster** hypotesen.

Signifikansniveauet α vælges sædvanligvis til 5% ($\alpha = 0.05$), men der er tale om et **arbitrært valg**.

Man bør derfor angive selve P-værdien, og allervigtigst:

Angiv estimat med konfidensinterval!

Her blev det udregnet til (-2.93, 3.41), – så vi kunne med det samme have set, at 0 *var* en rimelig værdi for middelværdien

65 / 100



Test vs. konfidensinterval

Der er **ækvivalens**, i den forstand, at:

- ▶ Hvis konfidensintervallet (sikkerhedsintervallet) indeholder 0, er testet ikke signifikant
- ▶ Hvis konfidensintervallet (sikkerhedsintervallet) *ikke* indeholder 0, er testet signifikant

Her fik vi konfidensintervallet (-2.93, 3.41), med tilhørende P-værdi $P=0.88$, så vi kan ikke forkaste hypotesen om middelværdi 0 for differenserne.

Men var det alt, hvad vi gerne ville vide?

Nej, vi vil gerne vide, hvor store forskellene typisk er....

66 / 100



Limits-of-agreement

Hvor store afvigelser vil man typisk se mellem de to metoder for **individuelle personer** (enkelt individer)

Limits of agreement er en speciel betegnelse for **normalområdet for differenser**, dvs. $\text{gennemsnit} \pm 2 \text{ spredninger}$, for differenserne

$$\bar{D} \pm 'ca. 2' \times SD = 0.24 \pm 2 \times 6.96 = (-13.68, 14.16)$$

Disse grænser er **vigtige** for at afgøre om to målemetoder kan erstatte hinanden. Det er nemlig **ikke nok**, at der ikke er nogen systematisk forskel!!

Og her er normalfordelingen vigtig!

67 / 100



Repetition: De to slags spredninger

SD: Spredningen i populationen

SEM: Standard error of the mean

$$SEM = SD(\bar{x}) = \frac{SD(x)}{\sqrt{n}}$$

(eller mere generelt blot **standard error**)

SD bruges til beskrivelser
SEM bruges til sammenligninger

68 / 100



SD til beskrivelser (“Tabel 1”)

Variable	Antal	Gennemsnit $\bar{X}$	Spredning SD
Alder	100	45	10
Immunoglobulin	100	0.80	0.47

Her tænker man:

- ▶ Patienterne er nok ca. 25-65 år
- ▶ og har immunoglobulinværdier på ca.
UPS: De kan være negative!!
(så der burde være transformeret, eller....)

Her er **normalfordelingen vigtig!**
fordi vi udtaler os om **enkeltobservationer**

69 / 100



Eksempel fra litteraturen

Malhotra, Welch, Rosenbaum & Poesz:

American Journal of Clinical Oncology • Volume 39, Number 3, June 2016

Efficacy and Toxicity of Docetaxel

TABLE 1. Demographic and Clinical Characteristics of Metastatic Prostate Cancer Patients (n=41) Treated With 3 Dosing Regimens of Docetaxel

Characteristics	Total Group	q1w (n=12)	q2w (n=14)	q3w (n=15)
Age (mean [SD]) (y)	69.7 (8.9)	72.3 (7.7)	70.2 (10.0)	67.2 (8.5)
Site of metastasis (n [%])				
Bone	36 (87.8)	11 (91.7)	12 (85.7)	13 (86.7)
Lymph node	10 (24.4)	3 (25)	4 (28.6)	3 (20)
Liver*	3 (7.3)	0	3 (21.4)	0
Lung	2 (4.9)	0	1 (7.1)	1 (6.7)
Bladder/ureter	5 (12.2)	1 (8.3)	3 (21.4)	1 (6.7)
PSA (mean [SD])				
Start	321.9 (744.4)	378.8 (458.8)	271.6 (812.1)	323.3 (894.6)
Nadir	94.6 (191.9)	141.9 (201.7)	68.1 (192.2)	81.3 (190.0)
Total dose (mean [SD]) (mg/m ²)	705.9 (479.0)	601.9 (286.9)	841.1 (649.5)	663.0 (411.7)
Total # cycles (mean [SD])*	14.8 (11.1)	19.0 (11.6)	17.1 (13.4)	9.3 (5.2)
Response (n [%])				
Yes	27 (66)	7 (58)	10 (71)	10 (67)
No	14 (34)	5 (42)	4 (29)	5 (33)
Follow-up status (n [%])				
Dead	31 (76)	11 (92)	10 (71)	10 (67)
Lost to follow up	1 (2)	0	0	1 (6)
Alive	9 (22)	1 (8)	4 (29)	4 (27)

The median age for the total group was 71.5 years with a range of 49.9 to 88.5 years.

Median PSA at start was 67.4 with a range of 0.4 to 3528.

Median PSA at nadir was 20.9 with a range of 0 to 733.9.

Median total dose of docetaxel was 540 mg/m², with a range of 100 to 1980 mg/m².

*P<0.05; remainder of comparisons were not significantly different from one another across treatment groups.

PSA indicates prostate-specific antigen; q1w, weekly regimens of docetaxel; q2w, 2 weekly regimens of docetaxel; q3w, 3 weekly regimens of docetaxel.

70 / 100



Hvad bør man så gøre?

Hvis fordelingen er **tydeligt skæv**
eller på anden måde afviger tydeligt fra normalfordelingen, bør man ikke engang angive gennemsnit og spredning, men snarere:

- ▶ fraktiler:
 - ▶ median
 - ▶ inter-quartile range, IQR: intervallet mellem 25% og 75% fraktil

For **helt små materialer** angives evt.

- ▶ median og range

..og så laver man ikke statistik, men **kasuistik**

71 / 100



SEM til sammenligninger (“Tabel 2”)

Variable	Gruppe 1 (n=35)		Gruppe 2 (n=65)	
	Gennemsnit $\bar{X}_1$	SEM ₁	Gennemsnit $\bar{X}_2$	SEM ₂
Alder	43	1.7	46	1.2
Immunoglobulin	0.63	0.08	0.89	0.06

Her tænker man:

- ▶ De to grupper ser ens ud rent aldersmæssigt
- ▶ men har måske nok forskellige niveauer af immunoglobulin

Her er **normalfordelingen ikke så vigtig**,
fordi det er **gennemsnit**, vi udtaler os om

72 / 100



Central grænseværdisætning

Fordelingen af et gennemsnit er **pænere** end fordelingen af de individuelle observationer (mere normalfordelt)

Jo flere observationer, der indgår i gennemsnittet

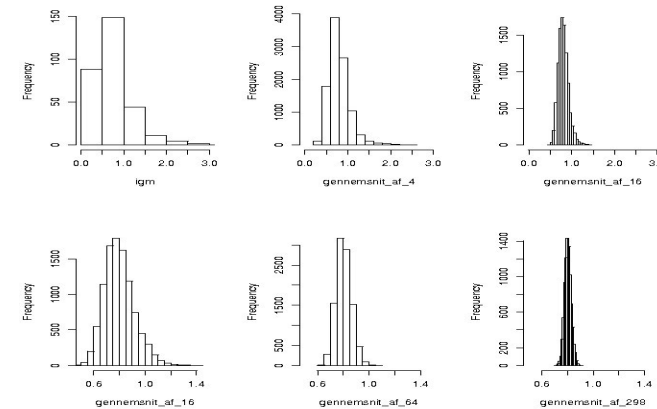
- ▶ des mere normalfordelt ser det ud
- ▶ des mindre spredning har dets fordeling (standard error of the mean, SEM), dvs. jo mere præcist fanger vi den sande middelværdi

Når man har mange observationer, gør det altså ikke så meget med fordelingsantagelsen - **så længe man ser på gennemsnit!**

73 / 100



Gennemsnit af flere og flere - immunoglobulin



Øverste linie:

Oprindelig fordeling, samt gennemsnit af 4 og 16

Nederste linie (i ny skala): gennemsnit af 16, 64 og 298

74 / 100



Hvis man **ikke** har mange observationer

- ▶ kan man **ikke** kontrollere normalfordelingsantagelsen
- ▶ og man bliver **ikke** reddet af *den centrale grænseværdisætning*

Man kan sige, at

- ▶ Muligheden for at kontrollere (forkaste) normalfordelingsantagelsen **vokser** med antallet af observationer
- ▶ Vigtigheden af normalfordelingsantagelsen **falder** med antallet af observationer

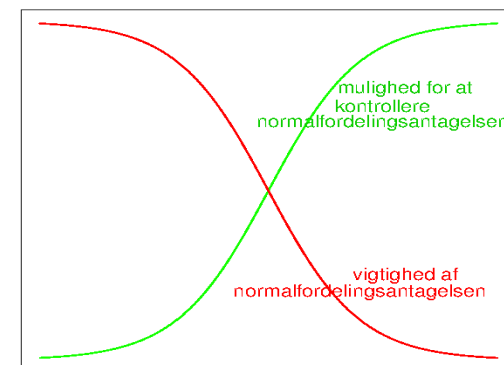
Så ved små studier kan man blive nødt til at benytte non-parametriske metoder

75 / 100



Kontrol af normalfordeling

Lidt af et dilemma:



sample size

76 / 100



Non-parametriske metoder - tests

Tests, der *ikke* bygger på en normalfordelingsantagelse

– men de er *ikke* forudsætningsfri

Ulemper

- ▶ tab af efficiens (sædvanligvis lille)
- ▶ uklar problemformulering
 - manglende model, og dermed ingen fortolkelige parametre
- ▶ **ofte ingen estimer!** – og ingen konfidensintervaller
- ▶ kan kun anvendes i simple problemstillinger
 - med mindre man har godt med computerkraft

77 / 100



Nonparametrisk one-sample test

af middelværdi 0 (parret two-sample test)

- ▶ **Sign test**, fortegnstest
 - ▶ udnytter kun observationernes fortegn, ikke deres størrelse
 - ▶ ikke særligt stærkt
 - ▶ invariant ved transformation
- ▶ **Wilcoxon signed rank test**
 - ▶ udnytter observationernes fortegn, **kombineret** med rangordenen af de numeriske værdier
 - ▶ stærkere end sign-testet
 - ▶ kræver at man kan tale om 'store' og 'små' forskelle
 - ▶ **kan påvirkes af transformation**

Men vi får hverken estimat, konfidensinterval eller limits of agreement...

78 / 100



Nonparametriske parrede tests i praksis

Brug Analyze/Nonparametric Tests/

Legacy Dialogs/2 Related Samples, se mere s. 100

Wilcoxon Signed Ranks Test

		Ranks		
		N	Mean Rank	Sum of Ranks
sv - mf	Negative Ranks	12 ^a	8,58	103,00
	Positive Ranks	7 ^b	12,43	87,00
	Ties	2 ^c		
Total		21		

a. sv < mf

b. sv > mf

c. sv = mf

Test Statistics^a

	sv - mf
Z	-,322 ^b
Asymp. Sig. (2-tailed)	,747
Exact Sig. (2-tailed)	,760
Exact Sig. (1-tailed)	,380
Point Probability	,007

a. Wilcoxon Signed Ranks Test

b. Based on positive ranks.

Disse giver kun en P-værdi, og hverken *estimat*, *konfidensinterval* eller *limits of agreement*...

Forskellige programmer benytter lidt forskellige teststørrelser! (og benytter approksimationer, som regel for $n > 25$)

79 / 100



APPENDIX

Vejledninger svarende til diverse slides:

- ▶ Indlæsning af vitamin D datasæt, s. 81-93
- ▶ Tegninger vedrørende vitamin D, s. 84-86
- ▶ Udregning af summary statistics og fraktildiagram, s. 87-92
- ▶ Indlæsning af MF-SV data, s. 94
- ▶ Tegninger vedrørende MF-SV, s. 95 - 98
- ▶ Parret T-test, s. 99
- ▶ Nonparametrisk test, s. 100

80 / 100



Det oprindelige Vitamin D datasæt

De første 5 linier (4 observationer):

```
country;category;vitd;age;bmi;sunexp;vitdintake
1;1;22.400;11.888;19.254;2;7.188
1;1;37.000;12.441;17.567;3;1.186
1;1;12.900;13.025;17.700;3;1.480
1;1;13.600;13.501;16.953;3;1.612
```

- ▶ Indlæsning ses på de næste sider
- ▶ Datasættet vitamind.txt ligger på hjemmesiden

81 / 100



Indlæsning

Slide 6

Vi indlæser filen fra nettet ved at benytte

File/Open/Internet Data, hvorefter man skriver stien

<http://publicifsv.sund.ku.dk/~lts/basal/data/VitaminD.txt>

i Web location... samt det ønskede navn på datasættet i

Dataset Name to Assign.

Derefter går man til File/Open/Data og sætter Files of Type til All Files og følger derefter instruktionerne.

I SPSS indeholder variablen country de numeriske værdier 1,2,4 og 6, medens der er defineret Value Labels (se s. 83) svarende til de mere sigende navne: DK, SF, EI, PL

82 / 100



Value labels

For at få relevante navne på værdierne for landene, kan man gå til Data/Variable View og klikke i Values under den relevante variabel (her country).

Herved fremkommer Variable Labels-box, hvor man successivt udfylder Value med de aktuelle værdier (her 1,2,4,6) og de dertil hørende labels (DK, SF, EI, PL).

Efter hvert par klikkes Add.

83 / 100



Histogram

Slide 9

Benyt Graphs/Chart Builder og vælg Histogram

(det enkleste, dobbeltklik det op i det store felt).

Sæt så vitd over på X-aksen, og tryk OK, så figuren fremkommer.

For at lægge en normalfordelingskurve oveni, dobbeltklikker man på figuren, klikker på show distribution curve-ikonet, afkrydser Normal og trykker Apply/Close.

84 / 100



Box-plot

Slide 11

Boxplottet skal vise alle landene, så vi må ind i Data/Select Cases, hvor der afkrydses i If og rettes til category=2 (vi satte den til noget andet s. 93).

Herefter benytter vi Analyze/Descriptive Statistics/Explore, hvor vi sætter vitd i Dependent List, country i Factor List samt sætter hak i Plots

85 / 100



Scatter plot med linier

Slide 12

Benyt Graphs/Chart Builder og vælg Scatter (det, der er nummer to fra venstre), og dobbeltklik det op i det store felt. Sæt bmi over på X-aksen, vitd over på Y-aksen, og country over i Set Color.

For at lægge regressionslinier oveni, dobbeltklikker man på figuren, klikker på Add Fit line at Subgroups-ikonet og afkrydser Linear.

Jeg plejer også at fjerne fluebenet i Attach Label to Line, før jeg trykker Apply.

86 / 100



Udregning af summary statistics

Slide 17

Der er (i hvert fald) 3 forskellige procedurer (under Analyze/Descriptive Statistics) til at producere summary statistics, og de har hver fordele og ulemper:

1. Frequencies: Her kan man vælge rigtigt mange forskellige summary statistics, bl.a. median og kvartiler. Opdeling på forskellige grupper (lande) kræver Data/Split File. Hvis man har at gøre med en kvantitative variabel, bør man fjerne fluebenet i Display frequency tables for ikke at producere en meget lang tabel over enkeltværdier.

87 / 100



Udregning af summary statistics, II

Slide 17

2. Descriptives: Denne giver et overskueligt output, men tillader kun få summary statistics. Her skal også bruges Data/Split File ved opdeling på grupper. Denne er anvendt på s. 17, se beskrivelse s. 89
3. Explore: Denne minder om Frequencies, men skal have grupperne i Factor?? i stedet for Data/Split File. Her kan fås en overskuelig tabel over fraktiler.

88 / 100



Udregning af summary statistics

Slide 17

For at få opdelt udregningerne i de enkelte lande, går vi først ind i Data/Split File, vælger Compare groups og sætter country over i Groups Based on.

Herefter benyttes

Analyze/Descriptive Statistics/Descriptives, hvor vi sætter vitd over i Variable(s), fjerner fluebenet i Display Frequency Tables, og under Statistics vælges Mean, Std.deviation, Min og Max

89 / 100



Frekvenstabel

Slide 22

For at lave en simpel frekvenstabel, kan vi benytte Analyze/Descriptive Statistics/Frequencies, hvor vi sætter lav over i Variable(s)

90 / 100



Udregning af specielle fraktiler

Slide 23

Nu skal vi igen kun se på Irske kvinder, så vi må tilbage i Data/Select Cases, hvor der afkrydses i If og skrives land=4 & category=2.

Herefter benyttes

Analyze/Descriptive Statistics/Frequencies, hvor vi sætter vitd over i Variable(s), fjerner fluebenet i Display Frequency Tables, og under Statistics vælges Percentiles, hvor vi skriver 2,5, klikker Add og skriver 97,5 og klikker Add/Continue

91 / 100



Fraktildiagram

Slide 26

Til fraktildiagrammet benyttes

- Analyze/Descriptive Statistics/Q-Q Plots, hvor vi sætter vitd over i Variables og der trykkes OK
- Analyze/Descriptives/Explore, hvor vi sætter vitd over i Dependent.
Klik herefter Plots og sæt flueben ved Normality plots with tests

Bemærk, at der er byttet om på X- og Y-akse i forhold til, hvad SAS gør

92 / 100



Transformation og udvælgelse

Slide 29

Der transformeres ved hjælp af Transform/Compute Variable, idet man f.eks. skriver $\log_{10}(\text{vitd})$ i Target Variable og $\text{Lg10}(\text{vitd})$ i definitionsfeltet.

Tilsvarende med $\log_2$, hvor man skriver $\log_2(\text{vitd})$ i Target Variable og $\text{Lg10}(\text{vitd})/\text{Lg10}(2)$ i definitionsfeltet.

Når man kun skal se på lrske kvinder ($\text{land}=4$ og $\text{category}=2$), benyttes Data/Select Cases, hvor der afkrydses i If og skrives $\text{land}=4 \ \& \ \text{category}=2$.

93 / 100



Datafilen vedr. MF og SV

Slide 40

Indlæsning: se dele af vejledningen s. 82

Her indlæses 21 linier og 2 kolonner (evt 3, hvis man vil have personidentifikationen med).

Definition af to ny variable:

Brug Transform/Compute Variable, idet man sætter dif i Target Variable og skriver mf+sv i definitionsfeltet.

Tilsvarende med gennemsnittet, hvor man sætter average i Target Variable og skriver $(\text{mf}+\text{sv})/2$ i definitionsfeltet.

94 / 100



Scatter plot med identitetslinie

Slides 41

Benyt Graphs/Chart Builder og vælg Scatter (det simple længst til venstre), og dobbeltklik det op i det store felt. Sæt sv over på X-aksen og mf over på Y-aksen.

For at lægge identitetslinien oveni, dobbeltklikker man på figuren, klikker på Properties-ikonet og vælger Add a reference line from Equation, hvorefter man i Custom Equation skriver $y=1*x+0$ og herefter Apply.

Evt kan man vælge liniens farve under Linear.

95 / 100



Omstrukturering til langt datasæt

Slide 42

Dette er vanskeligt at beskrive i SPSS....

Man skal benytte Data/Restructure og følge anvisningerne, herunder vælge

- ▶ person som betegnelse for Case group identification
- ▶ flow som betegnelse for Target Variable
- ▶ metode som betegnelse for Index Variable

96 / 100



Box-plot og Spaghetti-plot

Slide 42

Box plot (det **dårlige** valg):

Herefter benyttes Graphs/Chart Builder/Boxplot (det længst til venstre), hvor flow sættes på Y-aksen og metode på X-aksen.

Spaghettiplot (det **gode** valg):

Herefter benyttes Graphs/Chart Builder/Line (det opdelte, nr. 2 fra venstre), hvor flow sættes på Y-aksen, metode på X-aksen og person som Set color.

97 / 100



Bland-Altman plot

Slides 44

Benyt Graphs/Chart Builder og vælg Scatter (det simple længst til venstre), og dobbeltklik det op i det store felt. Sæt average over på X-aksen og dif over på Y-aksen.

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

Evt kan man vælge liniens farve under Linear.

98 / 100



Parret T-test

af MV vs. SV, med differenser dif

Slide 62 og 63

Benyt Analyze/Compare Means/Paired Samples T-test, marker både mf og sv samtidig og før dem over til Pair1.

Klik OK

Alternativt kan man benytte

Analyze/Compare Means/One Sample T test og sætte differensen $dif = mf - sv$ over i Test Variable(s).

99 / 100



Parret non-parametrisk test

af MV vs. SV, med differenser dif

Slide 79

Brug Analyze/Nonparametric Tests/ Legacy Dialogs/2 Related Samples, sæt mf over i Variable1, sv over i Variable2 og gå dernæst ind i Test Type og afkryds Wilcoxon samt i Exact og afkryds Exact

Jeg har også forsøgt at benytte Analyze/Nonparametric Tests/Related Samples, men **den giver noget helt uforståeligt** (P-værdi på 747....??)

100 / 100

