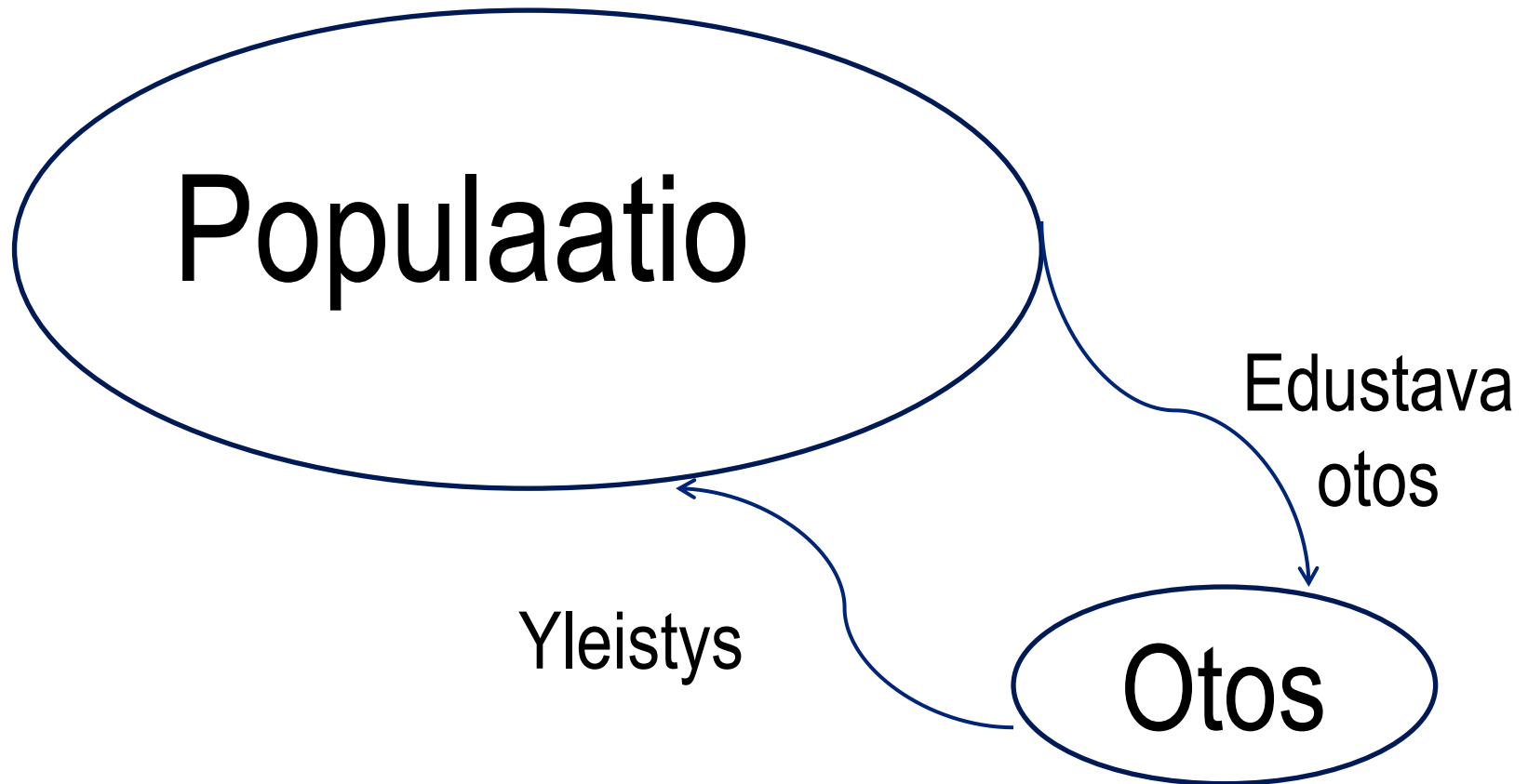


Tilastollinen päättely ja tilastolliset testit

Eliisa Löyttyniemi, 2015

Tavoite



Tausta

- Otoksen perusteella tehdään populaatiosta tai tarkasteltavaa ilmiötä koskevia päätelmiä.
- Muista, että jokainen otos antaa eri tuloksen
- Päätelmiin sisältyy aina epävarmuus
 - Ero voi jäädä havaitsematta
 - Ero, jonka löydät voi olla ei-todellinen
- Tilastotiede on päättelyä epävarmuuden vallitessa
- Päättely perustuu havaintoaineistoon

Keskiarvon luottamusväli

- Teet tutkimuksen ja lasket keskiarvolle 95% luottamusvälin
- $\approx 95\%$ todennäköisyydellä oikea arvo on tuolla välillä
- Kun koetta toistetaan 'äärettömän' monta kertaa, niin 95% tapauksissa Oikea populaatiokeskiarvo on luottamusvälillä
 - Eli 5% sattuu vain sellainen otos kohdalla, jossa näin ei ole
 - Koskaan et kerätessäsi dataa tiedä, oletko 95% vai 5% joukossa

Tutkimuksen suunnittelu

- Pyritään minimoimaan se virhe, että sanomme että eroa on, vaikka sitä ei ole (useimmiten 5%)
- Pyrimme maksimoimaan, että löydämme eron, jos ero on oikeasti olemassa (usein 80-90%)
- Tämä tehdään kunnollisella otoskokolaskuilla

- Tilastotieteellä pyritään selittämään ja ennustamaan reaalimaailman ilmiöitä
- Tilastotiede tavoittelee objektiivista ratkaisua ongelmiin
 - Antaa todennäköisyysmitan tuloksille, ei perustu siis vain tutkijan uskoon
 - Kuitenkin esim mallin valinta on subjektiivinen

Tilastojen/tilastotieteen väärinkäyttö

- Otos huono, valikoitu, ei ymmärretä tai tiedetä kokeen suunnittelusta tarpeeksi
- Kysymykset johdattelevia/huonoja
- Ei suunnitella tutkimusta kunnolla, vaan katsotaan 'mitä vaan'
- Kuvat piirretään vääristäen

- "One of the most common errors in reporting and interpreting medical research is the failure of distinguish between clinical and statistical significance. In general, **clinically important finding** is a conclusion that has implications for patient care. A **statistically significant finding**, is a conclusion based on probability" from book How to report Statistics in Medicine, Lang et al.

- "Statistical significance essentially reflects the influence of chance on the outcome; clinical importance reflects the biological value of outcome"
 - Small differences in large dataset can be statistically significant but clinically meaningless
 - Large differences in small datasets can be clinically important but not statistically significant

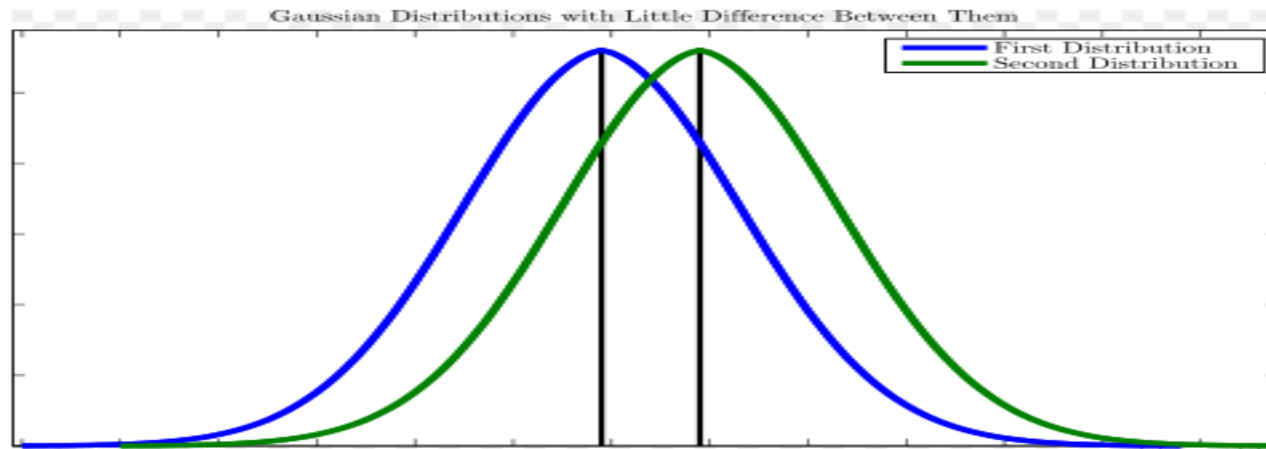
Tutkimustulos

- Nyt tutkimus on kuitenkin tehty ja haluamme tutkia tuloksia
- Kertaamme vielä tutkimuksen tavoitteet ja muotoilemme niistä statistisen hypoteesin

Tilastollinen päättely

- Hypoteesin testauksen vaiheet
 - 1) hypoteesien valinta,
 - 2) sopivan tilastollisen testin valinta,
 - 3) oletusten tarkastelu,
 - 4) testin suorittaminen ja
 - 5) lopullisen päätelmän tekeminen

- Yleensä haluamme tutkia, eroavatko kaksi ryhmää keskimäärin toisistaan, onko siis jakauma 'siirtynyt' vai ovatko ne 'päällekkäin'
- Perusoletus on jakaumien sama muoto ja sama varianssi



Hypoteesin muodostaminen

- On kiinnostuksen kohde, esimerkiksi mahdollinen ero keskiarvoissa lääkkeiden A ja B välillä
- Silloin Nollahypoteesiksi merkitään
 - $H_0: \mu_A = \mu_B$
- Ja vastakkainen vaihtoehtoinen hypoteesi on silloin
 - $H_1: \mu_A \neq \mu_B$
 - Tämä on ns kaksisuuntainen testaus (etsitään sekä pienempää, että suurempaa keskiarvoa)

Hypoteesin muodostaminen

- Joissakin erikoistapauksissa voi vaihtoehtoisen hypoteesi olla vain yksisuuntainen 'pienempi kuin' tai 'suurempi kuin'
- Jos
 - $H_1: \mu_A < \mu_B$
 - Voimme havaita eron vain tähän suuntaan, koska olemme yhdistäneet nollahypoteesiin 'yhtä kuin' ja 'suurempi kuin'

Hypoteesin muodostus

- Nollahypoteesi ja vaihtoehtoinen hypoteesi ovat siis aina toisensa poissulkevia
- Muista:
 - Tällä hypoteesin asettelulla voit osoittaa eron keskiarvoissa, et koskaan, että ne ovat samat
 - Haluat siis pyrkiä saamaan todisteita vaihtoehtoiselle hypoteesille
 - Jos haluat osoittaa 'samankaltaisuutta', niin tarvitaan erilainen hypoteesiasettelu

Hypoteesin testaus

- $H_0: \mu_A = \mu_B$
- $H_1: \mu_A \neq \mu_B$
- Testeillä testataan, onko data sopusoinnussa nollahypoteesin kanssa
- Jos datan perusteella todennäköisyys, että hypoteesi oli voimassa on pieni (alle 0.05 eli 5%), niin olemme saaneet riittävästi todisteita, että uskomme vaihtoehtoisen hypoteesiin

Päätelmä

- Silloin sanomme 'Lääkeryhmät A ja B erosivat tilastollisesti merkitsevästi toisistaan ($p=0.034$)'
- Jos emme saaneet todisteita erolle, päätelmämme ovat
- 'Lääkeryhmät eivät eronneet tilastollisesti merkitsevästi toisistaan ($p=0.35$)'.

Päätelmän kuvaus

- Pelkkä p-arvojen raportointi ei riitä, ero täytyy kuvailla esim keskiarvojen avulla 'Lääkeryhmässä A keskiarvo oli 45 (SD 1.1) ja ryhmässä B 35 (SD 1.3)'
- Tai vielä paremmin:
- 'Lääkeryhmässä keskiarvon luottamusväli oli 42-48 ja ryhmässä B 33-37.'
- Ja vielä paremmin:
- 'Lääkeryhmässä keskiarvon luottamusväli oli 42-48 ja ryhmässä B 33-37. Tämä ero on kliinisesti merkittävä'

Raportointi

- RESULTS:
- Hypoteesin **päätelmät**
- Tunnuslukujen tai luottamusvälin avulla **kuvailu** tuloksista
- DISCUSSION:
- Pohdinta tulosten **merkittävyydestä käytännössä**

P-arvo

- P-arvo on siis todennäköisyyssmitta sille, että sisältääkö data todisteita nollahypoteesin suuntaan vai pystytäänkö vaihtoehtoisen hypoteesin väite todistamaan riittävällä uskottavuudella
- Käytännössä 0.05 pidetään aina tilastollisesti merkitsevyyden rajana (poikkeuksena kuitenkin esim monivertailutapaukset)

Miten p-arvo lasketaan?

- Jokaisessa tilastollisessa testissä on erilainen kaava käytössä
- Yleisesti voidaan sanoa, että testiin vaikuttaa
 - Datan oletettu jakauma (tai että jakaumaoletusta ei haluta tehdä),
 - Havaitut datan tunnusluvut (esim keskiarvo, keskivirhe, frekvenssit, järjestys jne)
 - Otoskoko
 - Analyysissä mukana olevat muut tekijät

Testin oletukset

- Jos testin oletus on normaalijakauma, ja datassa selkeästi poikkeavia havaintoja, kaikki testistä tehtävät päätelmät ovat vääriä!
- Tekemällä jakaumaoletuksia saadaan voimakkaita testejä, mutta on siis tärkeää tarkistaa oletusten paikkansapitävyys

Epävarmuus

- Jos pidämme 0.05 tilastollisen merkitsevyyden rajana, meillä on siis aina 5% mahdollisuus, että meille on osunut sellainen otos, joka näyttäisi siltä, että ero on olemassa, vaikka ei oikeasti ole

Tilastollinen testi/malli

- Mallilla koitetaan mallintaa dataa, tutkia eri tekijöiden vaikutusta vasteeseen [ovatko esim miesten ja naisten datan jakaumat eri tasolla] ja siten selittää datassa olevaa vaihtelua
- Mikään malli ei ole täydellinen, eikä pysty kuvaamaan kaikkea vaihtelua
- Pitää pyrkiä löytämään kelpo malli, joka kuvaaisi ilmiötä riittävän hyvin
- Malli on syytä pitää mahdollisimman yksinkertaisena, pieni data ei voi selittää 100 tekijän vaikutusta
- Jokaisesta eri mallista tulee eri tulokset
- Mallin sopivuutta, ja oletuksia voi tarkastella

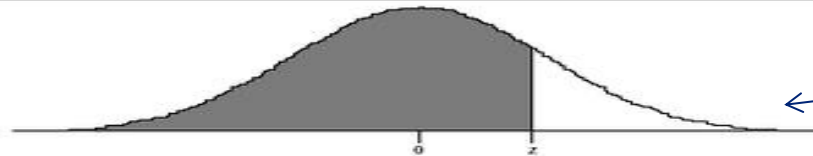
Testit

- Tilastollisissa testeissä otetaan usein huomioon havaitut tunnusluvut (esim keskiarvot ja keskihajonnat, otoskoko), mallin monimutkaisuus (kuinka paljon tekijöitä) => näiden avulla lasketaan testisuure
- Tätä testisuuretta verrataan nollahypoteesin tilanteeseen ja saadaan todennäköisyysarvio (p-arvo)

Testit

- Tätä vertailua varten tarvitaan todennäköisyysjakauma
 - kokonaispinta-ala on 1
 - Mihin kohtaa testisuure asettuu x-akselille, lasketaan jäljellejäävä pinta-ala todennäköisyys jakaumasta (=p-arvo)
 - Jakauman muotoon vaikuttaa vapausaste, joka lasketaan otoskoon avulla (kuinka monta vapaata havaintoarvoa datassa on) ja montako tekijää on mallissa

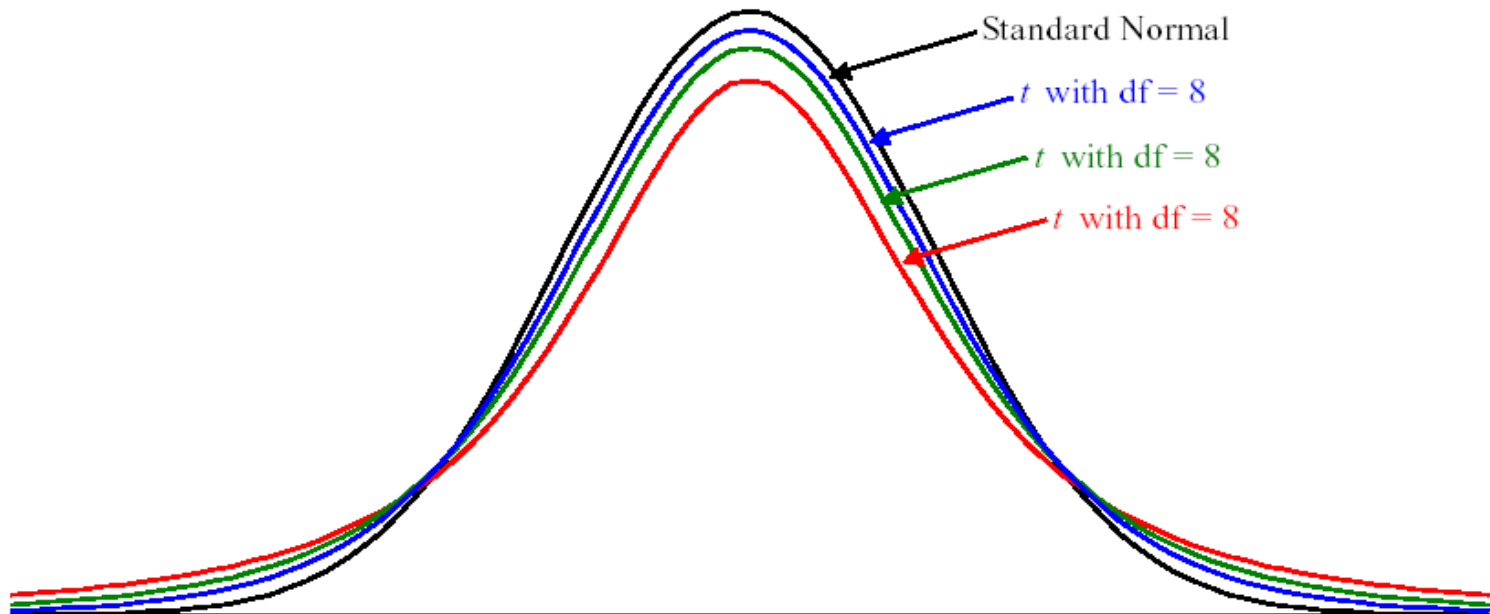
Jos testisuureen jakauma on standardisoitu normaalijakauma



Normal Deviate z	.00	.01	.02	.03	.04	.05	.06	.07	.08	.09
-4.0	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000
-3.9	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000
-3.8	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000
-3.7	.0001	.0001	.0000	.0000	.0000	.0000	.0000	.0000	.0000	.0000
-3.6	.0002	.0002	.0001	.0001	.0001	.0001	.0001	.0001	.0001	.0001
-3.5	.0002	.0002	.0002	.0002	.0002	.0002	.0002	.0002	.0002	.0002
-3.4	.0003	.0003	.0003	.0003	.0003	.0003	.0003	.0003	.0003	.0002
-3.3	.0005	.0005	.0005	.0004	.0004	.0004	.0004	.0004	.0004	.0003
-3.2	.0007	.0007	.0006	.0006	.0006	.0006	.0006	.0005	.0005	.0005
-3.1	.0010	.0009	.0009	.0009	.0008	.0008	.0008	.0008	.0007	.0007
-3.0	.0013	.0013	.0013	.0012	.0012	.0011	.0011	.0011	.0010	.0010
-2.9	.0019	.0018	.0018	.0017	.0016	.0016	.0015	.0015	.0014	.0014
-2.8	.0026	.0025	.0024	.0023	.0023	.0022	.0021	.0021	.0020	.0019
-2.7	.0035	.0034	.0033	.0032	.0031	.0030	.0029	.0028	.0027	.0026
-2.6	.0047	.0045	.0044	.0043	.0041	.0040	.0039	.0038	.0037	.0036
-2.5	.0062	.0060	.0059	.0057	.0055	.0054	.0052	.0051	.0049	.0048
-2.4	.0082	.0080	.0078	.0075	.0073	.0071	.0069	.0068	.0066	.0064
-2.3	.0107	.0104	.0102	.0099	.0096	.0094	.0091	.0089	.0087	.0084
-2.2	.0139	.0136	.0132	.0129	.0125	.0122	.0119	.0116	.0113	.0110
-2.1	.0179	.0174	.0170	.0166	.0162	.0158	.0154	.0150	.0146	.0143
-2.0	.0228	.0222	.0217	.0212	.0207	.0202	.0197	.0192	.0188	.0183
-1.9	.0287	.0281	.0274	.0268	.0262	.0256	.0250	.0244	.0239	.0233
-1.8	.0359	.0351	.0344	.0336	.0329	.0322	.0314	.0307	.0301	.0294
-1.7	.0446	.0436	.0427	.0418	.0409	.0401	.0392	.0384	.0375	.0367
-1.6	.0548	.0537	.0526	.0516	.0505	.0495	.0485	.0475	.0465	.0455
-1.5	.0668	.0655	.0643	.0630	.0618	.0606	.0594	.0582	.0571	.0559

+ 0
ence
st
fence.

Student's t -distribution



Vapausaste riippuu otoskoosta ja estimoitavien parametrien määrästä

Usein t -jakaumaa käytetään määrittelemään p -arvo, kun on pieni otos ($n < 20$)

Yhden otoksen t-testi

- One-sample t-test
- Vaste on numeerinen
- Oletuksena on normaalisti jakautunut data
- Meillä on joku arvo μ , jota vasten haluamme testin tehdä (muista sokeripaketti esimerkki), olkoon se $\mu=100$
- Nollahypoteesimme siis on
 - $H_0: \mu = 100$
- Ja vastakkainen, vaihtoehtoinen hypoteesi on silloin
 - $H_1: \mu \neq 100$

Testisuure

- Saadaan testisuure

$$t = \frac{\bar{x} - \mu}{SD/\sqrt{n}}$$

•

- $\bar{x}$ on vastemuuttujan keskiarvo, μ referenssiarvo, johon keskiarvoa verrataan,
 - SD vastemuuttujan keskihajonta ja n havaintojen lukumäärä.
 - Testisuure noudattaa Studentin t-jakaumaa vapausastein $df=n-1$ (df =degrees of freedom)
 - Studentin t-jakaumassa varianssi on suurempi kuin normaalijakaumassa. Kun havaintojen ja siten myös vapausasteiden määrä kasvaa, jakauma lähenee normaalijakaumaa
- (huom nyt puhutaan testisuureen jakaumasta, ei datan!)

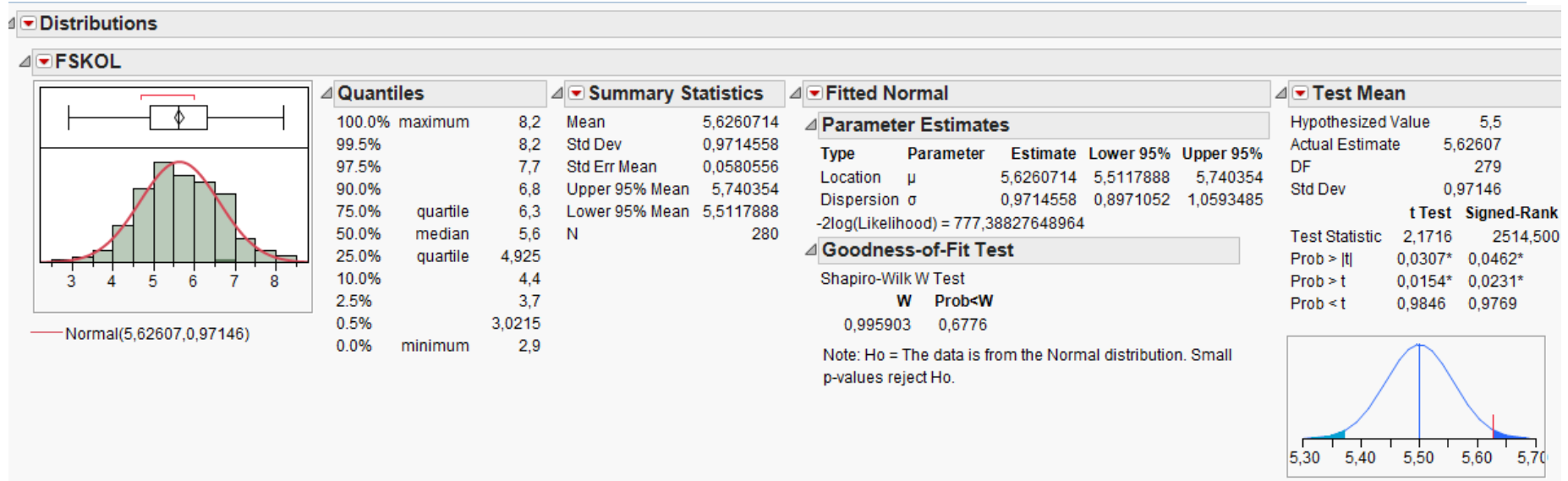
Tulos

- Jos testisuureen arvo riittävän iso tai riittävän pieni (ollaan todennäköisyysjakaumassa jommassakummassa hännässä), niin silloin havaittu keskiarvo on poikennut referenssikeskiarvosta (suhteutettuna hajontaan ja otoskokoon) niin paljon, ettei voida enää pitää todennäköisenä nollahypoteesin väittämää, joten päättelemme keskiarvon poikkeavan 100:sta (=ref arvo).

Referenssikeskiarvo?

- Joskus meillä on yhden ryhmän tilanteessa referenssiarvo, mihin haluamme verrata
- Usein vaste voi myös olla muutos (hemoglobiinin muutos alkutilanteen ja viikon välillä)
- Tällöin lasketaan jokaiselle henkilölle muutosmuuttujan arvo
- Tällöin luonnollinen testaus on, että onko tapahtunut muutosta, onko muutoksen keskiarvo=0

- Esimerkki. Halutaan Esim-datassa testat, eroaako FSKOL:n:n keskiarvo 5.5:sta?



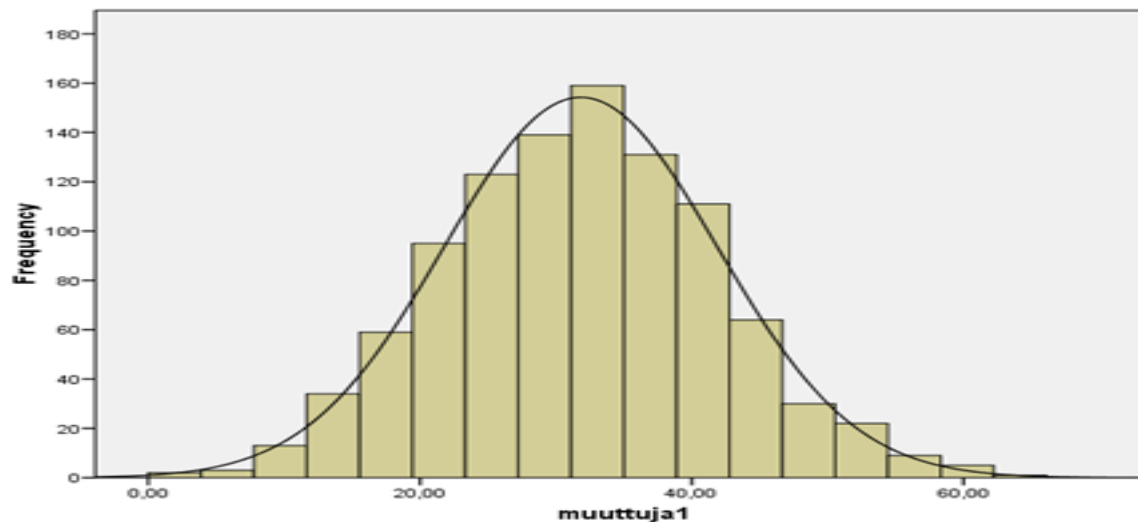
- JMP:
 - Analyze-distributions
 - Test mean – syötä referenssikeskiarvo

Normaalijakaumaoletus

- Yhden otoksen t-testi on esimerkki **parametrisesta** testissä, jossa oletetaan normaalijakauma vastemuuttujalle
- Tällöin keskiarvo ja hajonta ovat tehokkaimpia tunnuslukuja
-> tehokas testi
- Oletuksen paikkansapitävyys pitää tarkistaa

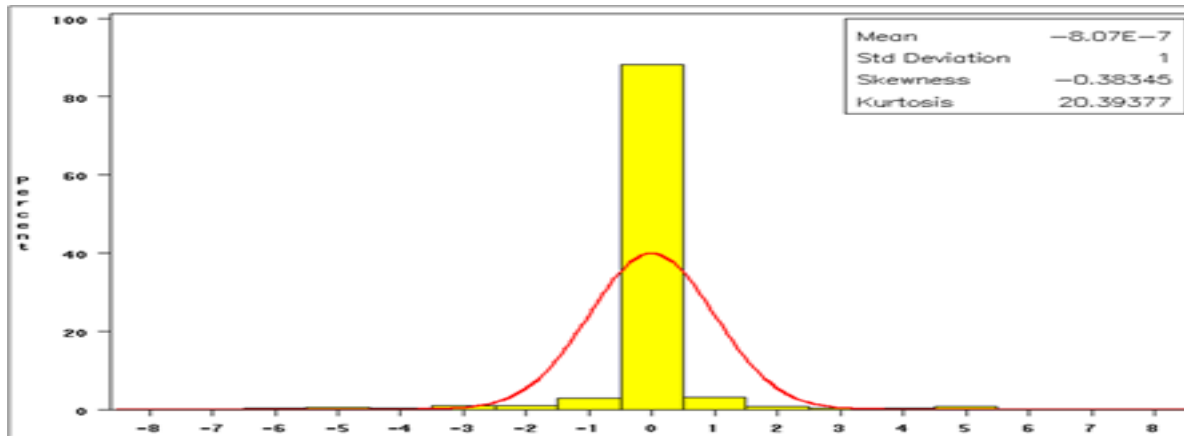
Normaalijakaumaoletus

- Alustava tarkistus histogrammi

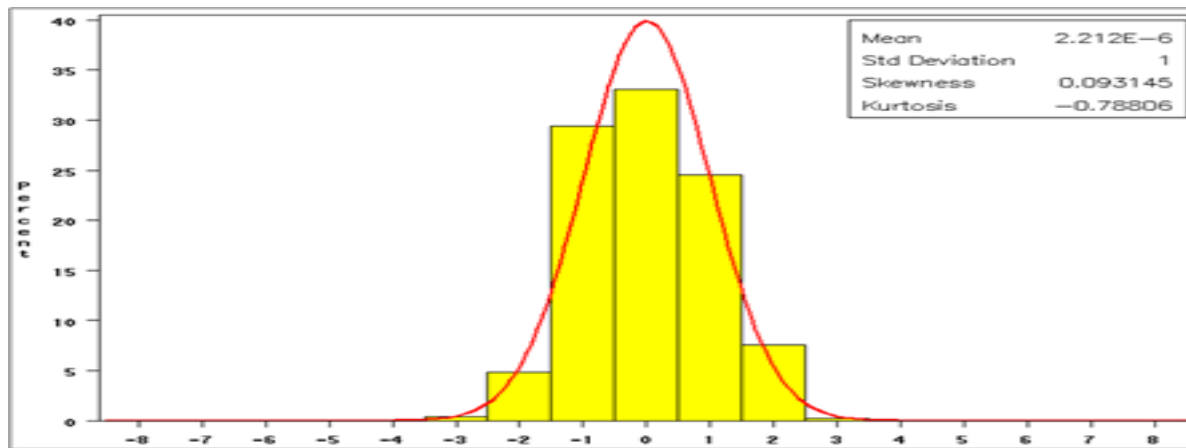


- Monessa ohjelmassa pystyy kuvaan piirtämään normaalijakauman
- Voi tarkistaa huipukkuuden ja vinouden (jos norm jak = 0)

Huipukkuus

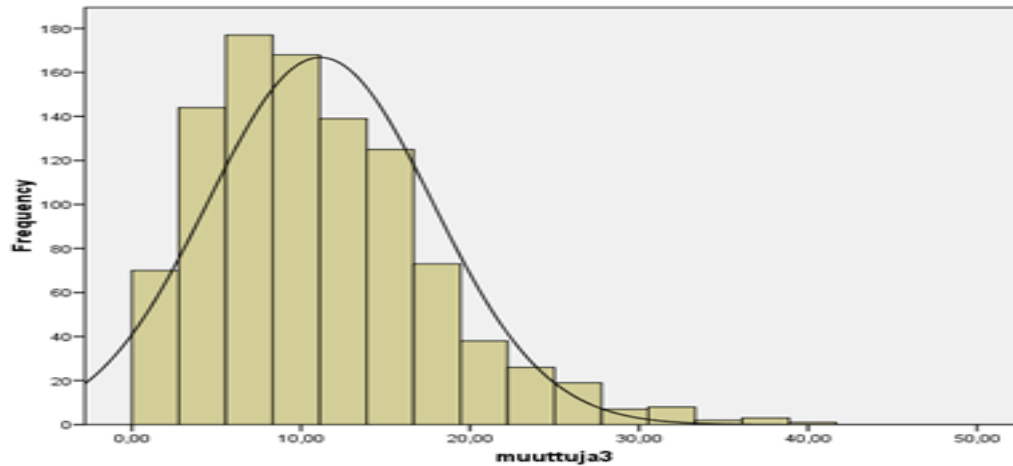


- Voimakkaasti positiivinen huipukkuuskerroin, havaintoja paljon enemmän keskiarvon lähellä kuin normaalijakaumassa



- Hieman negatiivinen huipukkuuskerroin, vähemmän havaintoja keskiarvon lähellä kuin normaalijakaumassa

Vinous



Muuttujan jakauma on oikealle eli
positiivisesti vino

$Vinous > 0$

Keskiarvo $>$ mediaani

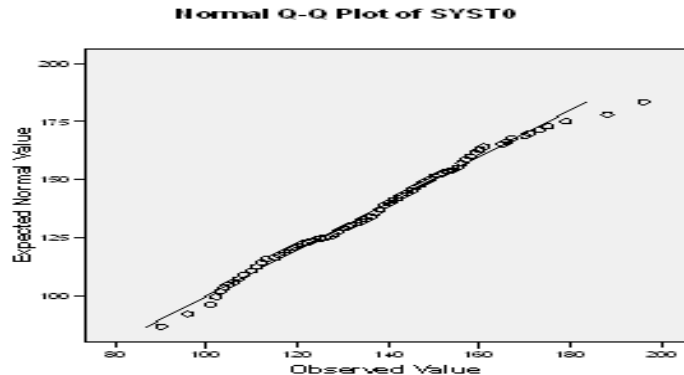
Normaalijakaumaoletus

- Histogrammissa pylvään korkeus kuvaa siinä luokassa olevien havaintojen määrää
- Histogrammi on hyvin herkkä luokkien valinnalle
- Ohjelmistoissa myös testi, onko data normaalisti jakautunut
 - Shapiro-Wilks ($n < 50$) [JMP käyttää rajan $n = 2000$]
 - Kolmogorov-Smirnov ($n > 50$)
 - Huom. Usein pienelle datalle testi ei ole voimakas -> poikkeamaa normaalijakaumasta ei löydetä ja vastaavasti isolle datalle pienetkin poikkeamat löytyy
 - => käytä järkeä!

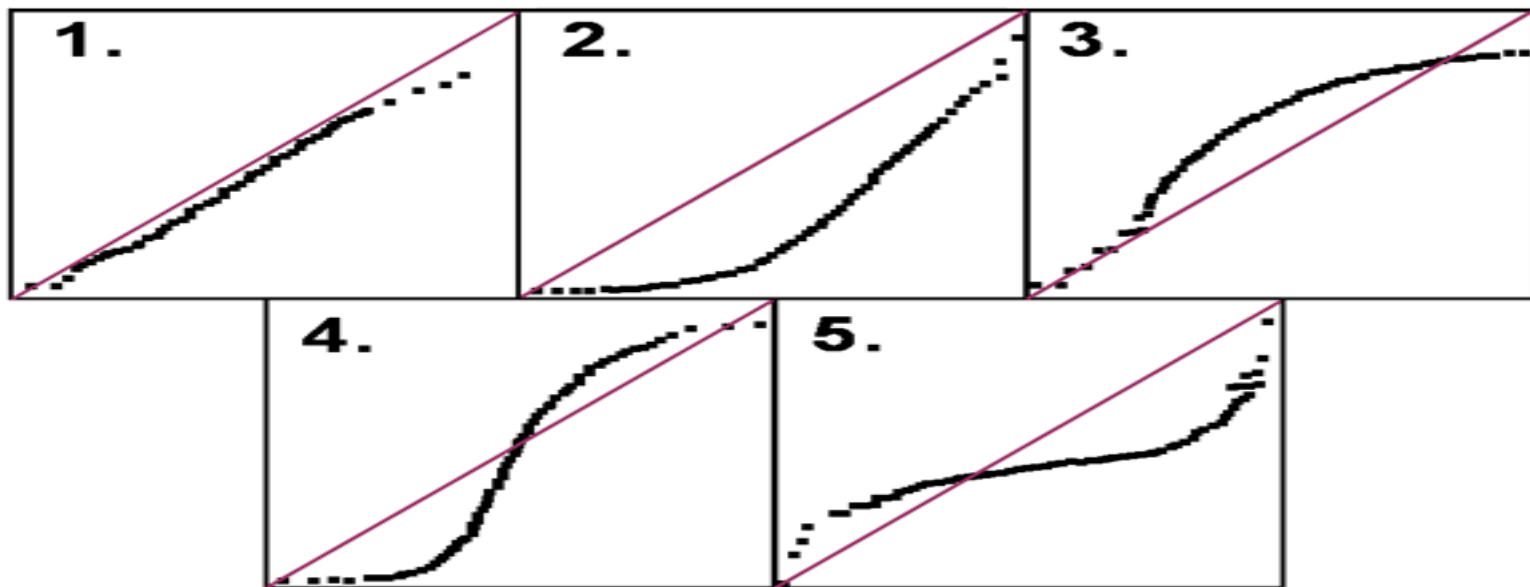
HUOM

- Hypoteesin asettelu normaaliajakaumatestissä
- H_0 : data on normaalisti jakautunut
- H_v : Data ei ole normaalisti jakautunut
- Nyt tulokseksi toivomme, että $p > 0.05$, että voisimme tehdä normaalijakaumaoletukseen perustuvan testin
- [aina kun vertailemme keskiarvoja tms toivomme pientä p-arvoa ($p < 0.05$)]

Kvantiilikuvio (Q-Q-plot)



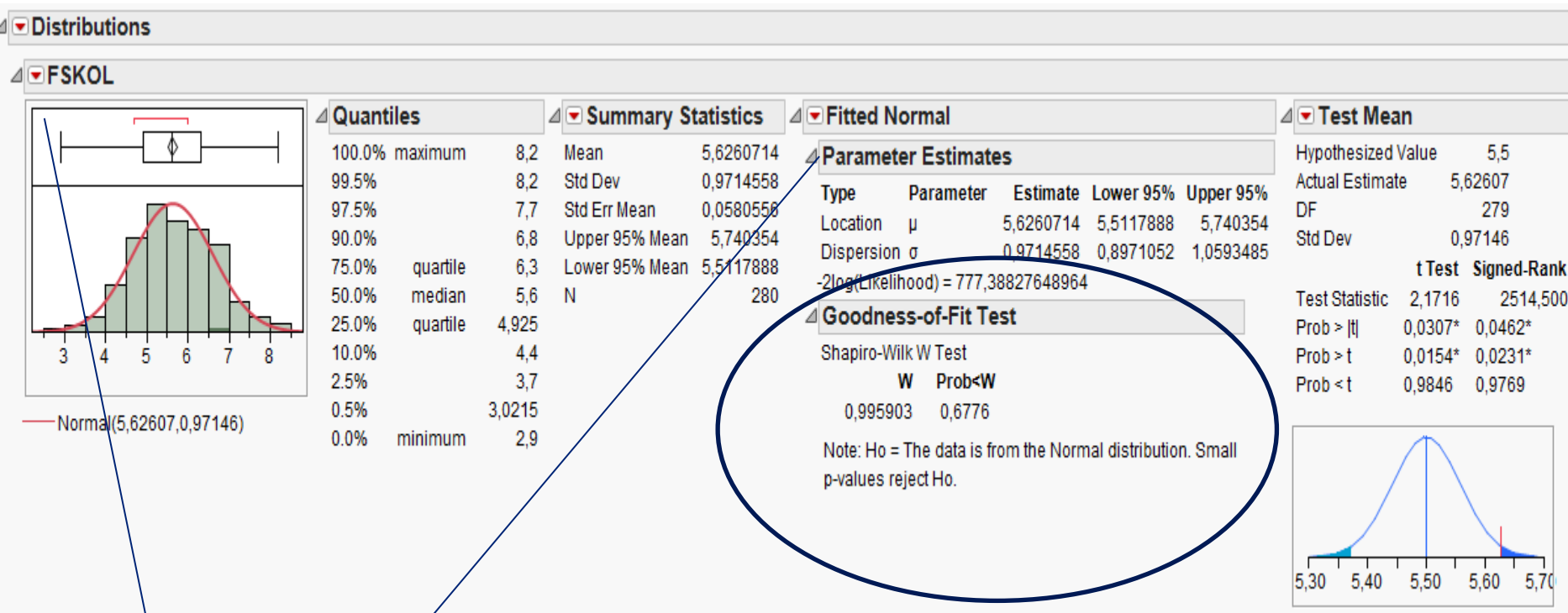
- Vaaka-akselilla (x-akselilla) on aineiston ***havaitut*** kvantiilien arvot ja pysty-akselilla (y-akselilla) normaalijakauman mukainen ***odotettu*** kvantiilin arvo (SPSS)
 - Kvantiilit ovat jakaumasta säännöllisin välein valittuja prosenttipisteitä
- Jos jakauma olisi täydellisesti normaalisti jakautunut, niin pisteet sijoittuisivat viivan päälle



Useimmiten: Pystyakselilla havaitut arvot ja vaakakselilla normaalijakauman arvot (SAS):

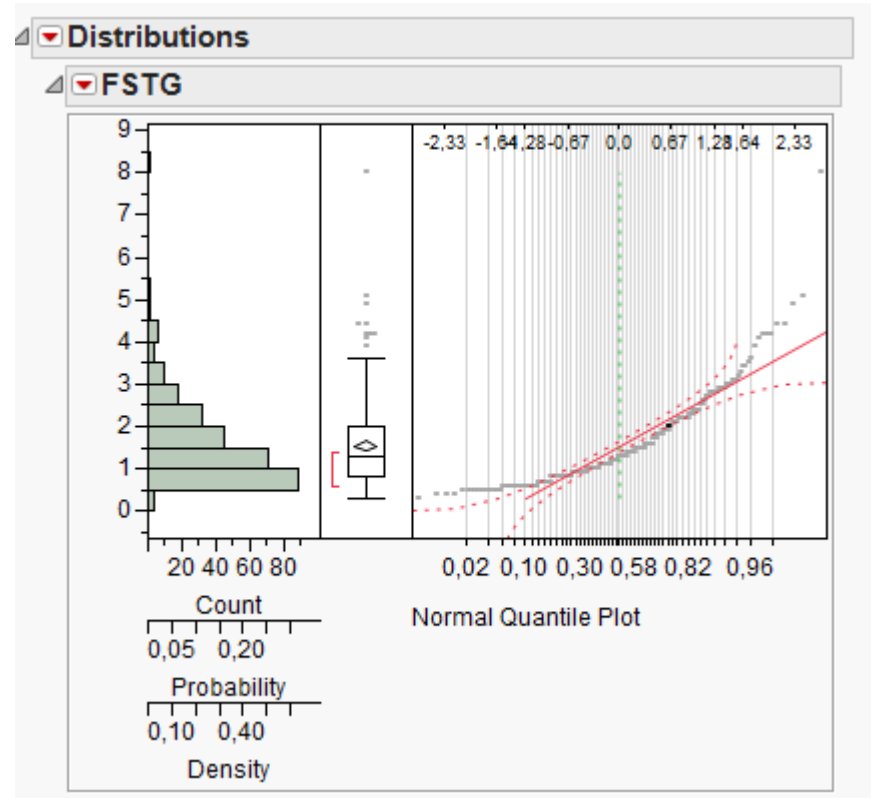
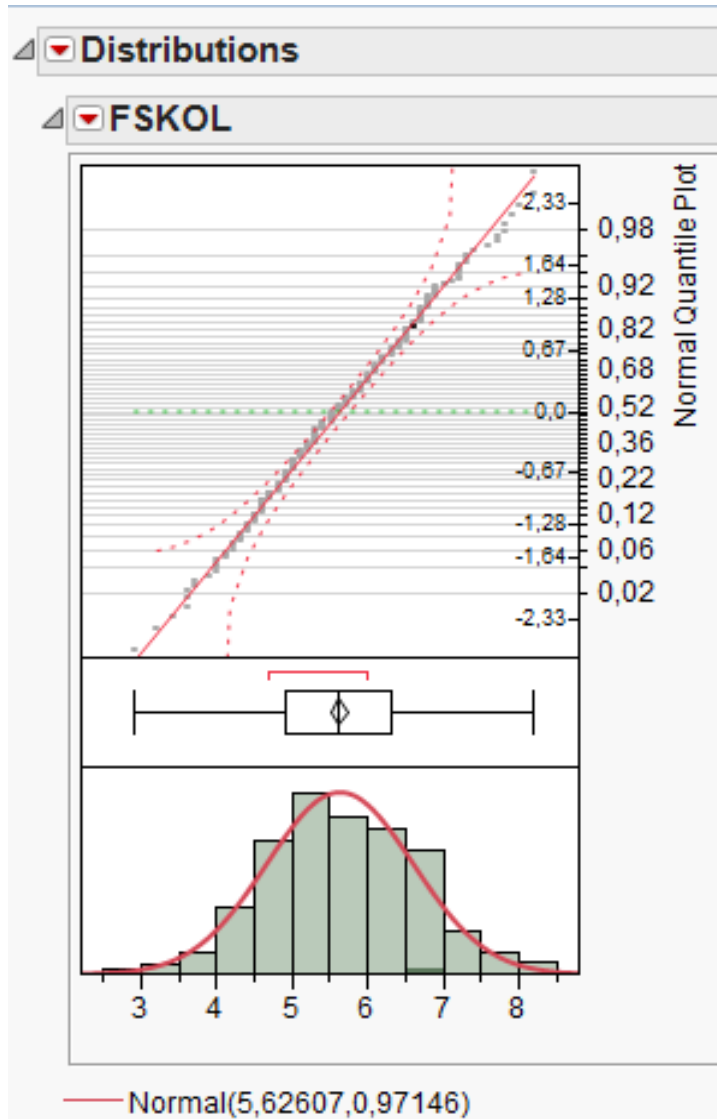
1. Normaalijakauma
2. Oikealle vino jakauma
3. Vasemmalle vino jakauma
4. Huipukas jakauma (positiivinen huipukkuuskerroin)
5. Lyhythäntäinen jakauma (negatiivinen huipukkuuskerroin) [liian tasainen vrt norm.jak]

Esimerkki



- JMP: Analyze-distribution
 - Continuous fit-Normal
 - Fitted normal-Goodness of fit

Normal Quantile plot-JMP



HUOM. JMP:ssa kuva täytyy olla näin päin, mikäli haluaa tulkita kuten yllä

Parittainen t-testi

- Paired t-test
- Laske '1 viikon arvo'-'lähtöarvo' kaikille uuteen sarakkeeseen/muuttujaan
- Vaste on numeerinen Muutos-muuttuja $\bar{d}$
- Muutos tapahtuu henkilön sisällä, kaksi havaintoa per henkilö!! Riippuvia havaintoja!!!
- Oletetaan taas normaalijakauma
- $H_0: \bar{d} = 0$
- $H_v: \bar{d} \neq 0$

Parittainen t-testi

- Nyt testisuure

$$t = \frac{\bar{d}}{SD_d / \sqrt{n}}$$

- Missä $\bar{d}$ on muutoksen keskiarvo ja SD_d muutoksen keskihajonta
- Ja testisuureesta saadaan taas p-arvo
- Muutosta kuvaamaan voidaan laskea muutoksen 95% luottamusväli

$$\bar{d} \pm 1.96 * SD_d / \sqrt{n} \text{ (kun } n > 20)$$

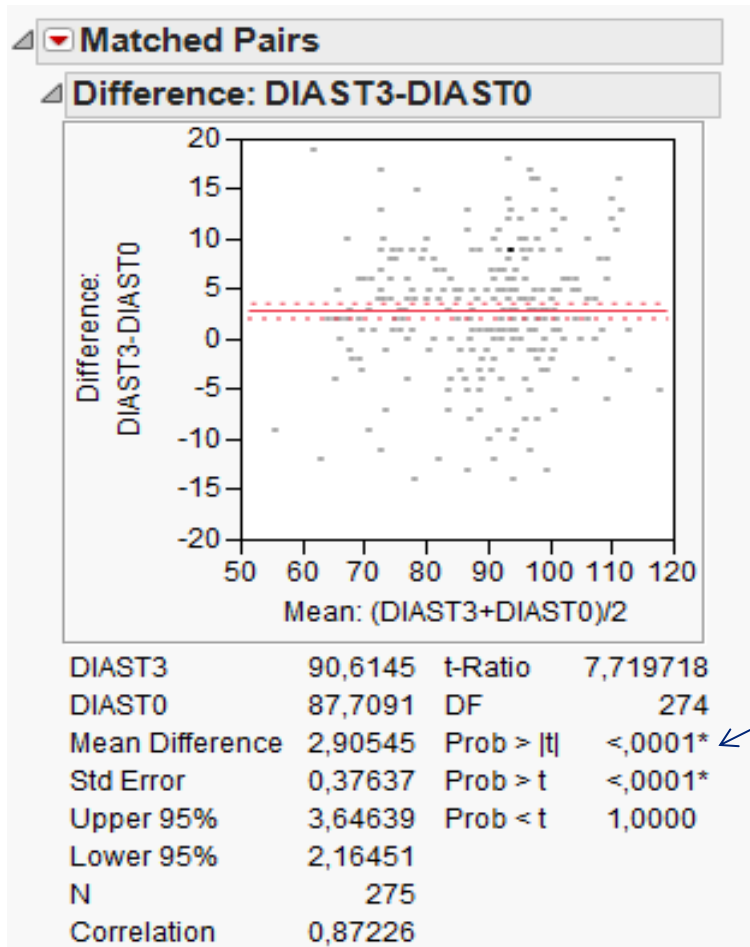
- **Esimerkki:** Tarkastellaan, onko systolisen verenpaineen arvoissa tapahtunut muutosta 9 kuukauden seurannan aikana. Testataan parittaisella t-testillä, onko alkutilanteen ja 9 kuukauden mittauksen keskiarvoissa tilastollisesti merkitsevää eroa eli testataan, poikkeako erotusmuuttujan arvo merkitsevästi 0:sta. Lasketaan lisäksi muutoksen keskiarvolla 95% luottamusväli populaatiossa.

Paired Samples Statistics					
		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	SYST0	135.19	117	17.446	1.613
	SYST9	131.31	117	15.848	1.465

Paired Samples Test									
		Paired Differences				t	df	Sig. (2-tailed)	
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
					Lower	Upper			
Pair 1	SYST0 - SYST9	3.880	9.670	.885	2.128	5.633	4.386	116	.000

- **Tulkinta:** Tositilanteessa tarkasteltaisiin ensin, että testin oletus muutosmuuttujan normaalijakautuneisuudesta pitää paikkansa (ei näytetä tässä). Systolisen verenpaineen lasku 3.88 mmHg (95% luottamusväli 2.13-5.63) on tilastollisesti merkitsevää ($p < 0.0001$, tarkempi $p = 0.00003$)

Muuttuuko diastolin 0 ja 3 kk välillä Esim- datassa?



- Diastolinen verenpaine on noussut keskimäärin 2.9 mmHg (SE 0.4; 95% CI 2.2-3.6). Verenpaineen nousu on tilastollisesti merkitsevää ($p < 0.0001$).
- Analyze-Matched pairs, annetaan molemmat muuttujat

Testit käytännössä

- Analyysitilanteessa ei tarvitse muistaa testisuureen kaavaa, ohjelmat hoitavat sen ja antavat p-arvon
- On kuitenkin oivallettava käyttää oikeaa testimenetelmää, ja tietää sen oletukset ja tarkastettava ne

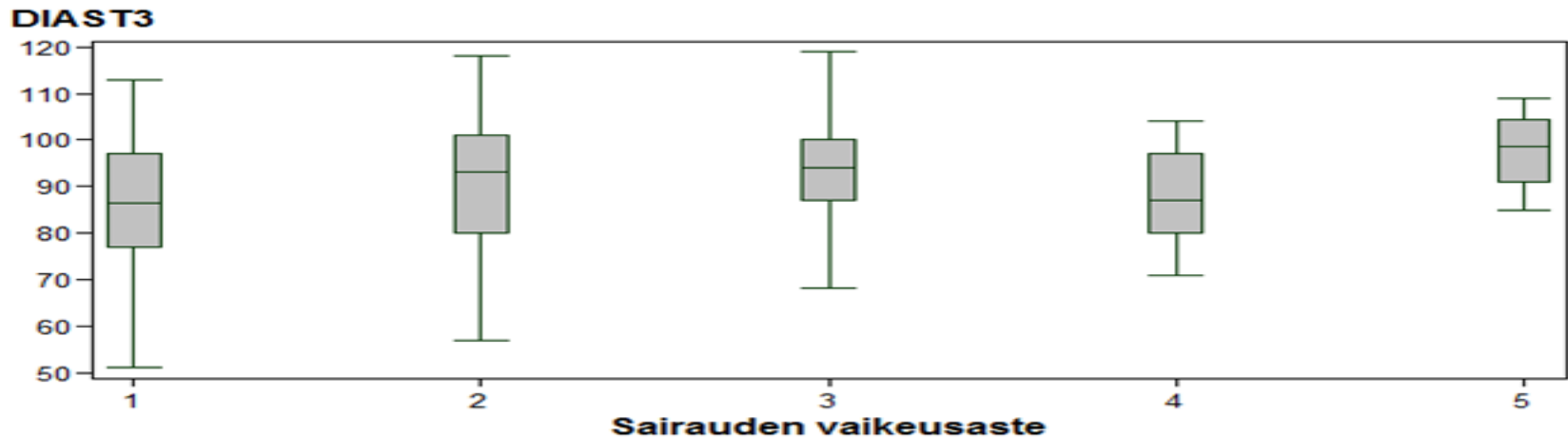
Kahden otoksen t-testi

- On jälleen ryhmät A ja B ja haluamme vertailla hemoglobiini keskiarvoja ryhmien välillä.
- On oivallettava, että nyt ryhmä on datan jakava tekijä, ja ryhmät ovat RIIPPUMATTOMIA toisistaan
- Vaste numeerinen ja normaalisti jakautunut molemmissa ryhmissä
- Tekijä kategorinen ja kaksiluokkainen

Kahden otoksen t-testi

- Lisäoletuksena varianssi
 - Erisuuret ryhmissä tai
 - Sama varianssi molemmissa ryhmissä
- Riippuen tästä oletuksesta testisuureen kaava on hieman erilainen
- Varianssin yhtäsuuruutta voi testata Levenen testillä tai valita aina tuon erisuuret varianssioletuksen

Varianssien yhtäsuuruusoletus



- Varianssien yhtäsuuruuden arviointia voi tehdä graafisesti
- Varianssien yhtäsuuruuden tarkastelu on usein perusteltua myös tutkittavaan ongelmaan liittyvistä lähtökohdista

HUOM

- Hypoteesin asettelu
- H_0 : varianssit yhtä suuret
- H_v : varianssit erisuuret
- Taas toivomme suurta p-arvoa ($p > 0.05$), jota voimme tehdä varianssien yhtäsuuruusoletuksen
- [aina kun vertailemme keskiarvoja tms toivomme pientä p-arvoa ($p < 0.05$)]

Kahden otoksen t-testi

- Tutkimushypoteesi
 - $H_0: \mu_A = \mu_B$
 - $H_1: \mu_A \neq \mu_B$
- Saadaanko datasta tukea väittämälle, että ryhmien A ja B keskiarvot ovat erisuurat?
- Testisuure (varianssit erisuurat)

$$t = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{SD_1^2}{n_1} + \frac{SD_2^2}{n_2}}}$$

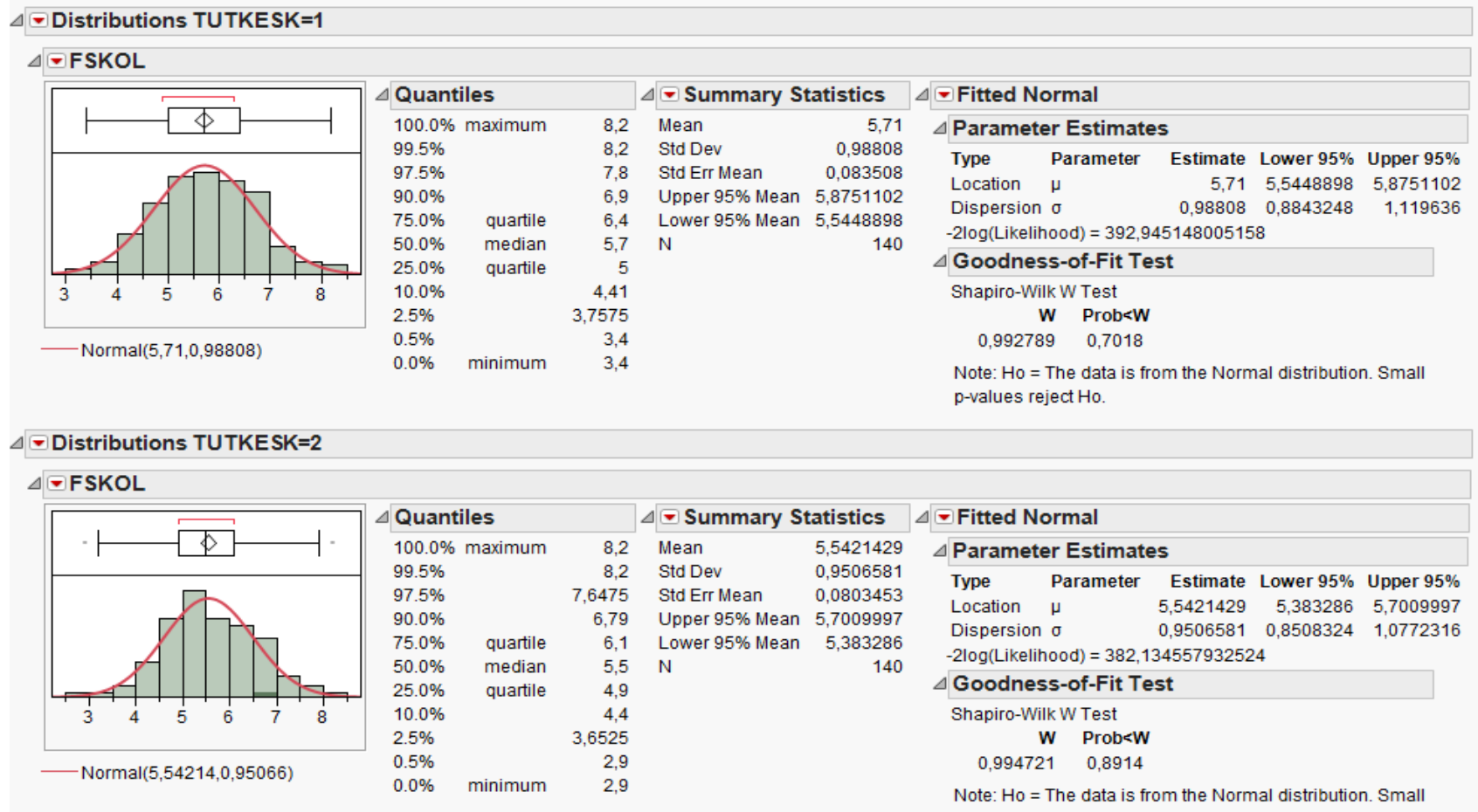
Kahden otoksen t-testi

- Jossa n_1 on ryhmän A koko ja SD_1 on ryhmän A keskihajonta ja $\bar{x}_1$ on ryhmän A keskiarvo jne
- Nyt testisuure noudattaa Studentin t-jakaumaa vapausasteella $n_1 + n_2 - 2$
- Saadaan taas p-arvo
- Voidaan laskea luottamusväli kahden ryhmän keskiarvojen erotukselle

$$(\bar{x}_1 - \bar{x}_2) \pm t_{95\%, va} \sqrt{\left(\frac{SD_1^2}{n_1} + \frac{SD_2^2}{n_2} \right)}$$

Esimerkki. Eroaako FSKOL keskiarvot tutkimuskeskusten välillä?

- Normaalisuus ryhmittäin ja hajonnan/varianssin katsominen



Analyze-Fit Y by X

-Means/Anova/pooled t (equal variances)

-T-test (unequal variances)

One-way Anova

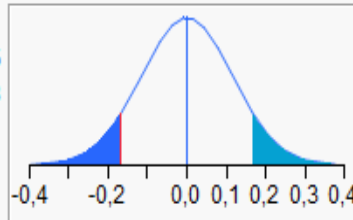
Summary of Fit

t Test

2-1

Assuming equal variances

Difference	-0,16786	t Ratio	-1,4485
Std Err Dif	0,11588	DF	278
Upper CL Dif	0,06026	Prob > t	0,1486
Lower CL Dif	-0,39598	Prob > t	0,9257
Confidence	0,95	Prob < t	0,0743



Analysis of Variance

Means for One-way Anova

Level	Number	Mean	Std Error	Lower 95%	Upper 95%
1	140	5,71000	0,08194	5,5487	5,8713
2	140	5,54214	0,08194	5,3808	5,7034

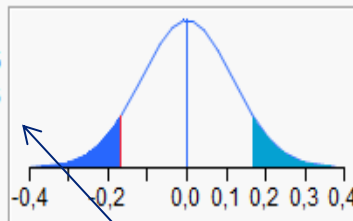
Std Error uses a pooled estimate of error variance

t Test

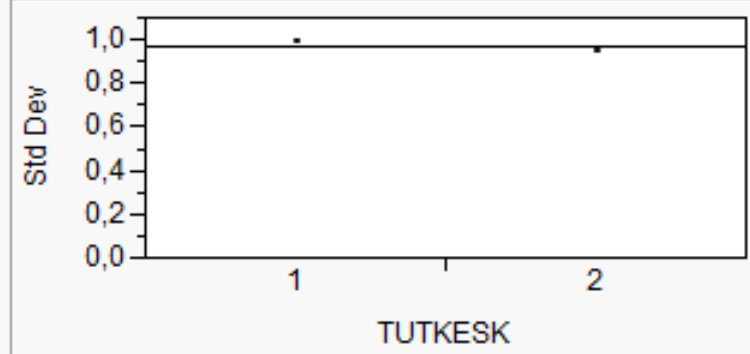
2-1

Assuming unequal variances

Difference	-0,16786	t Ratio	-1,4485
Std Err Dif	0,11588	DF	277,5866
Upper CL Dif	0,06026	Prob > t	0,1486
Lower CL Dif	-0,39598	Prob > t	0,9257
Confidence	0,95	Prob < t	0,0743



Tests that the Variances are Equal



Level	Count	Std Dev	MeanAbsDif	
			to Mean	to Median
1	140	0,9880800	0,8130000	0,8128571
2	140	0,9506581	0,7510918	0,7492857

Test	F Ratio	DFNum	DFDen	p-Value
O'Brien[.5]	0,2238	1	278	0,6366
Brown-Forsythe	0,8691	1	278	0,3520
Levene	0,8304	1	278	0,3629
Bartlett	0,2064	1	.	0,6496
F Test 2-sided	1,0803	139	139	0,6496

Reporting

- FSKOL mean was 5.7 (SD 1.0) for Group A and 5.5 (SD 1.0) for Group B.
- FSKOL means did not differ statistically significantly between Group A and Group B ($p=0.15$).
- Analysis was performed with two sample t-test with assumption of equal variances. Equality of variances was tested with Levene test. Normal distribution assumption was checked with Shapiro-Wilk's test.

Mitä jos ryhmiä onkin 3 tai enemmän?

- Nyt onkin ryhmät A, B ja C
- Ei voida silloin käyttää kahden otoksen t-testiä
- Vaste numeerinen (hemoglobiini) ja normaalisti jakautunut
- Tekijä kategorinen (ryhmä) ja ryhmät riippumattomat
- Varianssien yhtäsuuruus oletuksena
- $H_0: \mu_A = \mu_B = \mu_C$
- $H_1: \mu_i \neq \mu_j$ (joku keskiarvo eroaa jostakin $\mu_A \neq \mu_B$ ja/tai $\mu_A \neq \mu_C$ ja/tai $\mu_B \neq \mu_C$) jne

Yksisuuntainen varianssianalyysi

- 1-way ANOVA (analysis of variance)
- 1-way viittaa yhteen ryhmittelevään tekijään
- Varianssianalyysi sana viittaa siihen, että datassa esiintyvä vaihtelu jaetaan osiin; ryhmien sisäiseen, ryhmien välisiin ja niiden avulla muodostetaan testisuureet

- Jos ryhmä-tekijä on tilastollisesti merkitsevä, niin saamme siis selville, että jonkun ryhmän keskiarvo eroaa jostakin
- 'Hemoglobin means differed between the treatment groups statistically significantly ($p=0.017$)'
- -> tarvitaan lisäselvitystä eli parittaisia testejä/vertailuja A vs B, B vs C ja A vs C
- Nämä vertailut tehdään mallin sisällä, ja se on parempi kuin, että nyt tekisit 3 kahden otoksen t-testiä!

Monivertailut

- Ongelma alkaa kuitenkin ilmetä, että nyt siis tarvitaan 3 lisävertailua (A vs B, A vs C ja B vs C), jotta saadaan koko ilmiö selville
- => tyyppin I virhemahdollisuus kasvaa (toteamme eron, vaikka sitä ei todellisuudessa ole)
- Kiristämme merkitsevän p-arvon ehtoja, esim niin, että vasta p-arvo pienempi kuin 0.017 (Bonferroni-korjaus $0.05/\text{vertailujen määrä}$) on merkitsevä näissä parittaisissa vertailuissa

Monivertailut

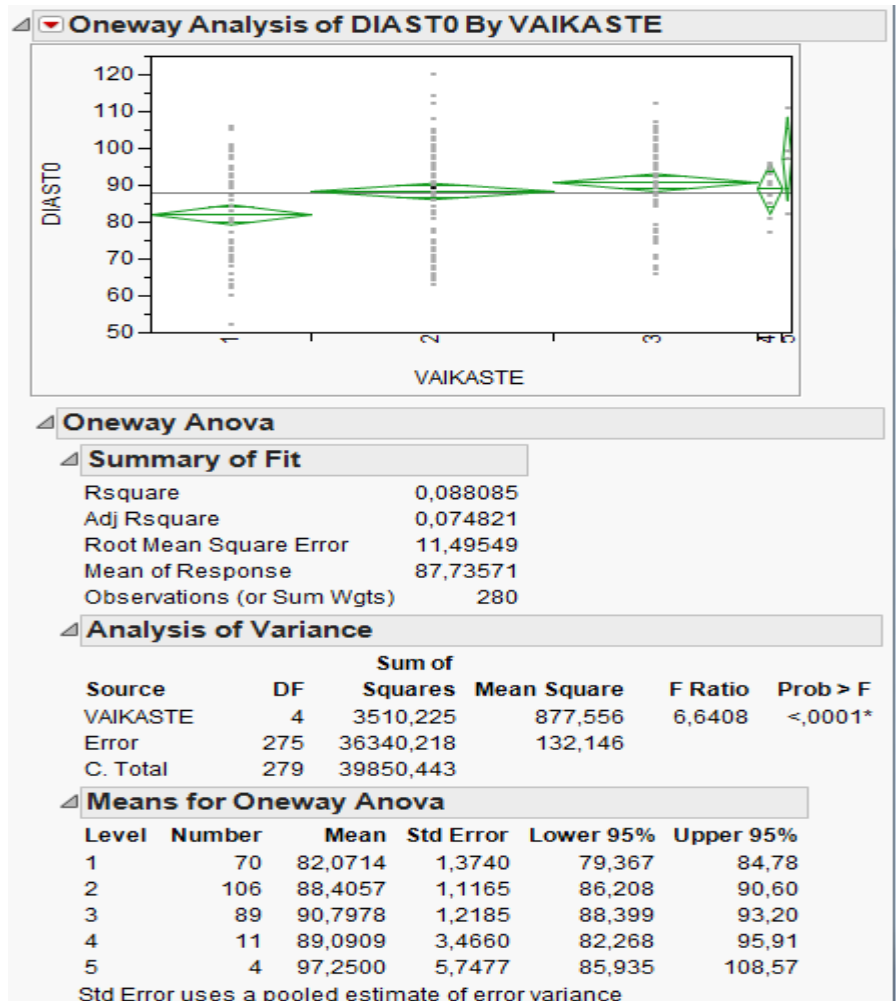
- Laskumetodeja on paljon, kaikki pyrkivät siihen, että tyypin I virhe säilyy tasolla 0.05 useammasta vertailusta huolimatta
- Bonferroni – merkitsevä p-arvo on $0.05/\text{vertailujen määrä}$
- Tukey HSD (honestly significant difference) vertailee kaikki parittaiset erot
- Dunnett – joku ryhmä tärkeä ja vain siihen verrataan (kontrolli, placebo)

Monivertailut

- Monivertailumahdollisuus usein valittavissa ohjelmistoissa
- Jos ryhmiä paljon, suuret erot eri metodeissa

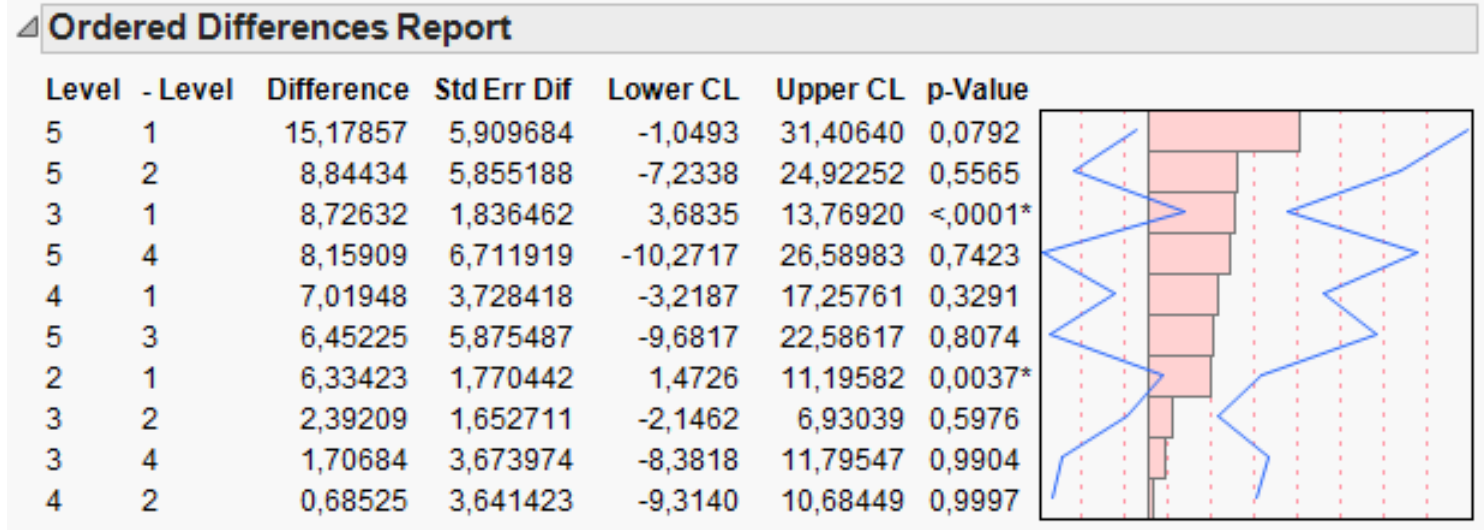
Tutkitaan, onko eroavatko diastolisen verenpaineen keskiarvot eri vaikeusasteen luokissa

Fit Y by X – Means/Anova



Monivertailut Tukeys-Kramer HSD

Compare Means-All pairs Tukey HSd



- Huomaa, että 5 vs 1 ja 5 vs 2 ei tule tilastollisesti merkitseväksi, vaikka ero suurin (5 oli se todella pieni ryhmä)

- Jos ei ole samaa kirjainta vierekkäin, eroavat
- Esim 3-1 eroavat, koska 3 vain A-kirjain ja 1 vain B-kirjain

Connecting Letters Report		
Level		Mean
5	A B	97,250000
3	A	90,797753
4	A B	89,090909
2	A	88,405660
1	B	82,071429

2-suuntainen varianssianalyysi

- 2-way ANOVA
- Vaste numeerinen (ja vasteita 1)
- 2 tekijää, kategorisia (tekijöissä voi olla vaikka kuinka monta luokkaa)
 - Ryhmät A ja B
 - Miehet ja naiset
- Oletetaan taas normaalijakauma vasteelle
- Samavarianssisuusoletus

2-suuntainen varianssianalyysi

- Kiinnostuksen kohteena
 - Eroavatko ryhmien A ja B keskiarvot ? (=päävaikutus)
 - Eroavatko miesten ja naisten keskiarvot? (=päävaikutus)
 - Onko lääkevaikutus erilainen miehillä kuin naisilla? (=yhdysvaikutus)

Malli

Keskiarvot	Ryhmä A	Ryhmä B	Total
Miehet	20	50	35
Naiset	50	20	35
	35	35	

Jos tarkastelisimme vain ryhmä ja sukupuoli päävaikutusta, molemmille tulisi p-arvo, ja päätelmämme olisi ettei mikään ero ole tilastollisesti merkitsevä.

Siksi on tarpeen lisätä malliin yhdysvaikutus!

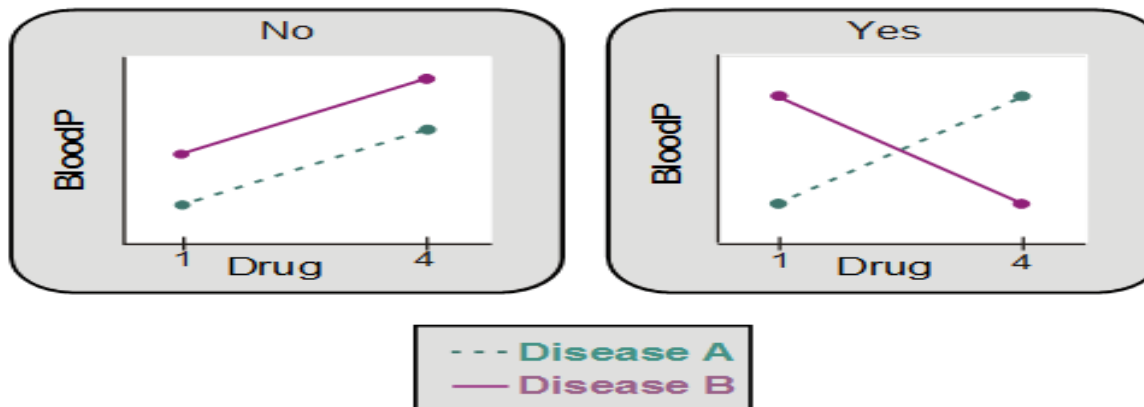
Yhdysvaikutus

- Malli on siis
- $Vaste = Ryhma + sukupuoli + ryhma * sukupuoli$
- Ryhma $p=0.89$
- Sukupuoli $p=0.91$
- Ryhma*Sukupuoli $p=0.012$
- Miehillä lääkevaikutus oli tilastollisesti merkitsevästi erilainen kuin naisilla siten, että miehillä lääke B nosti arvot 50 ja lääke A 20:een, kun taas naisilla...

Yhdysvaikutus

- Yhdysvaikutus kuvassa

Interactions



- Drug 1 vaikutus hyvin erilaista onko Disease B potilaalla vai ei

FSKOL=ryhma+vaikaste+ryhma*vaikaste

Fit Y by X

Response FSKOL

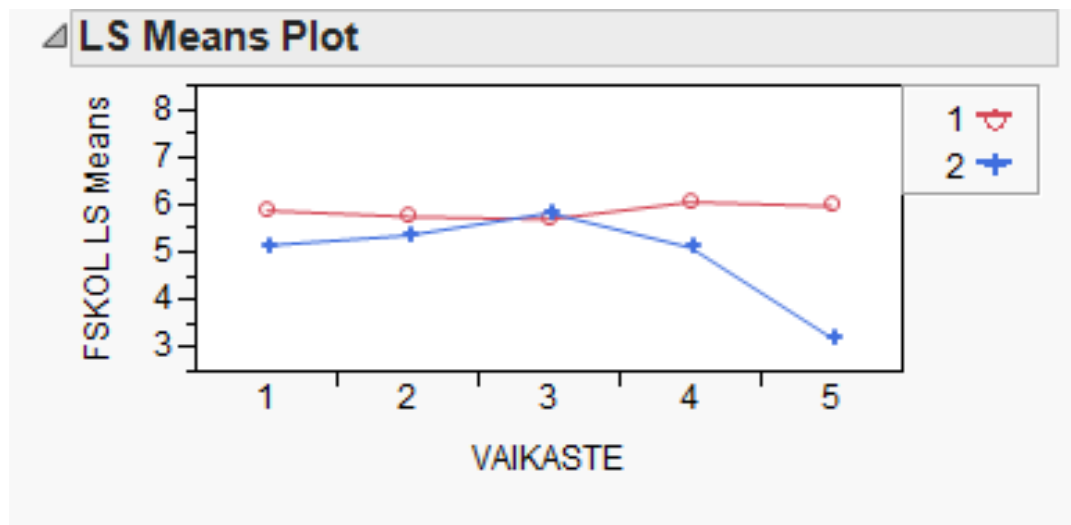
Whole Model

Effect Tests

Source	Nparm	DF	Sum of Squares	F Ratio	Prob > F
RYHMA	1	1	10,461639	11,7793	0,0007*
VAIKASTE	4	4	4,989633	1,4045	0,2327
RYHMA*VAIKASTE	4	4	10,132311	2,8521	0,0242*

- Tutkitaan FSKOL ryhmän ja vaikeusasteen vaikutuksia ja ryhmän ja vaikeusasteen yhdysvaikutusta = onko vaikeusasteen vaikutus FSKOLiin erilainen eri ryhmissä?

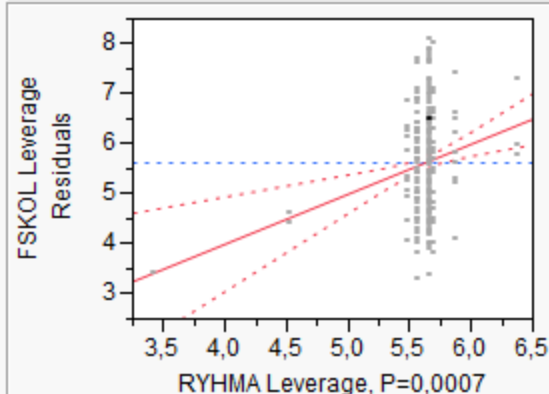
LS Means = Least Squares means



- LS Means kuvaa mallin sovitettuja keskiarvoja
- Kuvasta näkyy, että vaikeusasteen yhteys FSKOLiin on erilainen eri ryhmissä
- Jos olisi vain ryhmävaikutus, niin kuvassa kaksi vaakasuoraa eri tasoilla
- Jos vain vaikeusasteen vaikutus, keskiarvoviivat eivät olisi vaakasuoria (esim molemmat laskevia), mutta ryhmien keskiarvoviivat olisivat miltei päällekkäin
- Jos yhdysvaikutusta, niin keskiarvoviivat ovat erimalliset

RYHMA

Leverage Plot

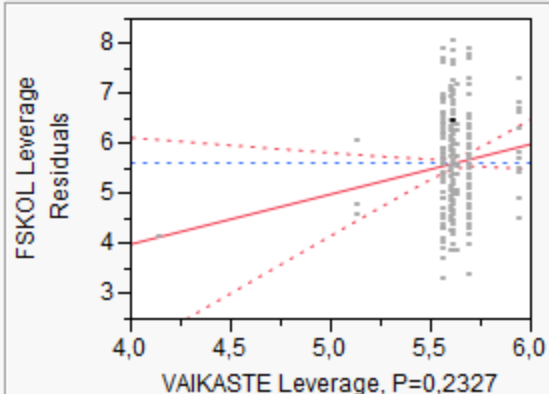


Least Squares Means Table

Least			
Level	Sq Mean	Std Error	Mean
1	5,8685281	0,13347091	5,76859
2	4,9256508	0,24012222	5,32022

VAIKASTE

Leverage Plot

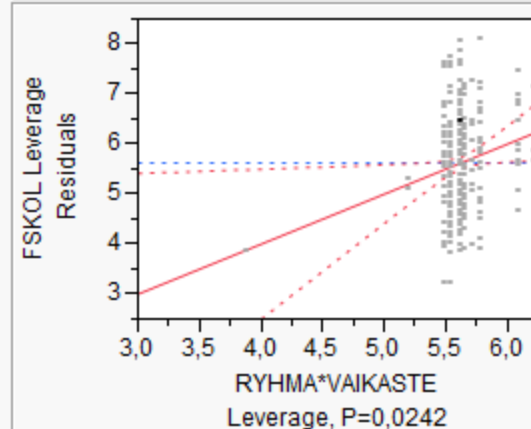


Least Squares Means Table

Least			
Level	Sq Mean	Std Error	Mean
1	5,5104575	0,11268583	5,50000
2	5,5539235	0,09731969	5,62075
3	5,7599550	0,13342724	5,71573
4	5,5777778	0,36835882	5,88182
5	4,5833333	0,54410203	5,27500

RYHMA*VAIKASTE

Leverage Plot



Least Squares Means Table

Least		
Level	Sq Mean	Std Error
1,1	5,8764706	0,16162239
1,2	5,7507042	0,11184377
1,3	5,6932432	0,10955321
1,4	6,0555556	0,31413746
1,5	5,9666667	0,54410203
2,1	5,1444444	0,15706873
2,2	5,3571429	0,15929676
2,3	5,8266667	0,24332983
2,4	5,1000000	0,66638617
2,5	3,2000000	0,94241237

Studentin t (ei monivertailukorjausta)

Level		Least Sq Mean
1,4	A	6,0555556
1,5	A B C	5,9666667
1,1	A	5,8764706
2,3	A B	5,8266667
1,2	A	5,7507042
1,3	A B	5,6932432
2,2	B C	5,3571429
2,1	C	5,1444444
2,4	A B C D	5,1000000
2,5	D	3,2000000

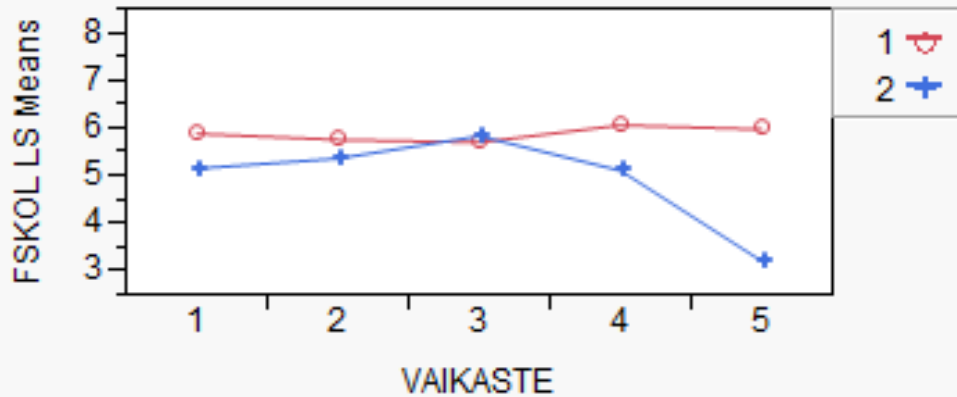
Levels not connected by same letter are significantly different.

		Least Sq Mean
1,4	A B	6,0555556
1,5	A B	5,9666667
1,1	A	5,8764706
2,3	A B	5,8266667
1,2	A B	5,7507042
1,3	A B	5,6932432
2,2	A B	5,3571429
2,1	B	5,1444444
2,4	A B	5,1000000
2,5	A B	3,2000000

Levels not connected by same letter are significantly different.

Tukey'n HSD eli monivertailukorjaus

LS Means Plot



- Tällä kurssilla emme ehdi opettelemaan miten tutkitaan yhdysvaikutuksia tarkemmin, eli purkamaan yhdysvaikutusta osiin, tärkeää kuitenkin testata ja kuvailla yhdysvaikutus

N-suuntainen varianssianalyysi

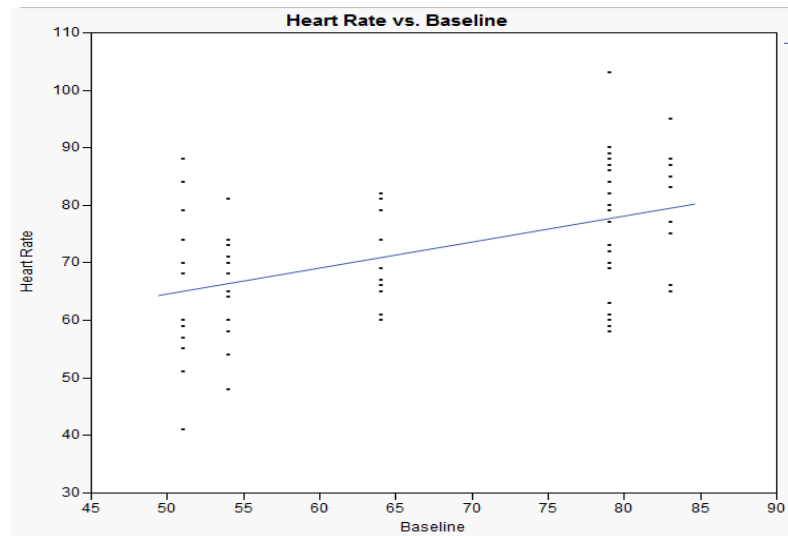
- Kun tekijöitä (kategorisia) tulee lisää, niin pitää harkita mitä yhdysvaikutuksia laitetaan malliin, sillä yhdysvaikutukset monimutkaistuvat
Ryhmä*Sukupuoli*Sairausaste jne
- Oltava tarkkana, jos jotkut tekijät ovat toisiinsa yhteydessä (esim siviilisääty ja asumismuoto), tekijöiden vaikutukset voivat 'syödä' toisiaan, sekoittua (tutki tekijöiden korrelaatioita)

Mallin rakennus

- Usein aloitetaan monimutkaisemmasta mallista, missä mukana joitakin kiinnostavia yhdysvaikutuksia
- Sitten mallia yksinkertaistetaan poistamalla aina yksi ei-merkitsevä yhdysvaikutus ja lopulta päävaikutuksia [malleja voi myös vertailla formaalisti, mutta se on jatkokurssin asia]
- Toinen strategia on valita huolella malli, pitää siinä myös kaikki ei-merkitsevät tekijät ja raportoida se
- Tärkeintä on, että tarkastaa mallin oletukset

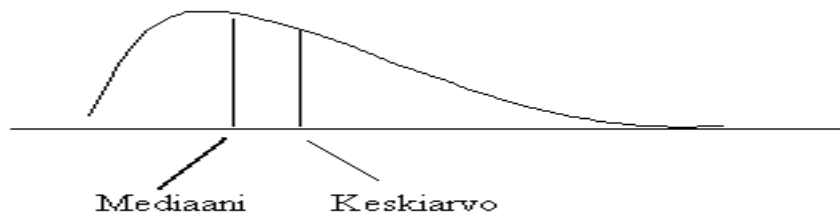
Kovarianssianalyysi

- Jos osa tekijöistä on numeerisia, niin analyysimetodia kutsutaan kovarianssianalyysiksi
- Silloin ei tule tulokseksi tietenkään LSMEANS tai keskiarvoeroja, vaan tutkitaan numeerista yhteyttä (regressiosuora tulee myöhemmin);



Norm jak oletus

- Jos normaalijakaumaoletus ei ole voimassa
 - Jos jakauma on (oikealle) vino, kokeile muunnosta
 - Logaritmi (luonnollinen)
 - Neliöjuuri
 - Eli ota logaritmi jokaisesta vasteen havainnoista ja luo siten uusi vastemuuttuja



Muunnos

- Jos data jakauma normaali muunnoksen jälkeen, tee analyysi kuten aiemmin on tehty
- On kuitenkin huomioitava, että mallin keskiarvot ovat logaritmi-asteikolla, joten voit kaivat niiden palauttamista alkuperäiseen asteikkoon, esim e^{mean}
- Luottamusvälin palautuksessa palauta alkupää ja loppupää e^{low} ja e^{high}
- Huom. Keskihajontaa tai keskivirhettä ei saa näin palauttaa!

Residuaalit

- Tarkkaan ottaen, varianssianalyysissä oletetaan residuaalien normaalijakaantuneisuus
- $\text{Residual} = \text{Orig arvo} - \text{mallin sovitettu arvo}$
- Tärkeää, mikäli monimutkaisempi malli

Epäparametriset testit

- Jos muunnoskaan ei auta, tai et halua käyttää normaalijakaumaoletusta, siirrytään **epäparametrisiin** testeihin
- Niissä ei käytetä keskiarvoja, ei hajontoja tms, ainoastaan järjestyslukuja
 - Vastedata siis järjestetään suuruusjärjestykseen, kaikille havainnoille annetaan järjestysnumero ja sitä käytetään alkuperäisen datan sijaan analyyseissä!

Järjestysluvut

Rank Scores

Treatment	A	A	A	A	A	B	B	B	B	B
Response	2	5	7	8	10	6	9	11	13	15
Rank Score	1	2	4	5	7	3	6	8	9	10
	SUM					SUM				
	19					36				

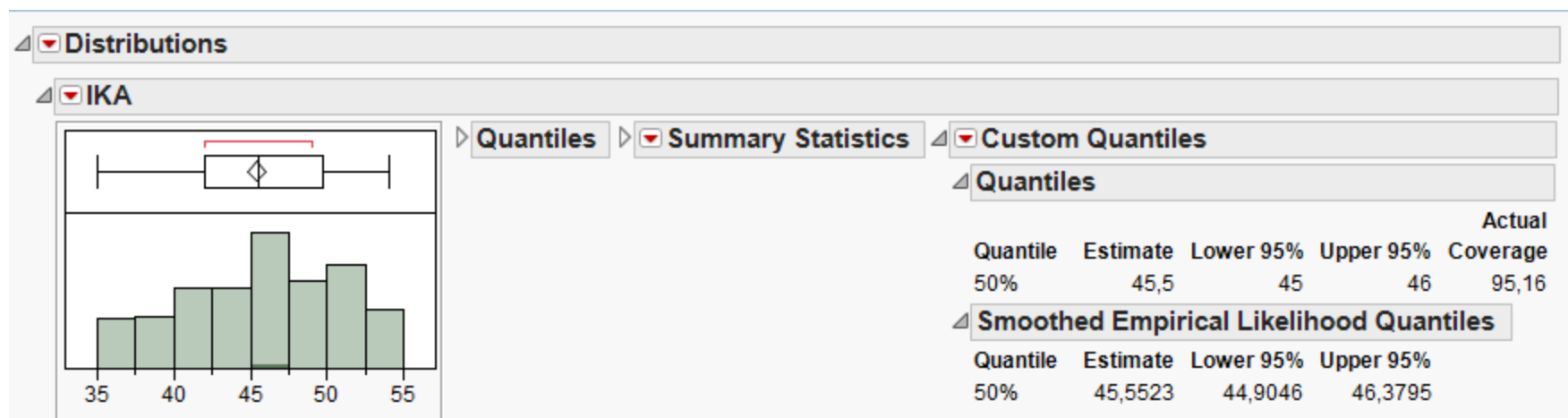
Epäparametriset testit

- Muuten ajatusmalli on sama
- Nollahypoteesi on useimmiten, että mediaanit ovat samat, eli tutkitaan jakaumien sijainnin erilaisuutta
- Oletuksena usein jakauman samanmuotoisuus (esim yksihuipukkuus toteutuu eri ryhmissä)
- Vaste on edelleen numeerinen
- Tekijät kategorisia

Epäparametrinen

- Tilanne A: 1 vaste, numeerinen, mutta normaalijakaumaoletus ei voimassa
- Ei tekijöitä
 - Usein järkevintä laskea mediaanin luottamusväli ja sen avulla tehdä päätelmät (tässä oletetaan, että keskiluku on se kiinnostuksen kohde)
 - Esim diagnostiikassa, tai laboratoriomuuttujien viitearvoissa kiinnostuksen kohde on usein esim 5% tai 1% persentiili

- Mediaanin (myös muut kvantiilit) luottamusväli



- Löytyy Analyze-Distributions
 - Muuttujan nimen kohdalta  Display Options - Custom Quantiles - 0.5

Muutos tms

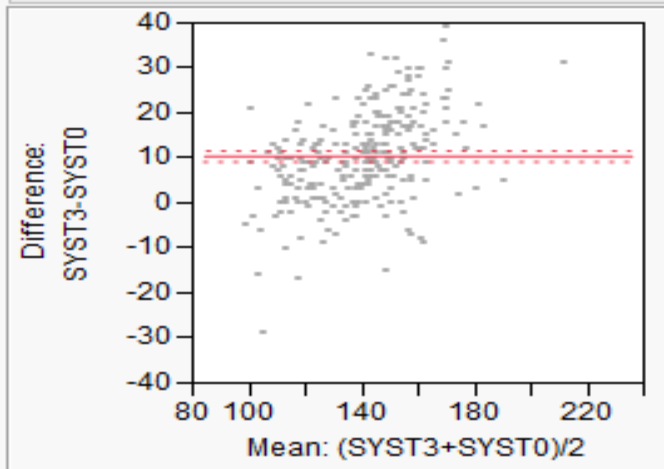
- Tutkitaan muutosta systolisen verenpaineen 0 ja 3 aikapisteen välillä. Ei haluta tehdä normaalijakaumaoletusta. Tutkitaan, eroaako muutos 0:sta.
- Se tehdään Wilcoxonin merkittyjen järjestyslukujen testin/Wilcoxonin merkkitesti avulla (Wilcoxonin signed rank test/Wilcoxon matched pairs signed ranks test)

Arvioija	Tuote		Erotus	Itseisarvo	Sijaluku
	A	B	A - B		
1	6	4	2	2	5
2	8	5	3	3	7.5
3	4	5	-1	1	2
4	9	8	1	1	2
5	4	1	3	3	7.5
6	7	9	-2	2	5
7	6	2	4	4	9
8	5	3	2	2	5
9	6	7	-1	1	2
10	8	2	6	6	10

- Lasketaan erotus, sen itseisarvo ja annetaan niille järjestysluvut

Matched Pairs

Difference: SYST3-SYST0



SYST3	145,498	t-Ratio	16,57407
SYST0	135,149	DF	274
Mean Difference	10,3491	Prob > t	<,0001*
Std Error	0,62441	Prob > t	<,0001*
Upper 95%	11,5784	Prob < t	1,0000
Lower 95%	9,11983		
N	275		
Correlation	0,8872		

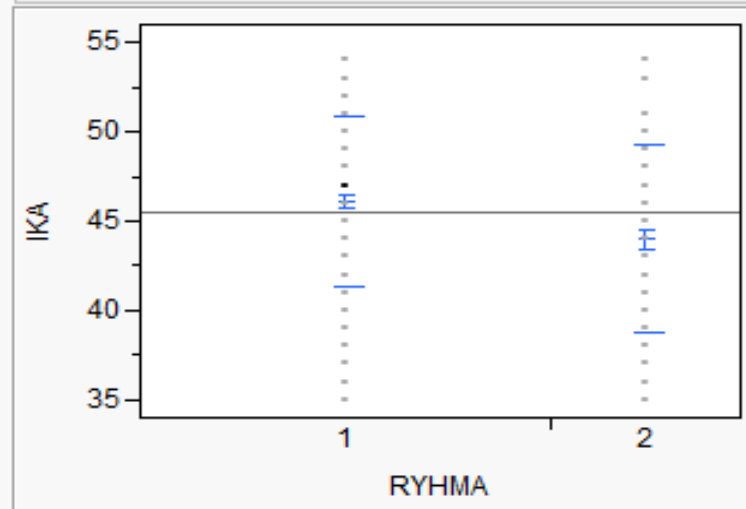
Wilcoxon Signed Rank

	SYST3-SYST0
Test Statistic S	15351,5
Prob> S	<,0001*
Prob>S	<,0001*
Prob<S	1,0000

Epäparametrinen

- Tilanne B: 1 vaste, numeerinen, ei normaalijakaumaoletus voimassa
- Tekijänä 1 kategorinen
- Silloin testissä verrataan jakaumien sijainnin eroja
- Oletuksena jakaumien samanmuotoisuus
- **HUOM kategorinen tekijä jakaa datan kahteen riippumattomaan osaan**
- Testi on Wilcoxonin rank sum test = Wilcoxonin järjestyssummatesti = Mann-Whitneyn U-testi

▲ Oneway Analysis of IKA By RYHMA



▲ Means and Std Deviations

Level	Number	Mean	Std Dev	Std Err		
				Mean	Lower 95%	Upper 95%
1	191	46,1047	4,73615	0,34270	45,429	46,781
2	89	44,0000	5,27214	0,55885	42,889	45,111

▲ Wilcoxon / Kruskal-Wallis Tests (Rank Sums)

Level	Count	Score Sum	Expected		(Mean-Mean0)/Std0
			Score	Score Mean	
1	191	28869,5	26835,5	151,149	3,229
2	89	10470,5	12504,5	117,646	-3,229

▲ 2-Sample Test, Normal Approximation

S	Z	Prob> Z
10470,5	-3,22911	0,0012*

▲ 1-way Test, ChiSquare Approximation

ChiSquare	DF	Prob>ChiSq
10,4323	1	0,0012*

▲ Analysis of Variance

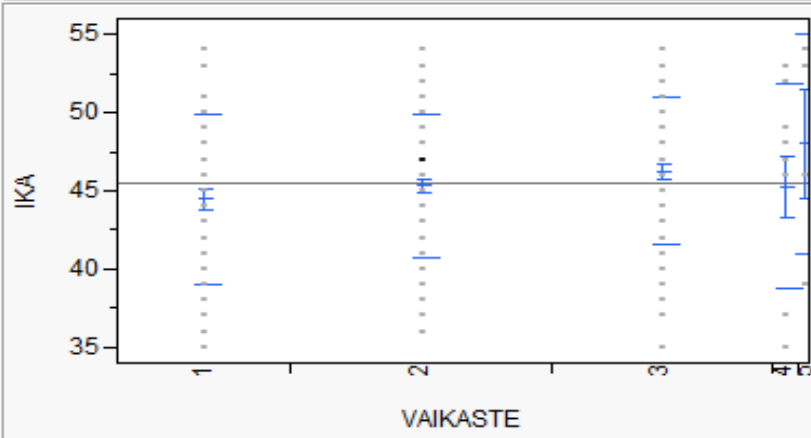
Source	DF	Sum of Squares	Mean Square	F Ratio	Prob > F
RYHMA	1	268,9371	268,937	11,1457	0,0010*
Error	278	6707,9058	24,129		
C. Total	279	6976,8429			

Epäparametrinen testi
ja 1-suuntainen
anova. Ei suurta eroa,
jos kutakuinkin
normaalijakautunut

Epäparametrinen

- Tilanne C. 1 vaste, numeerinen, ei normaalijakaumaoletusta
- Tekijöitä 1, mutta monta kategoriaa (esim ryhmät AA, BB, CC, DD), kategorinen
- Tutkitaan taas jakauman sijainnin eroa, jakaumat samanmuotoisia
- Silloin testi on Kruskal-Wallis test

▼ Oneway Analysis of IKA By VAIKASTE



▼ Means and Std Deviations

Level	Number	Mean	Std Dev	Std Err		
				Mean	Lower 95%	Upper 95%
1	70	44,4714	5,43665	0,6498	43,175	45,768
2	106	45,3113	4,63601	0,4503	44,418	46,204
3	89	46,2472	4,71767	0,5001	45,253	47,241
4	11	45,2727	6,51292	1,9637	40,897	49,648
5	4	48,0000	6,97615	3,4881	36,899	59,101

▼ Wilcoxon / Kruskal-Wallis Tests (Rank Sums)

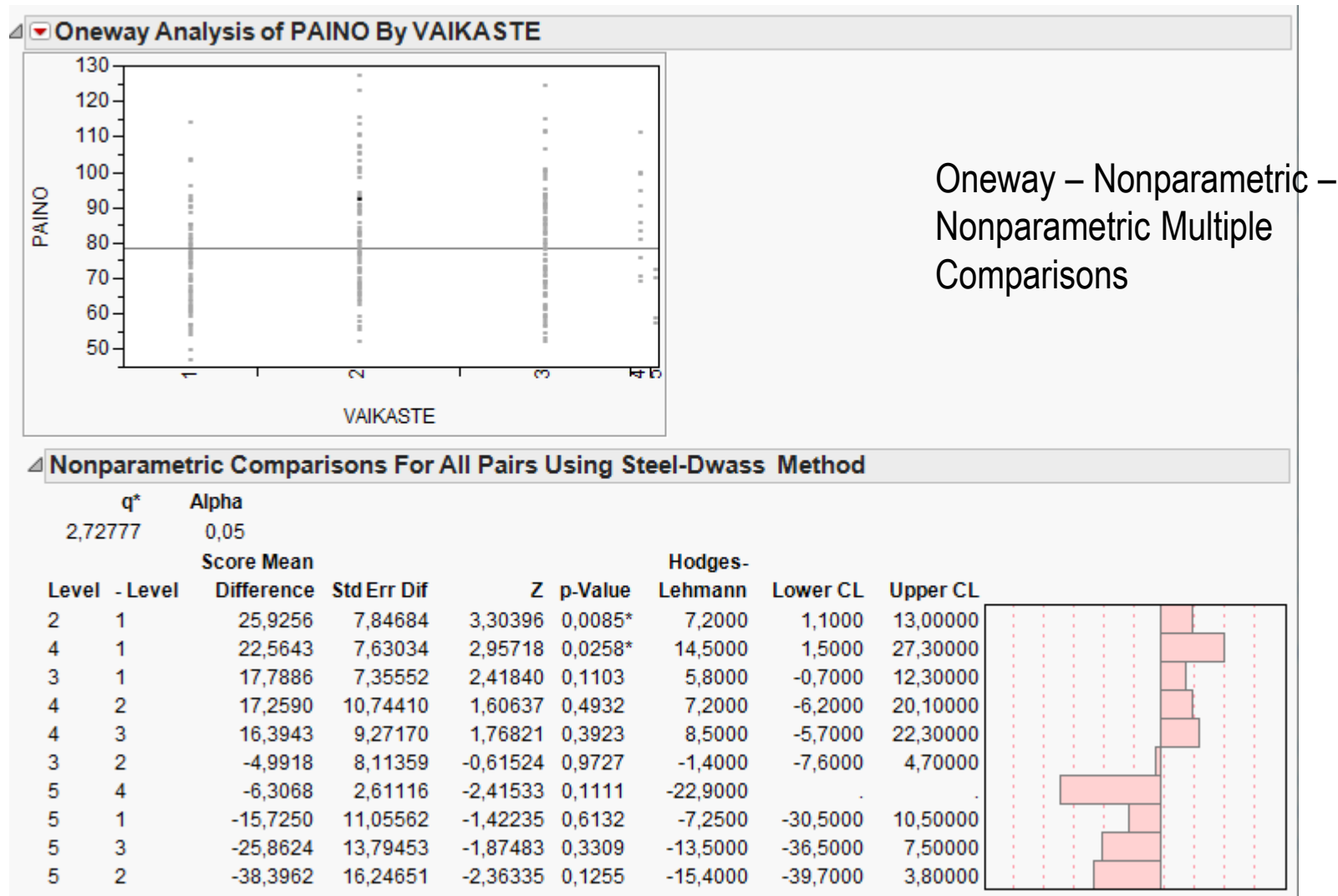
Level	Count	Score Sum	Expected		
			Score	Score Mean	(Mean-Mean0)/Std0
1	70	8715,50	9835,00	124,507	-1,911
2	106	14580,0	14893,0	137,547	-0,476
3	89	13709,0	12504,5	154,034	1,912
4	11	1608,50	1545,50	146,227	0,238
5	4	727,000	562,000	181,750	1,025

▼ 1-way Test, ChiSquare Approximation

ChiSquare	DF	Prob>ChiSq
6,4751	4	0,1664

- Ikä ei eronnut tilastollisesti merkitsevästi vaikeusasteen eri luokissa ($p=0.17$).

- Jos tulos tilastollisesti merkitsevä, niin JMP pystyy tekemään parittaiset testit



Epäparametriset testit

- Epäparametriset testit eivät ole niin voimakkaita kuin parametriset testit (esim ANOVA), eli useammin jää ero havaitsematta, siksi sitä ei kannata käyttää 'automaattisesti'
- Jos on epävarma normaalijakaumaoletuksesta, kannattaa tehdä epäparametrinen testi ns sensitiivisyysanalyysinä, eli sen kautta voi tarkastella poikkeavatko tulokset vähän/paljon toisistaan

Riippuvuus

- Kaksi numeerista muuttujaa, meitä kiinnostaakin muuttujien välien yhteys/riippuvuus
 - Esim luuntiheyden ja painon riippuvuus toisistaan
 - Pituuden ja painon
- Piirretään sirontakuvio

