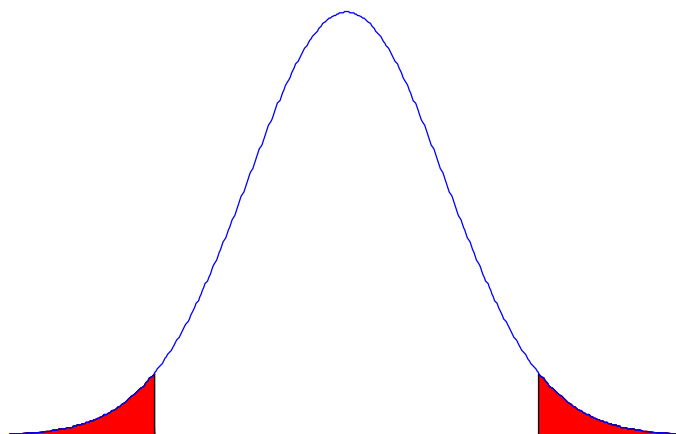


MINITAB i korthet

release 16

Jan-Eric Englund



Jan-Eric Englund är universitetslektor i statistik vid Enheten för statistik i Alnarp på Område Agrosystem.

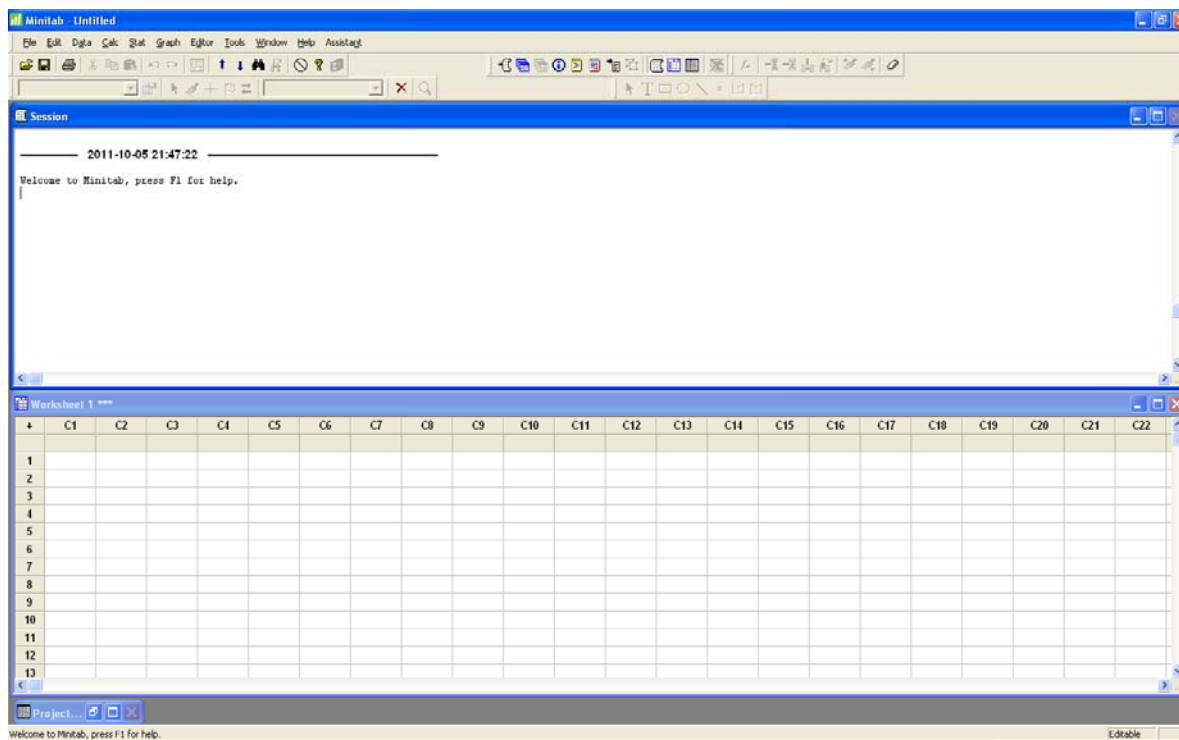
Minitab i korthet utgår från de statistiska metoder som brukar finnas i en grundkurs och avser inte på något sätt att vara en heltäckande beskrivning av Minitab.

Det är en uppdatering av tidigare kompendier för Minitab 14 och Minitab 15 och på några ställen är inte figurerna uppdaterade till den nya versionen.

1.	INLEDNING	1
2.	INLÄSNING AV OBSERVATIONERNA	4
3.	MINITAB TILLSAMMANS MED WORD OCH EXCEL	6
3.1.	<i>Från Excel till Minitab</i>	6
3.2.	<i>Från Minitab till Word</i>	7
3.3.	<i>Läsa in data i Minitab</i>	7
3.4.	<i>Editera Session Window i Minitab</i>	8
3.5.	<i>Editera Worksheet i Minitab</i>	8
	Ändra namn på Worksheet	8
	Kopiera delar av kolumner	8
	Använda matematiska funktioner	9
4.	BESKRIVANDE STATISTIK	10
4.1.	<i>Numeriska metoder</i>	10
4.2.	<i>Boxplot</i>	11
4.3.	<i>Scatterplot</i>	12
4.4.	<i>Histogram</i>	13
5.	STATISTISKA METODER FÖR ETT STICKPROV	14
5.1.	<i>Normalfördelning och ett stickprov</i>	14
	Normalfördelningspapper	16
5.2.	<i>Ej normalfördelning och ett stickprov</i>	17
	Teckentest	17
	Wilcoxon's tecken-rang-test	18
6.	STATISTISKA METODER FÖR TVÅ STICKPROV	19
6.1.	<i>Normalfördelning och två stickprov</i>	19
6.2.	<i>Test av lika varians</i>	21
6.3.	<i>Ej normalfördelning och två stickprov</i>	22
7.	STATISTISKA METODER VID FLERA STICKPROV	24
7.1.	<i>Normalfördelning och flera stickprov</i>	24
7.2.	<i>Ej normalfördelning och flera stickprov</i>	26
8.	BLOCKFÖRSÖK	26
9.	KORRELATION	30
10.	REGRESSION	32
11.	CHI-TVÅ-TEST	33

1. Inledning

När Minitab har öppnats finns det två fönster av intresse, det översta som kallas Session Window (översätts här med *resultatfönster*) och det nedersta som kallas Worksheet (översätts här med *datafönster*). Datafönstret påminner om ett Excelark, men det finns skillnader.



En körning i Minitab går i allmänhet till så att man först läser in data i datafönstret, sedan väljer en meny för att bestämma vad man skall göra för analys på data, och därefter får man utskriften i resultatfönstret. Om man också har gjort en figur så kommer den i ett eget fönster.

Kontrollera först om din version av Minitab vill ha decimalpunkt eller decimalkomma, om det blir fel framgår det av att kolumnen ändras till C1-T som visar att det är en textkolumn, och textkolumner går inte att använda till numeriska beräkningar¹. *I det följande antas att programmet är inställt för decimalpunkt.*

Exempel 1. Inledande exempel

Beräkna medelvärde och standardavvikelse för följande värden:

8.7, 8.1, 9.8, 6.1, 7.1.

(Kontrollera alltså om din version av Minitab använder decimalpunkt och inte decimalkomma.)

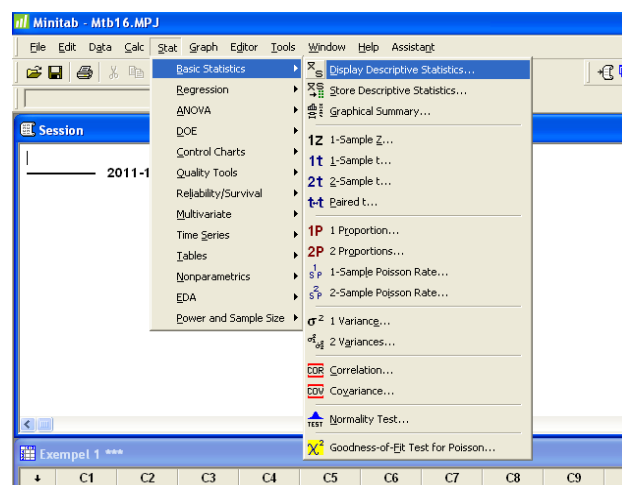
Läs in värdena i kolumn C1 enligt (variabeln kallas här alltså y)

¹ Om man av misstag har fått en kolumn med siffror att bli en textkolumn så kan man ändra det med kommandot Data ► Code ► Text to Numeric...

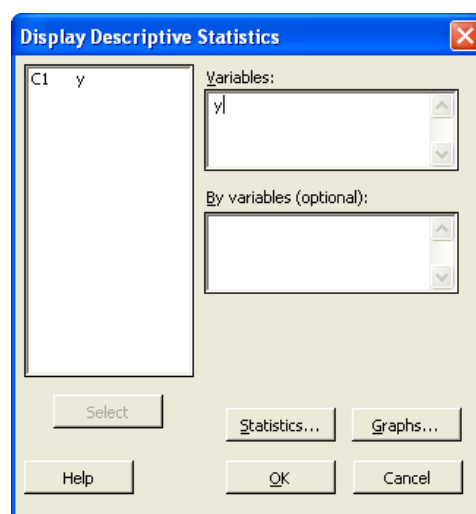
Exempel 1 ***			
	C1	C2	C3
	y		
1	8.7		
2	8.1		
3	9.8		
4	6.1		
5	7.1		
6			
7			

(I senare avsnitt beskrivs hur man ger datafönstret ett annat namn än Worksheet 1.) När nu värdena är inlästa kan man göra analysen genom att välja lämplig meny. I det här fallet väljs

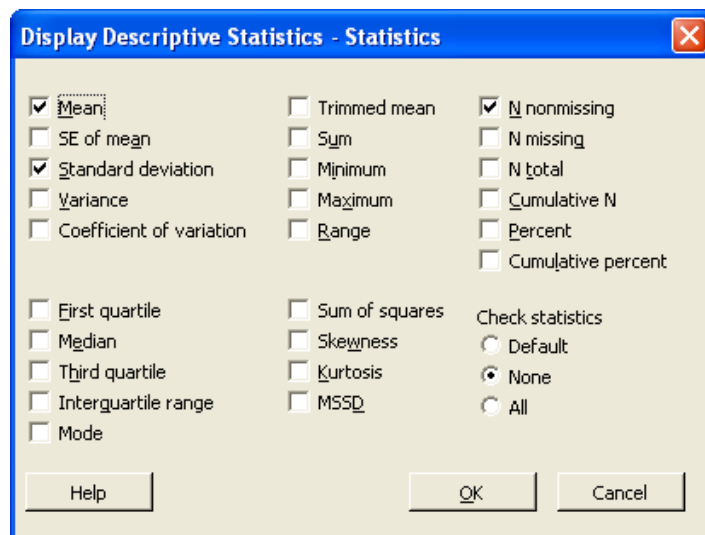
Stat ► Basic Statistics ► Display Descriptive Statistics...



På den meny som kommer upp finns nu lite olika sätt att välja den variabel man vill räkna på, men det enklaste är att först ställa markören i rutan märkt Variables: klicka en gång på raden märkt C1 y och därefter på **Select** (det fungerar också att dubbelklicka på C1 y).

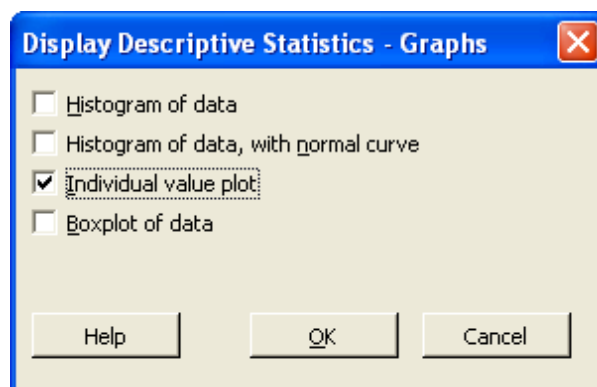


För att bestämma vilka statistiska mått man vill beräkna kan man klicka på **Statistics...** för att få upp menyn



I det här fallet väljer vi att kryssa för Mean (medelvärde), Standard deviation (standardavvikelse) och N nonmissing (antalet observationer som inte har saknat värde). Tryck **OK** när du är nöjd.

För att få olika typer av figurer som beskriver datamaterialet skall man välja **Graphs...**. I detta fall vill vi ha en "Individual value plot" och markerar därför enligt

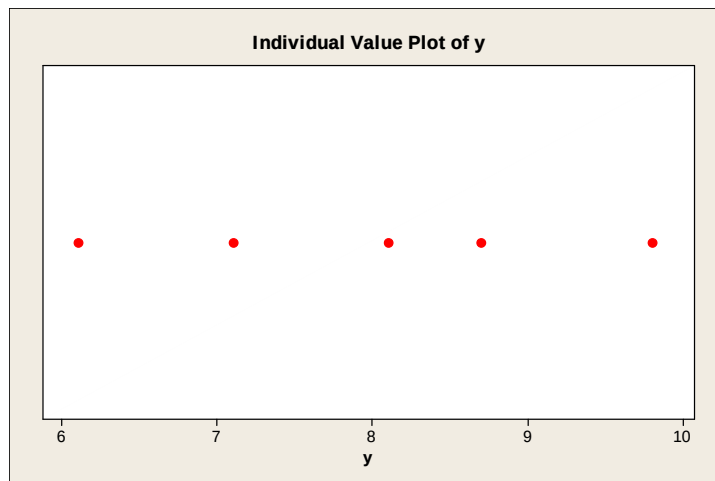


och trycker på **OK**. Välj sedan **OK** och studera resultatet. I resultatfönstret står nu bland annat

Descriptive Statistics: y

Variable	N	Mean	StDev
y	5	7.960	1.428

Under rubrikerna Mean resp. StDev (= Standard Deviation) står resultaten, 7.960 resp. 1.428. Man får också ett grafiskt fönster som ser ut enligt



Figuren hamnar alltså i ett eget fönster och kan redigeras. ■

Det finns alltså tre typer av fönster, Session som har "efternamnet" .TXT, Worksheet med "efternamnet" .MTW och Graphs med "efternamnet" .MGF. Dessutom kan hela projektet sparas med namnet .MPJ.

2. Inläsning av observationerna

Då man har samlat in observationerna och skall använda en dator för att hålla reda på eller analysera datamaterialet bör detta struktureras på ett sätt som gör att det fortsatta arbetet blir så enkelt som möjligt. Trots att det finns många olika datorprogram för statistisk bearbetning gäller i stort sett generella regler för hur datamaterialet bör vara upplagt. Det är därför klokt att redan från början använda dessa regler.

Materialet läses in i en matris, där raderna ("horisontellt") är observationer och kolumnerna ("vertikalt") är värdena på en av variablerna. Detta sätt att lägga in observationerna gör det enkelt att lägga till nya observationer efterhand som de kommer in.

	Variabel 1	Variabel 2	Variabel 3	...
	↓	↓	↓	
Observation 1 →				...
Observation 2 →				...
Observation 3 →				...
⋮	⋮	⋮	⋮	⋮

Exempel 2. Skolklass

Ur en skolklass har samlats in ett stickprov med fem elever där namn, kön, ålder och vikt har registrerats. Datamaterialet kan beskrivas av en matris med fem rader och fyra kolumner, och där ytterligare en rad har lagts in för att ange variabelnamnen.

NAMN	KÖN	ÅLDER	VIKT
Lisa	F	6	20
Stina	F	7	25
Lars	M	6	17
Peter	M	8	30
Anna	F	7	27

(Här har de svenska variabelnamnen använts, men de engelska är ibland att föredra eftersom vissa program inte accepterar å, ä och ö.) ■

När försöken upprepas kan också en eller flera variabler användas till att dela in observationerna i olika grupper. I följande exempel är behandlingen och veckan variabler som delar datamaterialet i grupper.

Exempel 3. Purjolök

Vid ett försök som syftar till att undersöka hur kväveinnehållet förändras i purjolök har man gjort ett försök där man jämför ett ogödslat fält med ett som innehållit perserklöver/havre. Purjolökens torrsubstansvikt i kg/ha samt kväveinnehållet i procent av torrsubstansen har mätts på tre platser per behandling och mätningen har upprepats 7 och 11 veckor efter planteringen.

För att förenkla inläsningen kvantifieras först behandlingen enligt

$$\begin{cases} 0 = \text{ogödslat} \\ 1 = \text{perserklöver / havre} \end{cases}$$

För att sedan få den rätta överblicken och lätt kunna få ut resultat ur ett statistiskt programpaket bör försöksresultatet skrivas enligt tabellen.

Om man vill läsa in nya variabler från ett annat problem så kan man fortsätta på samma ark och lägga in dem från C2 om man redan har något i C1, men man kan också skapa ett nytt datafönster genom att välja

File ► New ...

för att ha bättre ordning på de olika försöken. Här skapar vi ett nytt datafönster som får namnet Exempel 3.

	C1	C2	C3	C4	C5
	week	treat	dry	N	
1	7	0	232	2.63	
2	7	0	161	2.90	
3	7	0	148	2.99	
4	7	1	368	3.54	
5	7	1	218	3.30	
6	7	1	257	2.85	
7	11	0	1633	1.53	
8	11	0	2213	1.90	
9	11	0	972	*	
10	11	1	2560	2.58	
11	11	1	2430	*	
12	11	1	855	*	
13					
14					

Saknade värden har här betecknats med en asterisk, och detta är "missing value code" i Minitab. I andra program används andra tecken, SAS använder t.ex. punkt. Det är ibland de saknade värdena som försvårar överföringar mellan olika programpaket. ■

3. Minitab tillsammans med Word och Excel

3.1. Från Excel till Minitab

För att flytta datamaterial från Excel till Minitab är det viktigt att komma ihåg hur data skall läsas in. Gör man detta som i föregående avsnitt, har man ett rektangulärt schema där observationerna är i rader och variablerna i kolumner. Det är viktigt att komma ihåg att *inte* Excelarket får innehålla insprängda delmedelvärden och delstandardavvikelser; det är alltför många som har upptäckt att det är ett stort arbete att få bort dessa.

Det är ofta att föredra att man läser in större datamaterialet i Excel eftersom detta program för det mesta har något bättre editeringsmöjligheter.

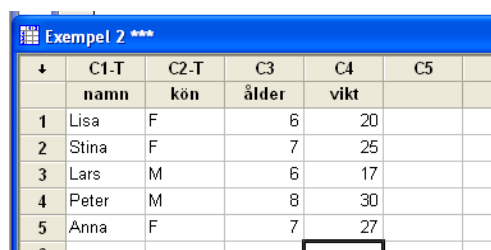
Exempel 2 (forts). Skolklass

För att läsa in data på skolbarnen i Excel gör man ett blad enligt

	A	B	C	D	E
1	namn	kön	ålder	vikt	
2	Lisa	F	6	20	
3	Stina	F	7	25	
4	Lars	M	6	17	
5	Peter	M	8	30	
6	Anna	F	7	27	
7					
8					

Man kan nu välja att öppna ett sparat Excelark direkt i Minitab, men är det små datamängder är det oftast lättare att klippa ut och klistra in det direkt i Minitab.

För att klistra in markeras alla observationerna, inklusive variabelnamnen och man väljer kopiera. Gå över till Minitab och ställ dig i den rad som skall innehålla variabelnamnen (här lägger vi in också dessa data i ett nytt datafönster). Klistra nu in datamaterialet i datafönstret. De variabler som inte innehåller siffror blir text-kolumner och får namnet C1-T resp C2-T.



Exempel 2 ***					
↓	C1-T	C2-T	C3	C4	C5
	namn	kön	ålder	vikt	
1	Lisa	F	6	20	
2	Stina	F	7	25	
3	Lars	M	6	17	
4	Peter	M	8	30	
5	Anna	F	7	27	
6					

Nu kan man arbeta vidare med detta material i Minitab. ■

Lägg speciellt märke till att ett väldigt vanligt skäl till att man helt plötsligt har fått en kolumn betecknad med "T" kan bero på att man har matat in decimalkomma istället för decimalpunkt (eller tvärtom beroende på hur programmet är inställt).

3.2. Från Minitab till Word

De resultat som erhålls i resultatfönstret eller i en graf kan skrivas ut direkt på skrivaren genom att man ställer sig i det fönster man vill skriva ut och väljer att skriva ut det. Då får man utskriften i A4-format, och ofta vill man lägga in resultaten i ett dokument så som har gjorts i exemplet ovan.

För att överföra en del av ett resultatfönster till Word markerar man den text som skall flyttas och kopierar eller klipper ut den på vanligt sätt. Sedan går man över till sitt Word-dokument utan att stänga Minitab och klistrar in texten.

Det är en bra övning att ta det tidigare resultatet, flytta över det i Word och se att texten blir likadan som i Minitab. Word behåller alltså de formateringar som finns i resultatfönstret².

Fönster med figurer ("Graph Window") är enklare att kopiera, det är bara att flytta sig till det grafiska fönster man vill kopiera och sedan gå över till Word och klistra in det. Här kan man det ofta vara bra att välja "Klistra in special..." och sedan "Bild" om man inte vill ha kvar kopplingen till Minitab.

3.3. Läs in data i Minitab

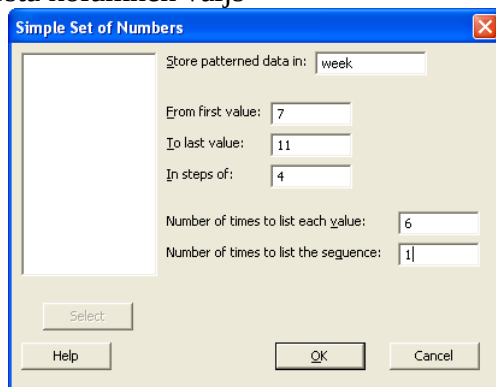
Det finns ett tillfälle då Minitab kan vara att föredra när man skall läsa in data, och detta är när man har värden som upprepas enligt ett regelbundet mönster. Som exempel väljer vi det tidigare exemplet med kväveinnehållet i purjolök.

Exempel 3 (forts). Purjolök

Här skall vi läsa in en kolumn med sex stycken 7 följt av sex stycken 11 samt en kolumn med tre 0 följt av tre 1 upprepat två gånger. Detta görs genom att välja

Calc ► Make Patterned Data ► Simple Set of Numbers...

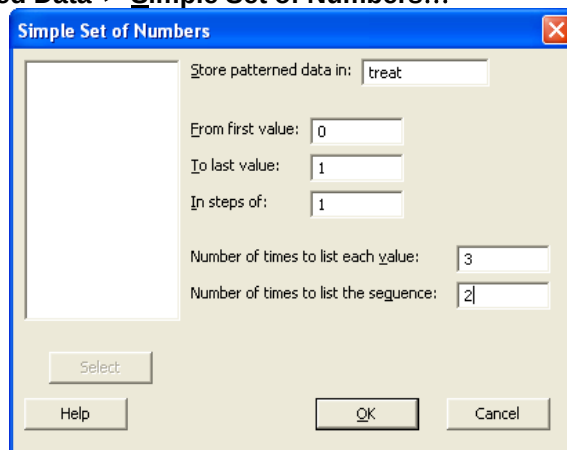
För att åstadkomma den första kolumnen väljs



² Resultatfönstret använder teckensnittet Courier som tar lika stor plats för alla bokstäver. Detta gör att man inte behöver använda tabbar utan det räcker att räkna blanksteg för att få raka marginaler.

För att göra kolumnen med 0 och 1 upprepad två gånger väljs återigen

Calc ► Make Patterned Data ► Simple Set of Numbers...



(Tänk själv efter (eller prova) att detta ger den önskade kolumnen.) ■

3.4. Editera Session Window i Minitab

Om man bara vill ha ut ett enskilt resultat ur resultatfönstret upptäcker man att resultaten läggs efter varandra då man gör flera analyser. Det är dock lätt att radera de utskrifter man inte vill ha ut på papperet. Om man vill undvika att av misstag sudda ut saker i fönstret kan man gå till "read-only" genom att markera

Editor ► Output Editable

när man har markören i resultatfönstret³.

3.5. Editera Worksheet i Minitab

Ändra namn på Worksheet

Det är väldigt bra att kunna ge namn på olika datafönster för att lättare hitta när man söker efter ett fönster under menyn Window. Detta gör man genom att klicka på Project Manager och sedan ge namn genom att klicka i filstrukturen på samma sätt som man gör i Windows för övrigt.

Kopiera delar av kolumner

Ibland vill man kanske ha ut en del av en kolumn. I exemplet med purjolök kanske man vill undersöka de lökar som är sju veckor gamla.

Exempel 3 (forts). Purjolök

För att kopiera de skörderesultat som hör till sju veckor gamla lökar till ett nytt datafönster väljer man

Data ► Subset Worksheet...

³ Läggs märke till att innehållet och menyerna förändras beroende på vilken typ av fönster man är i.

och fyller i de värden som man vill ha. Här finns ett alternativ som heter Brushed rows, och med detta menas att man i en bild kan markera de observationer som man tycker skall vara med. Man kan också använda

Data ► Copy ► Columns to Columns...

för att kopiera delar av datamaterial till nya kolumner eller datafönster. ■

Använda matematiska funktioner

Ibland vill man transformera observationerna eller använda en funktion som inte finns tillgänglig bland menyerna. Om man till exempel vill logaritmera värdena gör man detta genom att använda

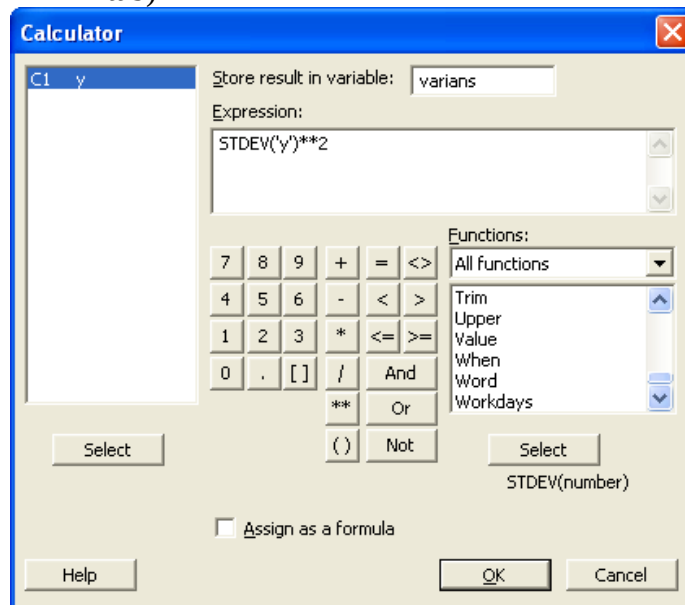
Calc ► Calculator ...

Exempel 1 (forts). Inledande exempel

I detta exempel har vi tidigare beräknat medelvärde och standardavvikelse, och man kan också få fram variansen. För att exemplifiera använder vi trots allt här de matematiska funktionerna för att beräkna denna. Välj

Calc ► Calculator ...

och om svaret skall skrivas i en kolumn som får namnet varians så fyller du i enligt (** betecknar "upphöjt till" i Minitab)



Lägg märke till att man kan skriva det variabelnamn som man vill ha direkt i rutan, men lägg också märke till att Minitab skriver över utan att varna om man väljer ett variabelnamn som redan finns i datafönstret.

Om man bara vill transformera observationerna i en kolumn eller lägga ihop två kolumner så är detta också enkelt gjort. Med det som har beskrivits ovan borde detta gå att göra utan större besvär. ■

4. Beskrivande statistik

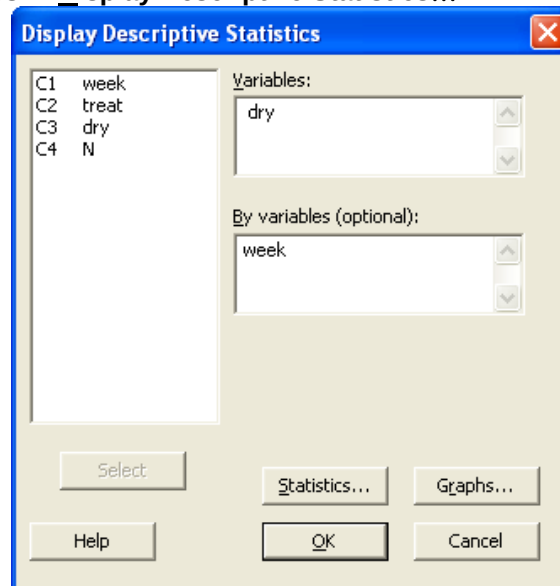
4.1. Numeriska metoder

För att undersöka hur skörden förändras från den sjunde till elfte veckan kan man göra lite numeriska mått för att få en uppfattning om detta. Här bryr vi oss inte om att de olika lökarna har fått olika behandling, syftet är ju bara att exemplifiera.

Exempel 3 (forts). Purjolök

De beskrivande måtten erhålls genom att man väljer

Stat ► Basic Statistics ► Display Descriptive Statistics...



Resultatet blir (Minitab minns vad som gjordes sist, så om inte rutan som begärde "Individual value plot" har avmarkerats så kommer det återigen en sådan och du får samma beskrivande mått som du tidigare valde med **Statistics...**)

Descriptive Statistics: dry

Variable	week	N	Mean	StDev
dry	7	6	230.7	79.3
	11	6	1777	741

Man kunde ju ha valt en hel del andra beskrivande mått. Här följer en förklaring till några av dessa olika mått:

N anger antalet observationer som används i räkningen.

N* anger antalet observationer som har "missing value" (kodade med *), men finns ej här.

Mean är det vanliga aritmetiska medelvärde. Torrvikten är givetvis högre efter elva veckor.

Median är som vanligt det "mittersta" värdet, eller medelvärde av de två mittersta om det är ett jämnt antal värden.

Tr Mean, är "trimmed mean" där man har tagit bort ett antal stora resp. små värden och därefter beräknat medelvärdet. Detta för att inte s.k. "outliers" skall få för stor betydelse i medelvärdesberäkningen. Minitab tar bort 5% av de största och 5% av de minsta värdena när detta medelvärde beräknas.

StDev är den vanliga standardavvikelsen. Det är betydligt större spridning bland de äldre lökarna.

SE Mean är Standard Error eller medelfel. Det är helt enkelt $StDev / \sqrt{N}$ och detta uttryck är i vissa sammanhang mer användbart än standardavvikelsen.

Min och Max är minimum och maximum och behöver inte förklaras ytterligare.

Q1 och Q3 är de bägge kvartilerna. Kvartil innebär att man lägger observationerna i storleksordning och därefter delar dem i fyra delar. De tre punkter som då delar materialet är nedre kvartilen, medianen och övre kvartilen. Kvartilerna illustreras tydligare i avsnittet om Boxplot.

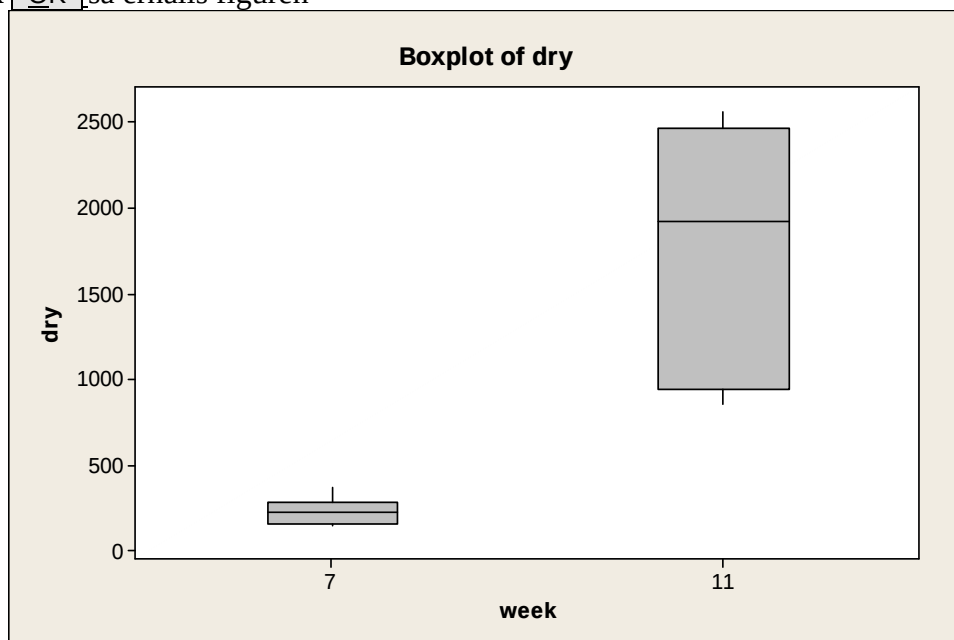
4.2. Boxplot

En mycket illustrativ grafisk presentation av medianen och kvartilerna ges av en s.k. Boxplot (eller Box-and-Whisker-plot, kallas på svenska ibland Lådagram). Denna får man enkelt genom att välja

Stat ► Basic Statistics ► Display Descriptive Statistics...

Välj sedan **Graphs...** och markera rutan "Boxplot of data".

Tryck **OK** så erhålls figuren



I figuren är strecket inne i lådan medianen, den undre lådkanten är nedre kvartilen och den övre lådkanten är övre kvartilen.

4.3. Scatterplot

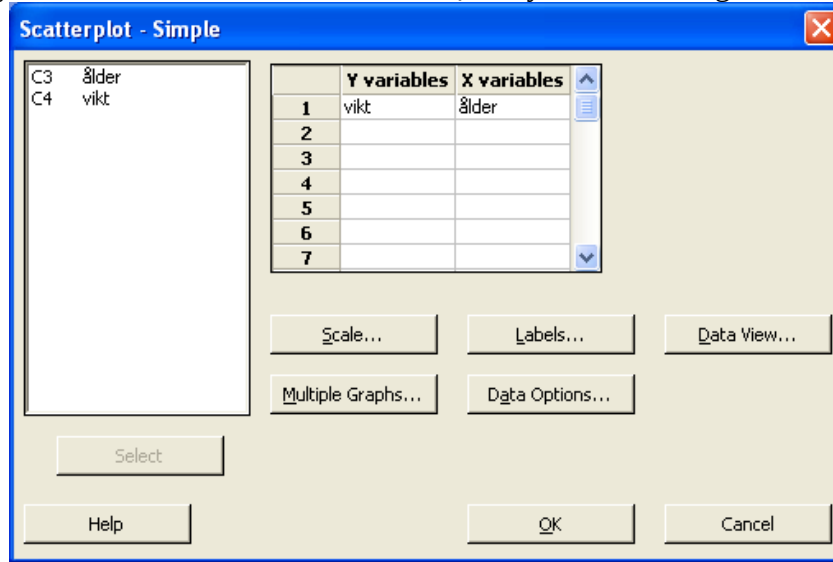
Ibland har man tvådimensionella variabler som man vill rita upp mot varandra för att se om det finns något samband. Den figur man får ur Minitab kan förbättras avsevärt om man har tid och lust till detta. Hur detta går till lämnas dock åt läsaren.

Exempel 2 (forts). Skolklass

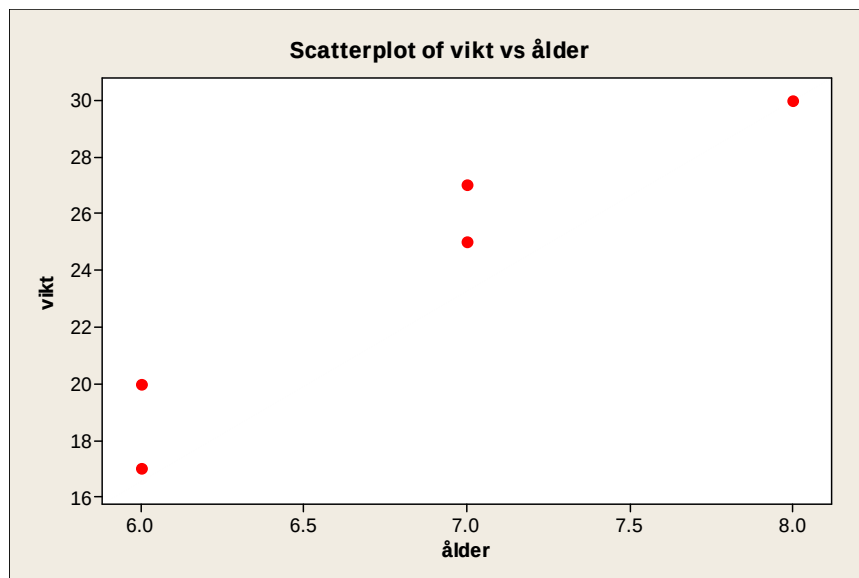
För att rita variabeln vikt mot ålder i en scatterplot (det finns inget bra svenskt namn) används

Graph ► ScatterPlot...

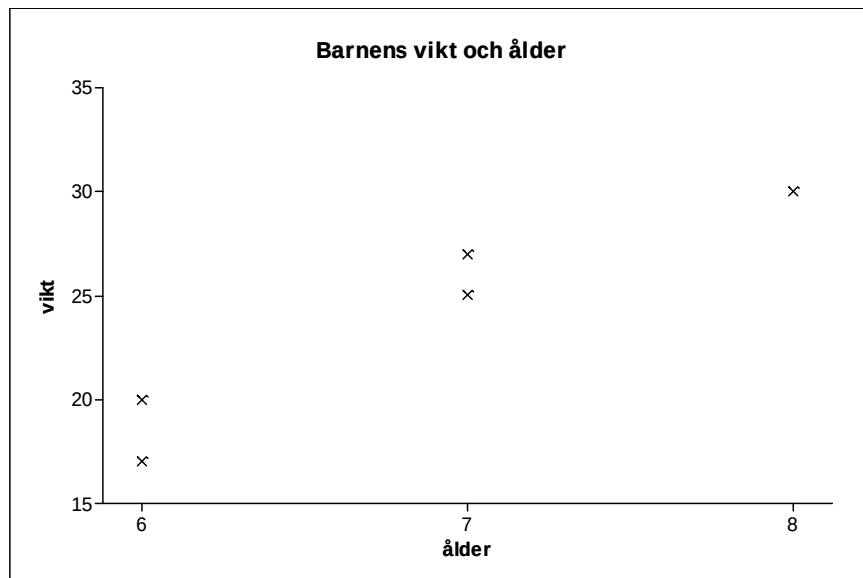
och sedan väljer man den förmarkerade rutan Simple. Fyll sedan i enligt



Tryck på **OK**, och resultatet blir



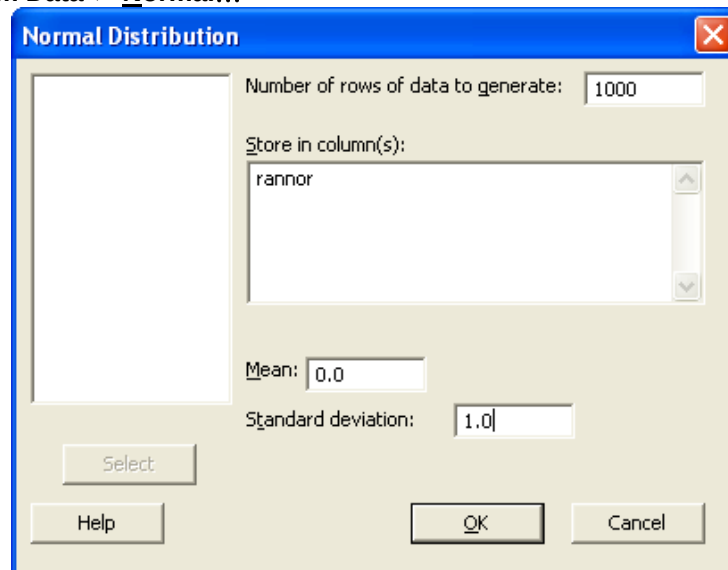
Genom att välja de olika optioner som finns kan man ändra utseende och skalor på figuren, och man kan också editera figuren. Detta lämnas dock till läsaren, men ett exempel på hur figuren kan förändras följer här:



4.4. Histogram

Histogram är lika enkelt som att göra andra figurer, bara man vet var man skall leta! Som ett exempel kontrollerar vi om slumpalsgeneratorn i Minitab ger värden som man kan tro är normalfördelade. Slumptal från normalfördelningen med populationsmedelvärde 0 och standardavvikelse 1 får man genom att välja

Calc ► Random Data ► Normal...



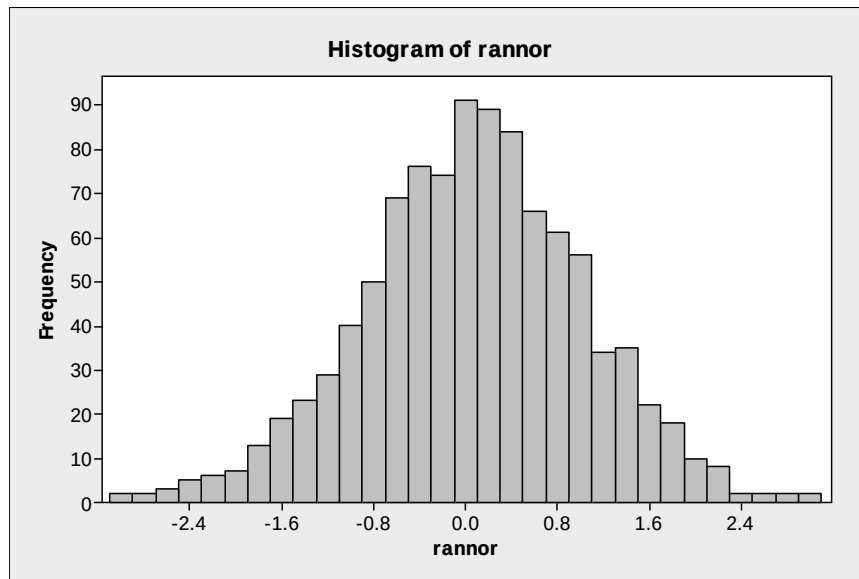
Nu har datorn slumpat 1000 observationer av en normalfördelning med populationsmedelvärde 0 och standardavvikelse 1 och lagt dessa värden i variabeln med namnet `rannor`. För att rita ett histogram över dessa värden kan man som tidigare gå via menyn för deskriptiv statistik men vi väljer här en annan väg som ger fler möjligheter att påverka utseendet på figuren (samma sak gäller för övrigt den Boxplot som vi gjorde tidigare):

Graph ► Histogram...

Och välj Simple. I rutan `Graph variables:` väljs `rannor`.

Tryck på **OK**.

Ett exempel på ett histogram av denna typ finns nedan, men eftersom det är slumpstal som är uppritade kommer utseendet att variera från gång till gång. En viss klockform borde man ju dock alltid få.



5. Statistiska metoder för ett stickprov

I de allra flesta fall skall man vid analysen jämföra två eller flera olika behandlingar. Ibland kan man dock ha ett stickprov som man vill uttala sig om, och det är i den ändan denna beskrivning börjar.

När man skall göra analyserna nedan är det viktigt att värdena utgör ett s.k. slumpmässigt stickprov, dvs. att värdena är oberoende av varandra.

5.1. Normalfördelning och ett stickprov

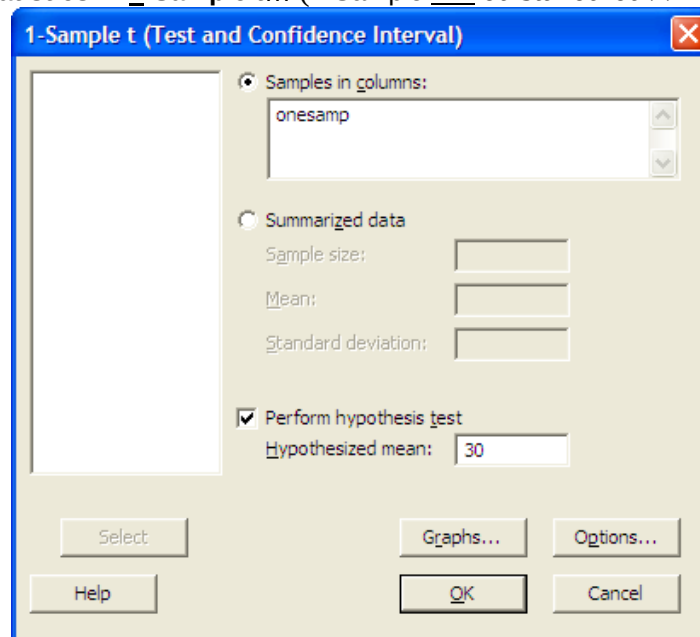
För det mesta antar man att stickprovet kommer från en normalfördelning, eller åtminstone att det är ungefär normalfördelat. Detta innebär att om man samlar in många observationer och ritar ett histogram så skall detta få formen av den berömda klock-kurvan. Baserat på ett stickprov som man antar är normalfördelat kan man antingen göra ett hypotestest eller ett konfidensintervall för populationsmedelvärdet. Detta görs mycket enkelt i Minitab, det enda man skall bestämma är om (populations-)standardavvikelsen är känd eller skall skattas ur det stickprov man har. I praktiken är det ju mycket ovanligt att standardavvikelsen är känd och därför behandlas här endast fallet med okänd standardavvikelse.

Exempel 4. Ett stickprov

Testa om stickprovet med värdena 33.5, 32.0, 32.5, 36.5 kan komma från en population med medelvärdet 30. Dessa värden antas inlästa i en kolumn med namnet `onesamp`.

Nollhypotesen i detta fall är att populationsmedelvärdet är 30 medan mothypotesen är att populationsmedelvärdet inte är 30. Välj

Stat ► Basic Statistics ► 1-Sample t... (1-Sample Z... då standardavvikelsen är känd)



För att eventuellt ändra mothypotesen eller konfidensgraden använder man rutan **Options...**, men i detta fall är det rätt inställt, not equal är ju precis vad man vill ha som mothypotes. Resultatet i resultatfönstret blir

One-Sample T: onesamp

Test of $\mu = 30$ vs not = 30

Variable	N	Mean	StDev	SE Mean	95% CI	T	P
onesamp	4	33.6250	2.0156	1.0078	(30.4178; 36.8322)	3.60	0.037

Det mest intressanta värdet här är P som är 0.037. Detta så kallade *p*-värdet anger om hypotesen kan förkastas eller ej, och den vanligaste "regeln" är att förkasta nollhypotesen om *p*-värdet är mindre än 0.05. I detta fall kan vi alltså förkasta nollhypotesen, vilket på vanlig svenska betyder att vi har anledning att misstänka att populationsmedelvärdet inte är 30. Som vanligt i statistiken kan vi dock inte vara helt säkra; vi har valt nivån så att man i 1 fall av 20 förkastar nollhypotesen även om den är sann. ■

Det man oftast ser är hypotestest som har beskrivits ovan, men hypotestest kritiserar ibland. Man kan då istället använda sig av konfidensintervall. Ett konfidensintervall med 95% konfidensgrad har erhållits automatiskt ovan som (30.42, 36.83), och om man vill ändra konfidensgraden så kan man välja **Options...**. Konfidensgraden 95% säger att konfidensintervallet har egenskapen att det med 95% sannolikhet (dvs. i 19 fall av 20) täcker det rätta värdet på populationsmedelvärdet.

Ju smalare intervallet är, ju bättre är det givetvis. Ett smalt intervall får man antingen om man har många observationer i stickprovet (och det kan man kanske påverka själv) eller genom att spridningen mellan värdena är liten (och det är ju svårt att påverka). I exemplet ovan är intervallet kanske för brett för att vara praktiskt användbart.

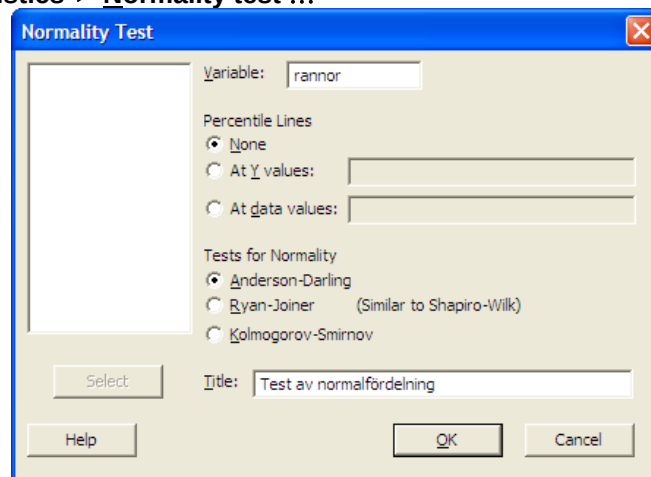
En förutsättning för att testet och konfidensintervallet ovan skall vara giltigt är att populationen är normalfördelad. Ett sätt att kontrollera detta är att använda s.k. normalfördelningspapper. Före datorernas tidevarv var detta något av det mest tidsödande man kunde sysselsätta sig med inom statistiken, nu är det mycket enkelt.

Normalfördelningspapper

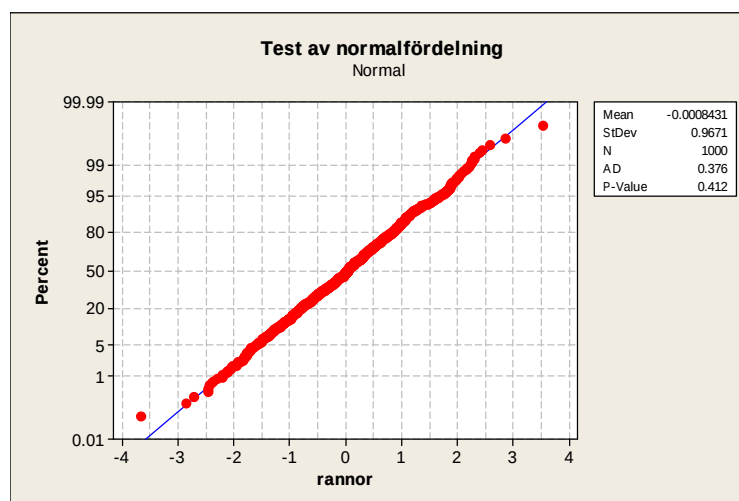
Ett normalfördelningspapper är så listigt ordnat att man har ändrat skalan på y-axeln så att när stickprovet från en normalfördelning ritas upp så skall de utgöra en rät linje. När man sedan gör test om huruvida de ligger intill linjen utnyttjar man asymptotiska resultat, dvs. resultat som bara gäller om man har många observationer. Vi använder det tidigare materialet med 1000 normalfördelade observationer som fanns i samband med histogrammet.

Materialet är redan inläst och finns i variabel `rannor`. Välj

Stat ► Basic Statistics ► Normality test ...



Resultatet i detta fall blir:



P-värdet (P-value) ger en fingervisning om hur man skall ställa sig till hypotesen om normalfördelning hos populationen. Ett värde under 0.05 gör att man förkastar hypotesen att populationen är normalfördelad, och ett värde över 0.05 gör att man inte kan förkasta hypotesen. Detta brukar innebära att man fortsätter arbeta som om populationen var normalfördelad. Lägg dock märke till att man på intet sätt har bevisat detta, man använder bara principen att tro på normalfördelningen tills någon annan bevisar motsatsen. I detta fallet förkastar vi inte hypotesen, och det stämmer ju väl med att data var simulerade från en normalfördelning. ■

Om man nu inte kan anta att populationen är normalfördelad har man två alternativ:

1. Om man har många observationer i stickprovet brukar det fungera i alla fall.
2. Använd en icke-parametrisk metod som inte förutsätter normalfördelad population.

5.2. Ej normalfördelning och ett stickprov

Icke-parametriska test heter på engelska "non-parametric tests" och därför finns de testen under menyn nonparametrics. I fallet med ett stickprov finns det två olika test, och det är värt att notera att icke-parametriska test inte testar medelvärde utan medianen. Att testa om populationen har medianen 30 är detsamma som att undersöka om hälften av populationen har värden under 30 och den andra hälften har värden över 30.

Teckentest

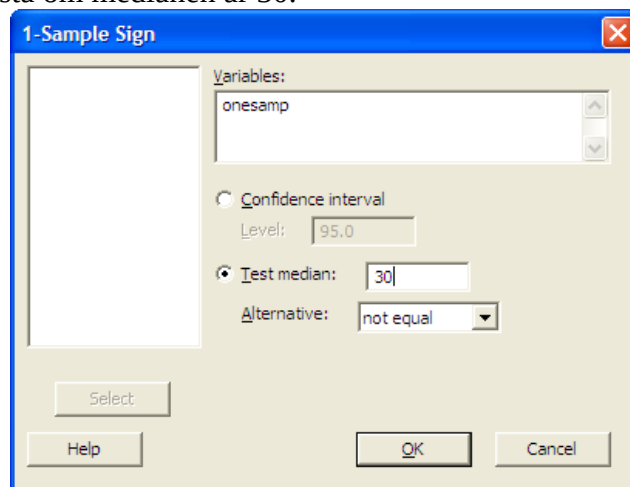
Teckentestet (eng. sign test) är mycket lätt att förstå men tyvärr inte så effektivt. Idén är helt enkelt att om medianen är 30 så borde ungefär hälften av stickprovets värde vara under 30 och resten vara över 30. Har man exempelvis 20 värden i stickprovet och alla 20 är under 30 är det ganska svårt att tro att medianen är 30.

Exempel 4 (forts). Ett stickprov

För att avgöra om medianen är 30 med ett hypotestest, välj

Stat ► Nonparametrics ► 1-Sample Sign...

och fyll i att du vill testa om medianen är 30.



Utskriften i resultatfönstret blir

Sign Test for Median: onesamp

Sign test of median = 30.00 versus not = 30.00

	N	Below	Equal	Above	P	Median
onesamp	4	0	0	4	0.1250	33.00

Trots att alla värdena är över 30 kan man inte förkasta hypotesen (p -värdet 0.1250 är ju större än 0.05). Teckentestet är alltså inte tillräckligt effektivt i detta fall. ■

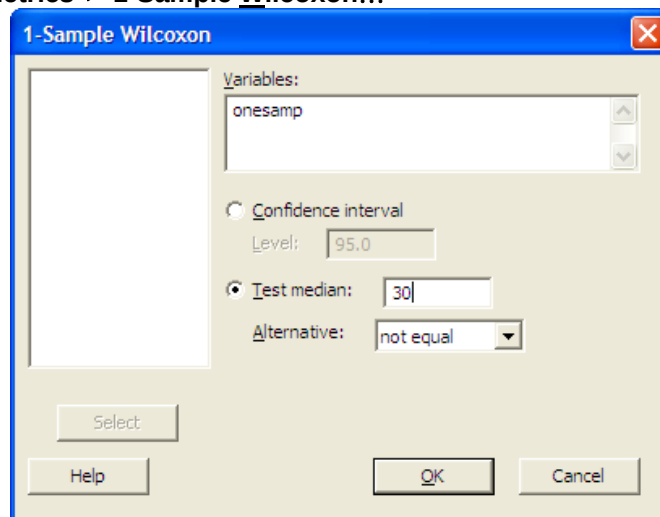
Wilcoxon's tecken-rang-test

Detta test är mer effektivt eftersom det också utnyttjar storleken på värdena, men man har här också kravet att populationen skall ha en symmetrisk fördelning. Detta leder ju dessutom till att medianen och populationsmedelvärdet sammanfaller.

Exempel 4 (forts). Ett stickprov

För att göra testa om medianen är 30, välj

Stat ► Nonparametrics ► 1-Sample Wilcoxon...



Wilcoxon Signed Rank Test: onesamp

Test of median = 30.00 versus median not = 30.00

	N	for	Wilcoxon	Estimated	
	N	Test	Statistic	P	Median
onesamp	4	4	10.0	0.100	33.25

Inte heller här kan man förkasta hypotesen även om p -värdet här är något mindre än i teckentestet. ■

Exemplet har gett det resultat man brukar få. Om man kan, skall man använda testet som bygger på normalfördelning eftersom det är mest effektivt.

Det finns ytterligare ett skäl till att ibland använda ett icke-parametriskt test. Skälet är att detta test är mer robust mot "dåliga värden". Om man av misstag har råkat läsa in 365 istället för 36.5 så kommer p-värdet vid normalfördelningsantagandet att ändras till 0.38. Det kanske förvånar att testet nu inte längre kunde förkasta nollhypotesen när man flyttade ett värde ännu längre ifrån 30, men anledningen är att man genom att höja värdet också ökade på den skattade standardavvikelsen ordentligt. Gör man teckentestet eller Wilcoxon's test med detta värde kommer inte p-värdet att förändras alls!

6. Statistiska metoder för två stickprov

Som nämnts tidigare är det vanligaste att man med en statistisk metod skall jämföra två eller flera populationer. I det förra avsnittet var ett av huvudresultaten att om man antog att det var en normalfördelad population så kunde man använda ett "bättre" test än om man inte kunde göra detta antagande. I fallet med två stickprov gäller precis samma sak, om man har normalfördelning hamnar man i det berömda t -testet.

6.1. Normalfördelning och två stickprov

Om man antar att de två populationer man skall jämföra är normalfördelade, finns det bara två parametrar som skiljer populationerna åt, populationsmedelvärdet och standardavvikelsen. Saker och ting blir ännu trevligare om man dessutom kan anta att standardavvikelseerna är desamma i de båda populationerna. Att testa om populationsmedelvärdena är desamma är ju då detsamma som att testa om de två stickproven är hämtade från precis samma population. Det finns också ett annat skäl till att man ofta antar att standardavvikelseerna är desamma: det blir teoretiskt mycket enklare att räkna i den situationen än om man har olika standardavvikelser. I ett datorprogram spelar det ingen roll om man har lika eller olika standardavvikelse, men man skall ha i minnet att lika standardavvikelser och normalfördelning innebär att man testar om de båda populationerna är lika.

Man kan också göra konfidensintervall för skillnaden mellan populationernas populationsmedelvärden. Här brukar man främst titta om intervallet "omsluter" 0, i så fall är nämligen inte populationsmedelvärdena signifikant skilda åt.

Exempel 5. Två stickprov

Testa om stickproven 33.5, 32.0, 32.5, 36.5 respektive 39.5, 36.0, 34.5, 36.5 kan komma från populationer med samma populationsmedelvärde.

Nollhypotesen i detta fall är att populationsmedelvärdena är desamma, mothypotesen är att populationsmedelvärdet är olika. Välj

Stat ► Basic Statistics ► 2-Sample t...

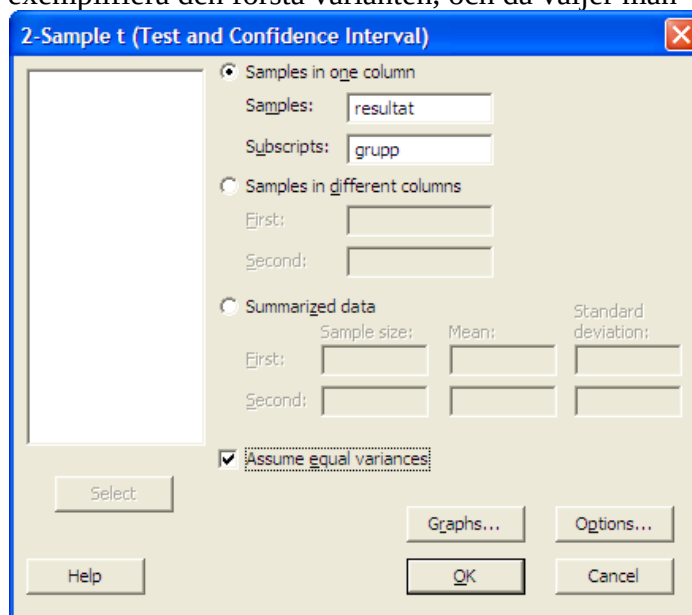
Nu hamnar man i en valsituation, antingen skall man markera att värdena är inlästa i en kolumn (Samples in one column) eller att värdena är inlästa i två kolumner (Samples in different

columns). Om man läser in data enligt de regler som angivits i tidigare kapitel så har man en kolumn, men man har också en kolumn som anger vilket stickprov den mätningen kommer ifrån. Datamaterialet ser då ut som till vänster, medan olika kolumner blir som till höger.

grupp	result
1	33.5
1	32.0
1	32.5
1	36.5
2	39.5
2	36.0
2	34.5
2	36.5

grupp 1	grupp 2
33.5	39.5
32.0	36.0
32.5	34.5
36.5	36.5

Här väljer vi att exemplifiera den första varianten, och då väljer man



Den alternativa hypotesen liksom nivån kan ändras genom att välja **Options...**. I den nedersta rutan kan man sedan välja om man vill att standardavvikelserna (varianserna) skall vara lika eller inte. Om man har valt att anta lika varians blir utskriften

Two-Sample T-Test and CI: result; grupp

Two-sample T for result

grupp	N	Mean	StDev	SE Mean
1	4	33.63	2.02	1.0
2	4	36.63	2.10	1.0

Difference = $\mu(1) - \mu(2)$
 Estimate for difference: -3.00000
 95% CI for difference: (-6.55820; 0.55820)
 T-Test of difference = 0 (vs not =): T-Value = -2.06 P-Value = 0.085 DF = 6
 Both use Pooled StDev = 2.0565

De "svåra" delarna av utskriften är raden med T-test. Det mest intressanta brukar vara p -värdet 0.085, som visar att man *inte* kan förkasta hypotesen att stickproven har samma populationsmedelvärde. Värdet på T är en teststorhet som inte har så stor betydelse och man har inte heller någon större nytta av DF. Eftersom man har antagit att varianserna (och därmed standardavvikelse) är lika⁴ så får man också en skattning av denna gemensamma standardavvikelse, och i detta fall är den 2.06. Här blir den precis medelvärdet av de båda stickprovens standardavvikelser, men det är bara en slump. Den "poolade" standardavvikelsen beräknas enligt en något mer komplicerad formel. Lägg också märke till att konfidensintervallet täcker 0, vilket är ytterligare en indikation på att hypotesen om lika populationsmedelvärden inte kan förkastas. ■

Det går också att göra ett test för om varianserna verkligen är lika, och detta behandlas i det följande avsnittet.

6.2. Test av lika varians

Exempel 5 (forts). Två stickprov

Testa om stickproven 33.5, 32.0, 32.5, 36.5 respektive 39.5, 36.0, 34.5, 36.5 kan komma från populationer med samma varians.

Nollhypotesen i detta fallet är att varianserna är desamma, mothypotesen är att varianserna är olika. Välj

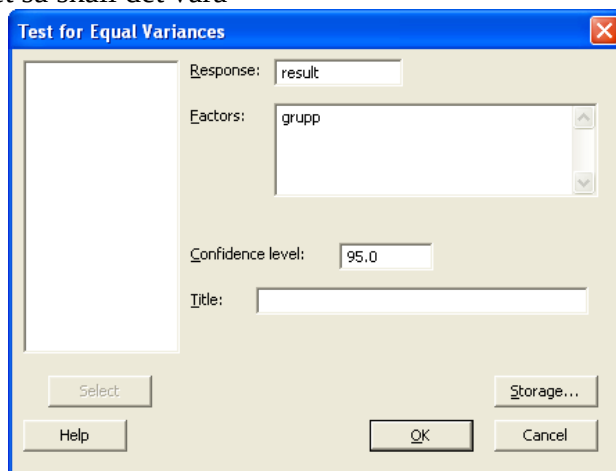
Stat ► ANOVA ► Test for Equal Variances...

(Har man bara två olika grupper kan man också använda

Stat ► Basic Statistics ► 2 Variances...

men det fungerar inte med fler grupper.)

Med mer än två grupper måste nu värdena vara inlästa i en enda lång kolumn och samtidigt ha en kolumn som anger vilket stickprov mätningen kommer ifrån. Med **grupp** och **result** som tidigare i exemplet så skall det vara



⁴ Det är lite konstigt att Minitab har förinställt så att man *inte* antar att varianserna är lika; det vanligaste är nämligen att man har antagit att varianserna är lika när man gör ett t -test.

När man kör proceduren får man dels en utskrift i resultatfönstret, dels en figur. Resultatfönstret blir

Test for Equal Variances: result versus grupp

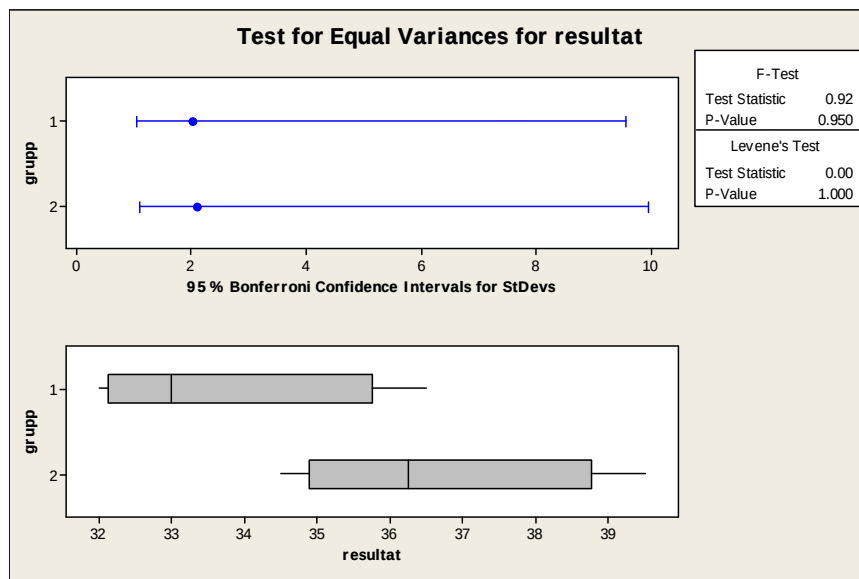
95% Bonferroni confidence intervals for standard deviations

grupp	N	Lower	StDev	Upper
1	4	1.05929	2.01556	9.54568
2	4	1.10190	2.09662	9.92958

F-Test (normal distribution)
Test statistic = 0.92; p-value = 0.950

Levene's Test (any continuous distribution)
Test statistic = 0.00; p-value = 1.000

Här får man två p -värden, ett som är ett F-test (med mer än två grupper så kallas detta Bartlett's test) som man använder vid normalfördelning, och ett som heter Levene's test och kan acceptera lite avvikelser från normalfördelning. Båda värdena är klart över 0.05, och man kan därför *inte* förkasta nollhypotesen att varianserna (standardavvikelsena) är lika. Figuren ger också en grafisk illustration av varianserna i de olika grupperna.



Lägg märke till att detta test inte är begränsat till två stickprov utan kan användas också då man jämför mer än två stickprov. Då ersätts emellertid F-testet med Bartlett's test, men detta behandlas inte i avsnittet med mer än två stickprov. ■

6.3. Ej normalfördelning och två stickprov

Det test man är hänvisad till i Minitab kallas Mann-Whitney's test. När man skall genomföra detta testet upptäcker man att Minitab är inkonsekvent när det gäller inläsning av data. Här skall man nämligen läsa in de två stickproven i olika kolumner, metoden med en kolumn fun-

gerar inte här! Har man läst in data i en enda lång kolumn kan man dock ganska enkelt dela upp det i två kolumner genom att använda

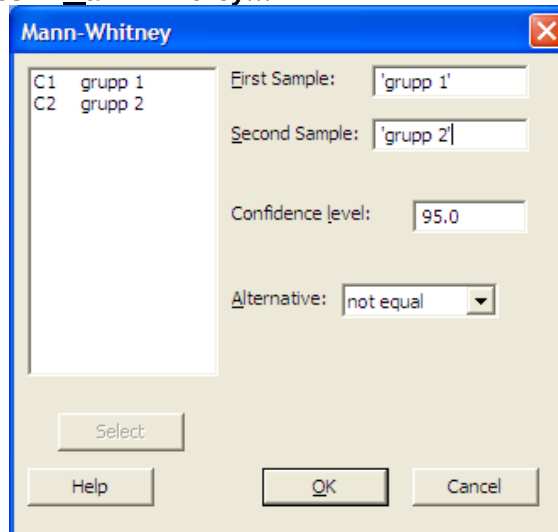
Data ► Unstack Columns...

Exempel 5 (forts). Två stickprov

Testa om stickproven 33.5, 32.0, 32.5, 36.5 respektive 39.5, 36.0, 34.5, 36.5 kan komma från populationer med samma median.

Nollhypotesen är i detta fall att medianerna är desamma, mothypotesen är att medianerna är olika. Välj

Stat ► Nonparametrics ► Mann-Whitney...



Ändra eventuellt hypoteser och signifikansnivå.

Mann-Whitney Test and CI: grupp 1; grupp 2

	N	Median
grupp 1	4	33.000
grupp 2	4	36.250

Point estimate for ETA1-ETA2 is -3.000
97.0 Percent CI for ETA1-ETA2 is (-7.501;2.001)
W = 12.5
Test of ETA1 = ETA2 vs ETA1 not = ETA2 is significant at 0.1489
The test is significant at 0.1465 (adjusted for ties)

Här får man ut två stycken p -värden, 0.1489 resp 0.1465 (adjusted for ties). Det senare betyder att man har modifierat testet något eftersom det innehåller två eller flera värden som är lika (= "ties"); det finns ju två värden som är 36.5. Här får man det lite konstiga påståendet att testet är signifikant på nivån 0.1465, men det innebär ju att man inte kan påstå att det finns någon skillnad mellan populationernas medianer eftersom $0.1465 > 0.05$. ■

Precis som fallet var vid ett stickprov är det så att man förlorar lite om man använder Mann-Whitney's test om man i själva verket har förutsättningar att göra ett t -test.

7. Statistiska metoder vid flera stickprov

Den allra vanligaste metoden i statistiken är att göra en variansanalys. Detta kallas ofta ANOVA efter engelskans Analysis of Variance. Variansanalysen i sin "standardform" förutsätter emellertid helst både normalfördelning och att det är samma varians i alla populationerna. Det är också viktigt att komma ihåg att det trots namnet inte är varianser man testat, det är populationsmedelvärdena som testas genom att man jämför variansskattningar.

7.1. Normalfördelning och flera stickprov

Variansanalys är en utveckling av *t*-testet, där man istället för två stickprov har fler stickprov. Variansanalysen kan sedan utvecklas i en mängd olika riktningar, men här behandlas bara de mest grundläggande problemen.

Exempel 6. Flera stickprov

Man har fyra stickprov, och man vill undersöka om dessa har samma populationsmedelvärde. Det finns skäl att tro att alla stickproven kommer från normalfördelade populationer med samma standardavvikelse för alla populationerna. Nollhypotesen om lika populationsmedelvärde är därför detsamma som att testa om alla observationerna är tagna ur samma population.

Stickprov 1: 33.5, 32.0, 32.5, 36.5

Stickprov 2: 39.5, 36.0, 34.5, 36.5

Stickprov 3: 32.0, 32.5, 33.5, 34.5

Stickprov 4: 36.0, 35.5, 33.5, 37.0

Precis som vid två stickprov kan man välja om man vill läsa in värdena i en kolumn, med en kolumn bredvid som anger stickproven, eller att läsa in stickproven i fyra olika kolumner.

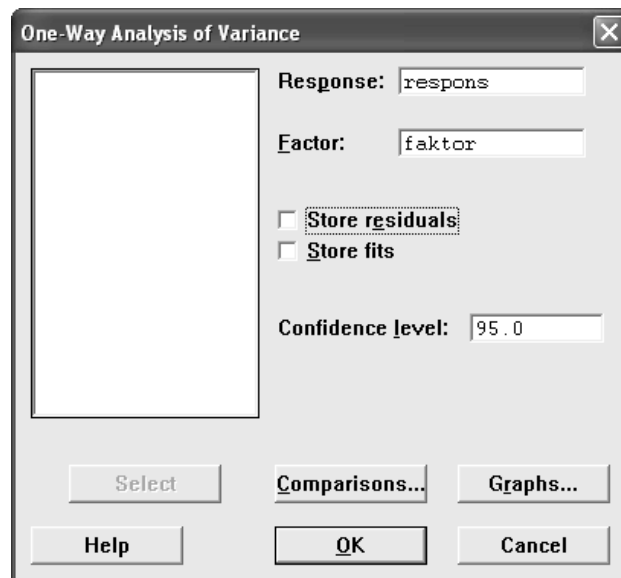
Om man har läst in i fyra kolumner skall man välja

Stat ► ANOVA ► Oneway (Unstacked) ...,

medan man skall välja

Stat ► ANOVA ► Oneway...

om man har läst in alla i en enda lång kolumn. Om man väljer varianten med en enda lång kolumn kan man utvidga det när man får mer avancerade modeller, och denna metod är därför att föredra. I detta exempel har faktorerna lagts i variabeln `faktor` och responsen i variabeln `respons`.



Resultatet som kommer ut i resultatfönstret blir

One-way ANOVA: respons versus faktor

Source	DF	SS	MS	F	P
faktor	3	31.92	10.64	3.59	0.046
Error	12	35.56	2.96		
Total	15	67.48			

S = 1.721 R-Sq = 47.30% R-Sq(adj) = 34.13%

				Individual 95% CIs For Mean Based on Pooled StDev	
Level	N	Mean	StDev		
1	4	33.625	2.016	(-----*-----)	
2	4	36.625	2.097	(-----*-----)	
3	4	33.125	1.109	(-----*-----)	
4	4	35.500	1.472	(-----*-----)	
				-----+-----+-----+-----	
				32.0 34.0 36.0 38.0	

Pooled StDev = 1.721

Den siffra som är viktigast är det p -värde 0.046 som visar att nollhypotesen om att stickproven kommer från samma population kan förkastas. Man får också konfidensintervall för varje enskilt stickprov utritade. Lägg dock märke till att detta inte är samma konfidensintervall som man får om man bara tittar på ett stickprov åt gången, eftersom man här har använt den gemensamma standardavvikelsen när man har gjort intervallen. ■

För att avgöra vilka nivåer som skiljer sig åt kan man använda **Comparisons**, och välja vilken typ av jämförelse man skall göra. Då får man bokstäver som visar vilka nivåer som skiljer sig åt men Minitab producerar inte det så kallade LSD-värdet.

Det nämndes tidigare i avsnittet med två stickprov hur man testar om varianserna i stickproven är lika, och den metoden går att använda även på fler stickprov.

För att se om det finns några konstiga värden i datamaterialet är det lättast att i menyn välja knappen **Graphs...** och sedan Four in one. Detta ger fyra olika bilder av residualerna⁵ som var och en på sitt sätt säger om antagandena för analysen är uppfyllda.

7.2. Ej normalfördelning och flera stickprov

Om man inte tror på normalfördelningen och vill göra ett icke-parametriskt test, så heter det test som motsvarar den ensidiga variansanalysen för Kruskal-Wallis test. Om man vet hur man gör ANOVA så är det enkelt att göra Kruskal-Wallis test på samma sätt. Välj

Stat ► Nonparametrics ► Kruskal-Wallis...

och fyll i de två rutorna. Detta ger

Kruskal-Wallis Test: respons versus faktor

Kruskal-Wallis Test on respons

faktor	N	Median	Ave Rank	Z
1	4	33.00	6.1	-1.15
2	4	36.25	12.4	1.88
3	4	33.00	4.9	-1.76
4	4	35.75	10.6	1.03
Overall	16		8.5	

H = 6.76 DF = 3 P = 0.080

H = 6.85 DF = 3 P = 0.077 (adjusted for ties)

* NOTE * One or more small samples

För att man skall kunna lita på Kruskal-Wallis test måste det vara ganska många värden i varje stickprov, och så är inte fallet här. Därför får man varningen på sista raden. Här är också två p -värden, där det andra har tagit hänsyn till att några värden är lika. Men p -värdet är således inte helt pålitligt med så här få observationer. ■

8. Blockförsök

Om man har ett blockförsök med normalfördelade observationer så kan man inte använda menyn för envägs ANOVA utan måste använda den något mer avancerade proceduren General Linear Model... Man måste också ha läst in sina värden i en kolumn på det sätt som angavs i början av detta kompendium för att det skall fungera.

Exempel 7. Blockförsök

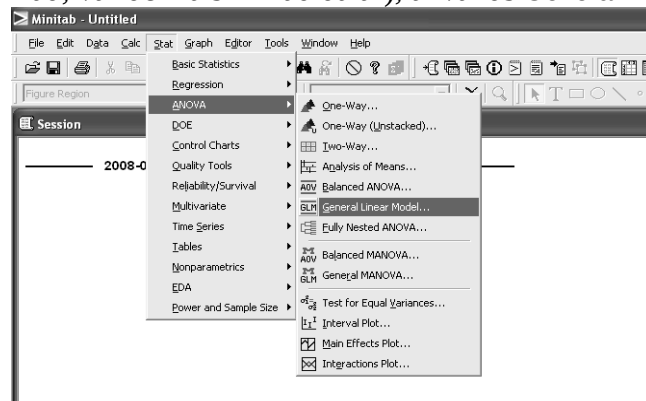
Kvävekoncentrationen i marken mättes på jordprover två veckor efter det att fyra organiska gödningsmedel (A, B, C eller D) hade tillsatts. Experimentet var upplagt som ett blockförsök med tre block (I, II, III). Resultat:

⁵ Det är viktigt att komma ihåg att man inte skall kontrollera de ursprungliga observationerna utan istället residualerna som kan beskrivas som det som är kvar att förklara när man har tagit bort det som förklaras av modellen.

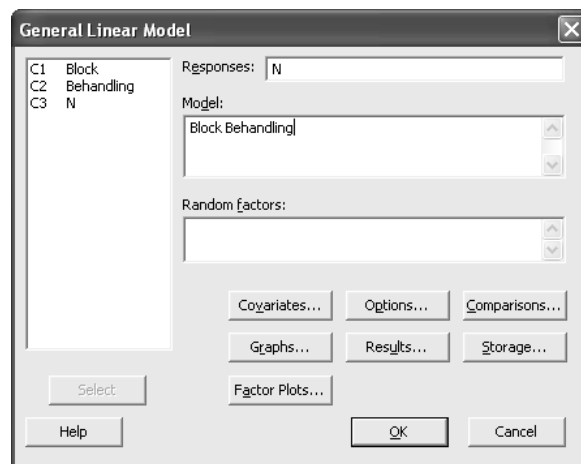
Block	Behandling	N
I	A	3.85
I	B	10.70
I	C	8.60
I	D	10.95
II	A	4.25
II	B	11.85
II	C	9.35
II	D	11.80
III	A	5.60
III	B	9.15
III	C	8.40
III	D	7.55

Materialet läses in på vanligt sätt i Minitab och då kommer **Block** och **Behandling** att bli textkolumner, men det fungerar i ett blockförsök.

För att nu göra en variansanalys för att se om det är någon skillnad mellan behandlingarna (och om det är en skillnad, var denna skillnad då är), används General Linear Model...



Menyn fylls i enligt



och om man bara vill undersöka om det är någon skillnad utan att få reda på var skillnaden är, och utan att rita upp datamaterialet, så räcker det att fylla i menyn på detta viset.

Resultatet blir

General Linear Model: N versus Block; Behandling

Factor	Type	Levels	Values
Block	fixed	3	I; II; III
Behandling	fixed	4	A; B; C; D

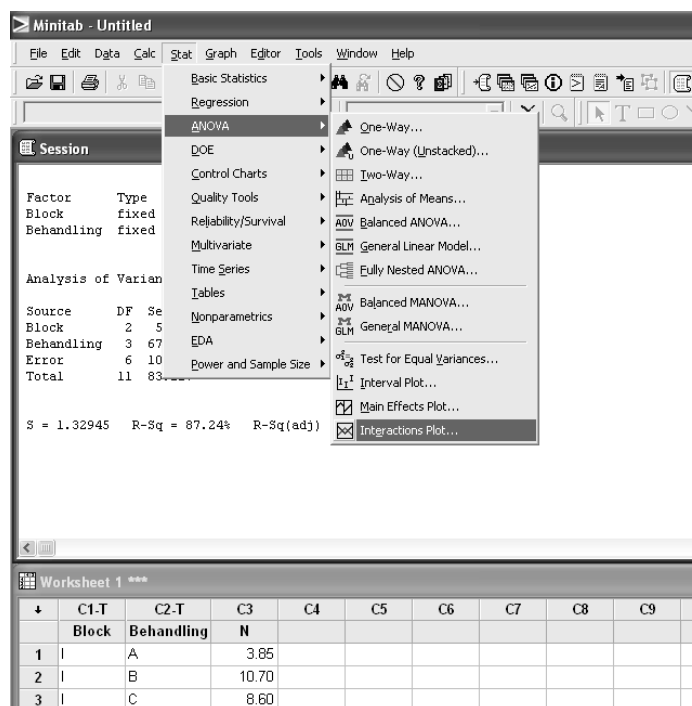
Analysis of Variance for N, using Adjusted SS for Tests

Source	DF	Seq SS	Adj SS	Adj MS	F	P
Block	2	5.365	5.365	2.683	1.52	0.293
Behandling	3	67.147	67.147	22.382	12.66	0.005
Error	6	10.605	10.605	1.767		
Total	11	83.117				

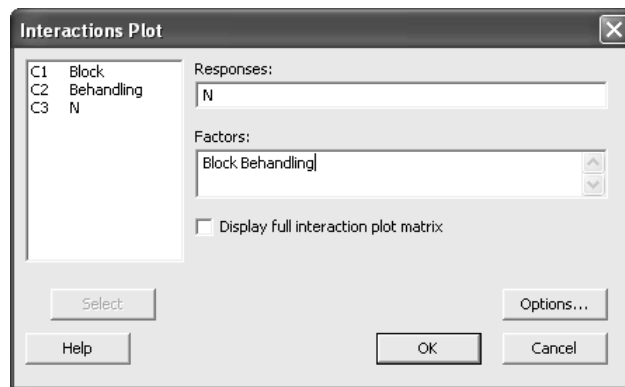
S = 1.32945 R-Sq = 87.24% R-Sq(adj) = 76.61%

Ur detta kan man utläsa att det är en tydlig skillnad mellan behandlingarna (p -värdet är 0.005 som gör att det är signifikant på nivån 0.01). Ur utskriften kan man också se att det inte var någon stor skillnad mellan blocken.

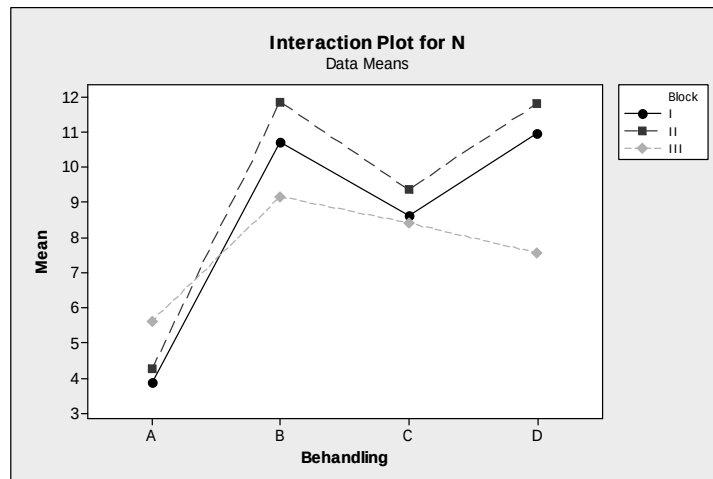
Om man vill ha en bild av datamaterialet så är det mest lämpliga vid blockförsök att rita en så kallad interaction plot som man hittar längst ned i ANOVA-menyn:



Det som skall fyllas i är nu

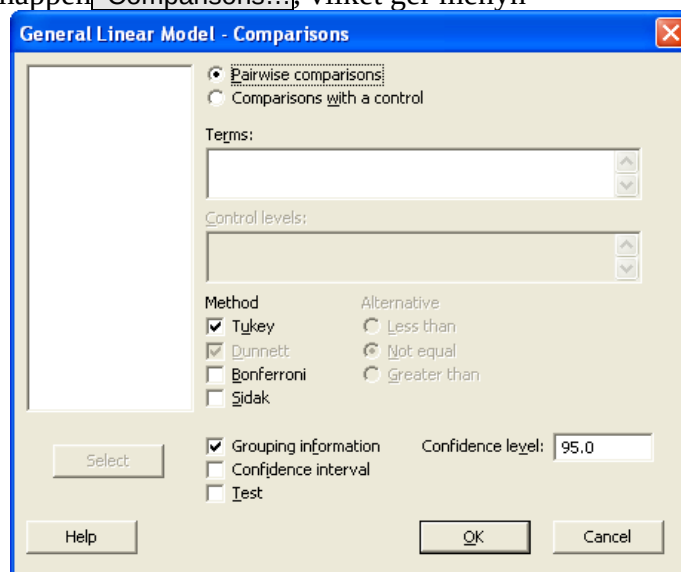


och resultatet blir

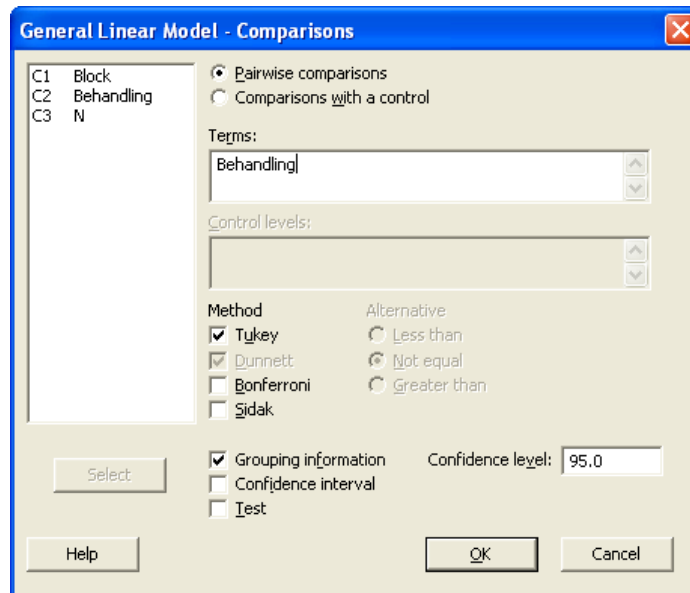


Även om man här har huggit av skalan på y-axeln så framgår det att det inte är så stor skillnad mellan blocken, däremot finns det i alla blocken samma mönster för skillnaderna mellan behandlingarna (dock en liten avvikelse i block III). Ser man figuren kan man gissa att behandling A kommer att visa sig ha signifikant lägre värde än de övriga, medan man inte säkert kan säga hur den övriga jämförelsen faller ut.

För att jämföra de olika nivåerna använder man återigen General Linear Model... men väljer denna gång också knappen **Comparisons...**, vilket ger menyn



Här måste man fylla i rutan vid Terms: för att få något resultat, och sedan kan man välja vilket test man vill ha (Dunnett's test är sammankopplat med jämförelse med en kontroll och är därför gråmarkerad). Lägg också märke till att Fisher's test inte är ett alternativ, Minitab tycker inte att man skall använda det i detta sammanhang. Gör man inget mer än att fyller i enligt



så blir resultatet (under den ANOVA-tabell som vi redan har presenterat) de två utskrifterna

Grouping Information Using Tukey Method and 95.0% Confidence

Behandling	N	Mean	Grouping
B	3	10.6	A
D	3	10.1	A
C	3	8.8	A
A	3	4.6	B

Means that do not share a letter are significantly different.

Sammanfattning: Det är behandling A som signifikant skiljer sig från de övriga, men däremot kan man inte visa någon signifikant skillnad mellan B, C och D.

9. Korrelation

Om man har tvådimensionella variabler kan man undersöka om det finns något samband mellan de båda variablerna. Det vanligaste måttet på detta samband är korrelationskoefficienten⁶ som brukar betecknas ρ (grekiska bokstaven "rå"). Det p -värde man kan få är ett p -värde för nollhypotesen att korrelationskoefficienten är 0 mot att korrelationskoefficienten inte är 0.

⁶ Den korrelationskoefficient som beräknas av Minitab brukar kallas Pearson's korrelationskoefficient. Man ser ibland istället Spearman's korrelationskoefficient som är ett robust alternativ, men denna korrelationskoefficient finns inte i Minitab.

Exempel 8. Bomull

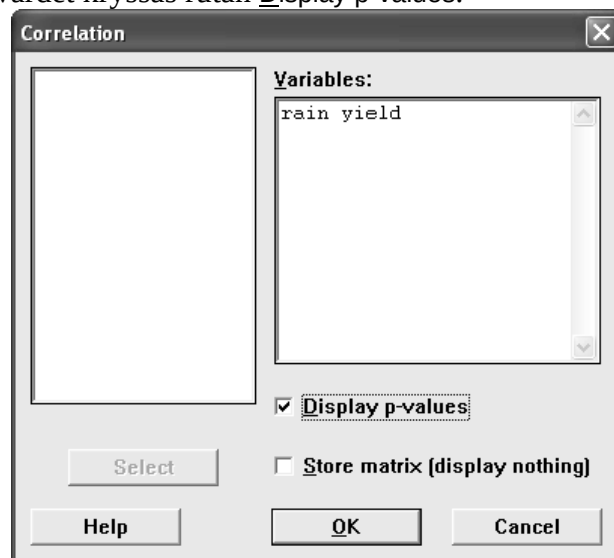
Man har värden på regnmängd och skörd av bomull. För att se om det finns något samband mellan dessa värden beräknas korrelationskoefficienten.

Regn (inches)	Skörd (lb/acres)
7.12	1037
63.54	380
47.38	416
45.92	427
8.68	619
50.86	388
44.46	321

Korrelationskoefficienten beräknas genom

Stat ► Basic Statistics ► Correlation...

och om man vill ha p-värdet kryssas rutan **Display p-values**.



Man kan stoppa in mer än två variabler i rutan **Variables**: och får då korrelationskoefficienten för alla par av variablerna. I detta fall med två variabler blir resultatet

Correlations: rain; yield

Pearson correlation of rain and yield = -0.836
P-Value = 0.019

Korrelationskoefficienten ligger alltid mellan -1 och 1, och här har vi en stor negativ korrelation. Det verkar alltså som om lite regn är relaterat till mycket bomull. ■

Korrelationen mäter linjärt beroende mellan variablerna men ibland kan konstiga värden påverka resultatet och därför bör man ta för vana att rita en scatterplot över sitt datamaterial när man beräknar korrelationer.

10. Regression

Regression är likartat med korrelation, men man skiljer det oftast genom att den ena variabeln här är något man kan styra själv (x -variabeln) medan responsen y är slumpmässig. Den varierar emellertid kring en rät linje. Genom att göra regression kan man få fram skattningen för linjen, och man kan också få en bild över den linje som är skattad. Ett exempel får räcka som exempel på detta.

Exempel 9. Koldioxid

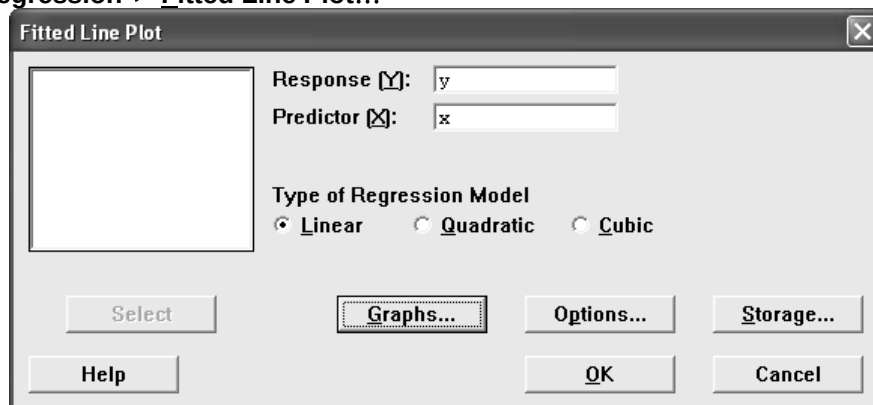
Följande data är för koldioxidkoncentration (x) och upptaget i blad av vete (y).

Konc (x)	75	100	100	120	130	160	160	190	200
Upptag (y)	0.00	0.65	0.50	1.00	0.95	1.80	2.10	2.80	2.50

Det första man bör göra är att använda plot-kommandot för att se om det är rimligt att ansätta en regressionsmodell, men vi lämnar det till läsaren. Minitab levererar ändå en bild över data när man gör regression enligt nedan.

För att få den ekvation som hör till linjen samt en del andra mått så väljs

Stat ► Regression ► Fitted Line Plot...



Resultatet kommer både i resultatfönstret och i en figur.

Regression Analysis: y versus x

The regression equation is
 $y = -1.671 + 0.02214 x$

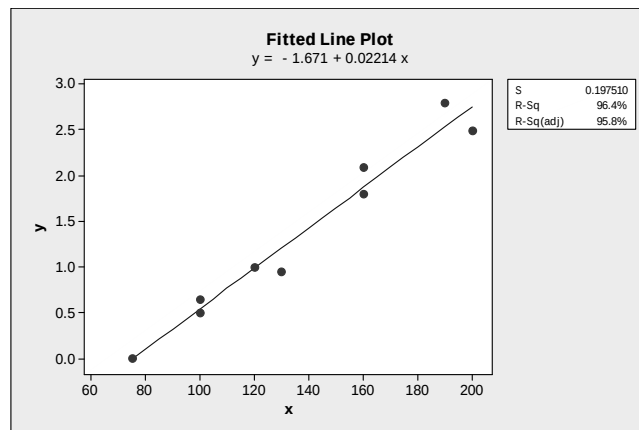
$S = 0.197510$ $R\text{-Sq} = 96.4\%$ $R\text{-Sq}(\text{adj}) = 95.8\%$

Analysis of Variance

Source	DF	SS	MS	F	P
Regression	1	7.23193	7.23193	185.38	0.000
Error	7	0.27307	0.03901		
Total	8	7.50500			

Regressionsekvationen anges till $y = -1.67 + 0.0221x$, och den andra viktiga siffran är p -värdet 0.000. Om regressionskoefficienten är 0 lönar det sig inte att använda x -värdet för att

förutse y -värdet. I detta fall är regressionskoefficienten signifikant skild från 0. Man får också en figur över regressionsekvationen enligt



11. Chi-två-test

Om man har diskreta värden som man vill räkna på något sätt så är det under avdelningen Tables man skall leta. Här finns möjlighet att göra lite av varje, men det viktigaste är nog möjligheten att göra så kallade Chi-två-test (χ^2 -test). Man brukar dela dessa test i "goodness-of-fit" och homogenitetstest. Goodness-of-fit förekommer till exempel när man har kastat en tärning många gånger och skall undersöka om den är falsk. Homogenitetstest används för att testa om två eller flera populationer är lika. Nu är emellertid inte populationerna normalfördelade utan innehåller element som kan klassificeras i någon klass. För att exemplifiera används samma exempel för båda testen.

Exempel 10. Tärningar

Det finns tre tärningar, en röd, en blå och en grön. För att undersöka om dessa tärningar är likvärdiga kastar man tärningarna vardera 500 gånger och får resultatet

	1	2	3	4	5	6
Röd	75	84	92	76	88	85
Blå	81	89	82	81	76	91
Grön	84	87	80	84	82	83

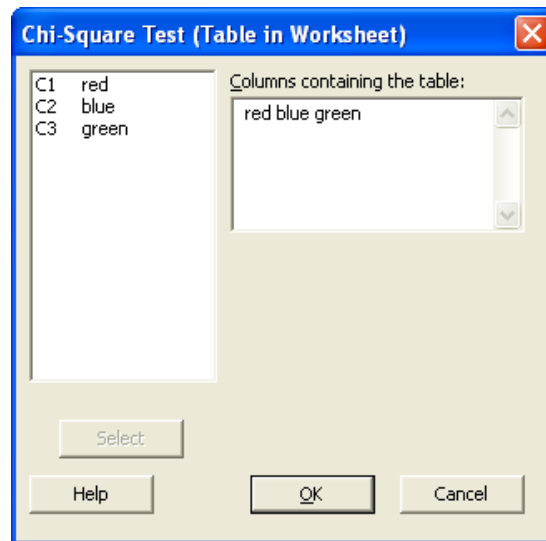
Man läser in data i tre kolumner enligt

red	blue	green
75	81	84
84	89	87
92	82	80
76	81	84
88	76	82
85	91	83

För att se om de är likvärdiga gör man ett homogenitetstest genom att välja

Stat ► Tables ► Chi-Square Test (Table in Worksheet)...

och fyller i enligt



Resultatet i resultatfönstret blir

Chi-Square Test: red; blue; green

Expected counts are printed below observed counts

Chi-Square contributions are printed below expected counts

	red	blue	green	Total
1	75	81	84	240
	80.00	80.00	80.00	
	0.313	0.013	0.200	
2	84	89	87	260
	86.67	86.67	86.67	
	0.082	0.063	0.001	
3	92	82	80	254
	84.67	84.67	84.67	
	0.635	0.084	0.257	
4	76	81	84	241
	80.33	80.33	80.33	
	0.234	0.006	0.167	
5	88	76	82	246
	82.00	82.00	82.00	
	0.439	0.439	0.000	
6	85	91	83	259
	86.33	86.33	86.33	
	0.021	0.252	0.129	

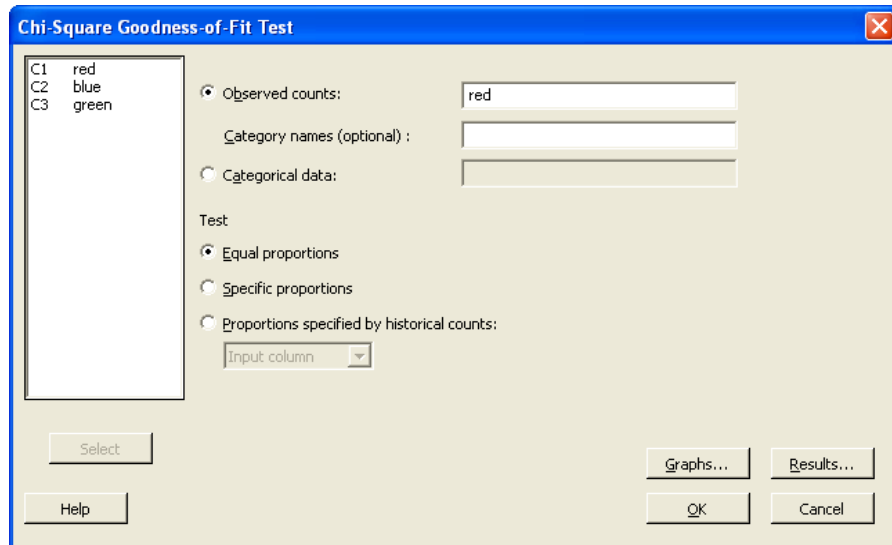
Total 500 500 500 1500

Chi-Sq = 3.334; DF = 10; P-Value = 0.972

Det intressanta värdet är det sista, 0.972, som visar att det inte är någon signifikant skillnad mellan de tre tärningarna. Lägg dock märke till man inte uttalar sig om huruvida tärningarna är falska eller ej, de kan mycket väl vara falska "åt samma håll".

Om man istället ställer frågan om den röda tärningen är falsk gör man ett goodness-of-fit-test (kallas ibland anpassningstest). Välj då

Stat ► Tables ► Chi-Square Goodness-of-Fit Test (One variable)...
och sedan



Resultatet blir nu (förutom några figurer som också produceras automatiskt)

Chi-Square Goodness-of-Fit Test for Observed Counts in Variable: red

Category	Observed	Test Proportion	Expected	Contribution to Chi-Sq
1	75	0.166667	83.3333	0.833333
2	84	0.166667	83.3333	0.005333
3	92	0.166667	83.3333	0.901333
4	76	0.166667	83.3333	0.645333
5	88	0.166667	83.3333	0.261333
6	85	0.166667	83.3333	0.033333

N	DF	Chi-Sq	P-Value
500	5	2.68	0.749

Slutsatsen av detta blir alltså att det inte finns några bevis för att den röda tärningen skulle vara falsk (p -värdet 0.749 är ju långt över 0.05). ■