

Statistik i Excel — en introduktion

Thommy Perlinger

När man använder statistik i Excel behöver man paketet **Data Analysis** som ligger under menyn **Verktyg** (Alt-y,d). Finns det inte där kan man installera det genom att i **Verktyg** välja **Tillägg** (Alt-y,l) och sedan bocka för **Analysis ToolPak**. Nu finns paketet tillgängligt.

1 Några tips om Excel

- Att beskriva ett statistiskt datamaterial innebär ofta att man behöver ange två eller flera variabler samtidigt. Excel kräver då att dessa kolumner angränsar till varandra. För att uppnå detta kan man dölja de för stunden irrelevanta kolumnerna genom att markera dem, högerklicka och sedan välja **Dölj**. De dolda kolumnerna görs synliga genom att markera de båda angränsande kolumnerna, högerklicka och välja **Ta fram**.
- En av fördelarna med Excel är *cellreferenser*, dvs möjligheten att i en cell använda information från andra celler genom det sk *funktionsbegreppet*. Man anger för Excel att cellen innehåller en funktion genom att börja inmatningen med ett likhetstecken, =. En formel kan t ex vara att i cellen C2 ange

$$= A2 * B2$$

När man använder datamaterial är det vanligt att man vill utföra samma operation på hela datamaterialet och det känns då onödigt att skriva in samma funktion om och om igen. Kanske vill vi i C3 ha $A3*B3$, i C4 ha $A4*B4$ osv. Detta uppnår vi genom att

1. Markera C2 och kopiera (CTRL-C).
2. Markera de fält i C-kolumnen dit vi vill kopiera formeln. Klistra in (CTRL-V).

Detta kallas för en *relativ cellreferens* eftersom referenserna är relativa i förhållande till formelns position. Ibland är funktionerna mer komplicerade och det kan vara så att det finns celler i referensen som skall användas i samtliga celler dit formeln kopieras. Detta måste man upplysa Excel om genom att göra en sk *absolut cellreferens*, vilket uppnås genom att låta ett dollartecken föregå bokstaven eller siffran, exempelvis \$A\$1. (Den del som dollartecknet föregår blir absolut.) Antag att vår formel i C2 är

$$= A2 * B2 + D2$$

där vi vill att innehållet i D2 skall användas i alla aktuella celler i C-kolumnen. Då måste man ändra formeln till

$$= A2 * B2 + D\$2$$

innan man kopierar den. Sammanfattningsvis gäller alltså att relativa referenser anpassas automatiskt när du kopierar dem, medan absoluta referenser alltid är desamma.

2 Beskrivande statistik

Vi förutsätter nu att vi har vårt datamaterial som ett arbetsblad i Excel med individer som rader och variabler som kolumner. För att finna de bäst lämpade metoderna för att beskriva detta datamaterial bör man börja med att notera

1. Vilken datanivå har den eller de variabler som vi skall beskriva?
2. Hur många variabler skall beskrivas samtidigt?

Vi delar därför in framställningen i två huvudavsnitt, dels metoder för kvalitativa variabler och dels metoder för kvantitativa variabler. Vi fokuserar på metoder där vi beskriver en variabel åt gången men anger även metoder för hur man beskriver två variabler samtidigt. Även om metoderna skiljer sig åt är principerna samma.

1. Strukturera datamaterialet med frekvenstabeller (en variabel) eller korstabeller (flera variabler).
2. Åskådliggör datamaterialet med någon form av diagram.

2.1 Kvalitativa variabler

2.1.1 Frekvenstabeller

För att konstruera frekvenstabeller för kvalitativa variabler (kategorivariabler) använder vi med fördel en *pivottabell* i Excel. Denna finner vi i **Rapport för pivottabell** under menyn **Data** (Alt-d,r). Här är det tre steg man skall ta sig igenom.

1. Markera vilken typ av data det är. Vi har en Excellista.
2. Markera det aktuella området inklusive etikett (variabelnamn).
3. Markera var frekvenstabellen skall placeras. Här finns dock även det viktigaste steget
 - **Layout:** Här får vi en schematisk uppställning av en frekvenstabell. Eftersom vi även har markerat etiketten finns den till höger i bilden. Dra denna till **Rad** vilket anger att detta är tabellens variabel. Dra den sedan även till **Data** vilket då eventuellt blir **Summa av kön**. Vi är dock intresserade av frekvenser (antal) varför detta skall ändras. Dubbelklicka på knappen **Summa av kön** och byt till **Antal**. I frekvenstabellen finns nu ett värdefält där det står "tom". Klicka på knappen **Kön** och bocka av denna.

Nu har vi fått en frekvenstabell som vi t ex kan använda för att konstruera diagram.

2.1.2 Diagram för att beskriva en variabel

Då man skall beskriva kvalitativa eller diskreta data med få variabelvärden använder man ofta *stapel*diagram (bar charts) eller *cirkel*diagram (pie chart). För att konstruera dessa diagram används diagramguiden i Excel som finns i **Diagram** under **Infoga** (Alt-i,i). Alternativt kan man direktklicka på diagramikonen. I diagramguiden skall man ta sig igenom fyra steg för att slutföra diagrammet.

1. Stapeldiagram kan i Excel antingen fås stående (vanligast) eller ligande. Välj stapel (stående) där det finns ett antal olika undertyper beroende på ändamålet. Även för cirkeldiagram finns ett antal olika varianter.
2. **Dataområde:** Välj pivottabellen som indataområde genom att vänsterklicka någonstans i tabellen.

3. Nu är det dags att börja snygga till diagrammet. Här finns ett stort antal möjligheter och förändringarna uppdateras direkt.
4. Här väljer man var diagrammet skall placeras.

Det finns även möjlighet att i efterhand justera diagrammet vilket man bör göra. Vi får bort de fula “knapparna” genom att högerklicka på någon av knapparna och välja **Dölj knappar för pivotdiagramfält**. Sedan är det bara att flytta över det till t ex Word.

Då det gäller kvalitativa variabler är det inte alltid uppenbart i vilken ordning variabelvärdena skall placeras. Ofta finns det dock en “naturlig” ordning som t ex vid veckodagar eller partiideologi etc. Det är dock inte självklart att Excel känner till alla sådana ordningsmöjligheter utan arbetar vanligtvis utifrån bokstavsordningen. Man kan dock berätta för Excel hur man vill ha värdena ordnade. Sådana “udda” ordningsmöjligheter finns under menyn **Verktyg–Alternativ** där man sedan väljer **Anpassa Lista** och **Ny Lista**. Placeras sedan in värdena i tur och ordning och skilj dem åt genom att trycka **Enter** mellan varje.

2.1.3 Korstabeller

En simultan frekvenstabell för två variabler kallas för en *korstabell*. Korstabeller är mycket viktiga både i beskrivande- och analytisk statistik. Även här använder vi med fördel en *pivottabell*, dvs **Rapport för pivottabell** under menyn **Data** (Alt-d,r). Enda skillnaden mot frekvenstabellen är att det blir två variabler inblandade. Markera för enkelhets skull hela datamaterialet (utan att få med några tomma celler). I **Layout** drar vi nu den ena variabeln till **Rad**, den andra till **Kolumn** (se nästa avsnitt) och sedan någon av dem till **Data**. Som ovan vänsterklickar man på denna knappen och byter **Summa** till **Antal**. Avsluta med att finna en lämplig placering av tabellen.

2.1.4 Diagram för att beskriva två variabler

För att beskriva en korstabell använder man lämpligen grupperat eller staplat stapeldiagram. Med ett sådant stapeldiagram blir variabelvärdena för den första variabeln uppdelade på värdena för den andra. Det är därför viktigt att man i pivotguiden väljer rätt variabel till **Rad** respektive **Kolumn**. Vill vi t ex se könsfördelning på olika utbildningar eller utbildningsfördelning för de båda könen? Förutom detta och att det finns flera stapelalternativ är tillvägagångssättet i princip identiskt med tidigare.

2.2 Kvantitativa variabler

2.2.1 Frekvenstabeller

För att på bästa sätt konstruera en frekvenstabell för en kvantitativ variabel bör vi ha en uppfattning om största och minsta värde. Därför ordnar vi materialet i storleksordning med avseende på den eller de variabler som vi är intresserade av.

- Markera datamaterialet. OBS! Det är viktigt att markera hela datamaterialet så att samtliga variabelvärden associeras med rätt individ även efter sorteringen.
 - Välj **Sortera** under menyn **Data** (Alt-d,o). Välj huvudvariabel och sorteringsordning. Man kan även välja att sortera efter undervariabler inom varje värde för huvudvariabeln.
- Nu har vi en uppfattning om min och max för vår variabel. Nästa steg blir att föra in klassgränserna i en egen kolumn i arbetsbladet. I cellerna skall endast *övre* klassgräns anges eftersom ett variabelvärde faller i den klass där det är mindre än eller lika med den utsatta gränsen. Fyller man inte i klassgränser konstruerar Excel automatiskt klasser. För kvalitativa variabler och kvantitativt diskreta variabler med få värden anger man samtliga variabelvärden som klassgränser. Vi skapar nu en frekvenstabell under **Verktyg–Data Analysis–Histogram**. Markera
 - Indataområde (med eller utan etikett. Ta hänsyn till detta genom att bocka för etikettrutan).
 - Klassgränser (Bin-range)(med etikett om detta markerades för indataområde). Anges antingen manuellt eller genom att markera en kolumn med färdiga (övre) klassgränser (se ovan).
 - Utdataområde. Här väljer man om man vill ha frekvenstabellen i samma arbetsblad, nytt arbetsblad eller ny arbetsbok. Vill man ha det i samma arbetsblad skall man markera den cell där man vill placera tabellens övre vänstra hörn.
 - Rutan längst ner, **Chart output**, för att få ett histogram. Detta placeras automatiskt strax intill frekvenstabellen.

2.2.2 Histogram

Histogrammet är till en början inte i publicerbart skick utan behöver antagligen en ordentlig putsning. Det finns dessutom två stora nackdelar med att

konstruera histogram i Excel jämfört med t ex Minitab. Dels måste klasserna vara lika breda och dels hamnar klassgränserna i mitten av histogramrektanglarna vilket gör att de ser ut som klassmitter. Vi snyggar till histogrammet genom att klicka på de olika delarna.

1. Rubriker bör ändras. Huvudrubrik samt rubriker på x - respektive y -axlarna. Ett tips är att rubriken på y -axeln, Frequency, och värdeskala kan raderas då denna ofta inte är så informativ. Istället kan man placera klassfrekvenser över respektive rektangel.
 - (a) Markera rubriken på y -axeln och radera. Markera motsvarande värdeaxel och välj **Mönster**. Markera **Inga** i samtliga tre fält för skalstreck.
 - (b) Dubbelklicka på någon av histogramrektanglarna. Välj **Dataetiketter** och markera **Visa värde**. Är det diskreta data med få variabelvärden så är vi antagligen nöjda med detta stolpdiagram men om vi har en kontinuerlig klassindelning bör det inte vara något glapp mellan rektanglarna. Detta justeras genom att välja **Alternativ** och i **Mellanrum i x-led** välja 0.

2.2.3 Spridningsdiagram

Då man vill undersöka om det finns något samband mellan två numeriska variabler använder man ofta ett *spridningsdiagram* (scatter plot eller xy -plot). Denna diagramtyp finns i diagramguiden, dvs i **Diagram** under **Infoga** (Alt-i,i), eller ett direktklick på diagramikonen. I diagramguiden skall man ta sig igenom fyra steg för att slutföra diagrammet.

1. Som **Diagramtyp** väljer vi **Punkt** och som **Undertyp** det diagram utan sammanbindande linjer.
2. I steg 2 klickar vi på **Seriefliken** och väljer där **Lägg till**. Då framträder tre indatafält.
 - (a) **Namn**: I det här frivilliga fältet anger vi diagramrubriken.
 - (b) **X-värden**: Markera x -kolumnen i arbetsbladet utan etikett.
 - (c) **Y-värden**: Markera y -kolumnen i arbetsbladet utan etikett.
3. Här har vi möjlighet att ändra rubriker. Huvudrubrik, x -axel och y -axel.
4. Här väljer man var diagrammet skall placeras.

Det finns även möjlighet att i efterhand justera diagrammet vilket man bör göra. Några förbättringsförslag är

- Ta bort förklaringsfältet till höger om diagrammet genom att markera, högerklicka och välj **Radera**.
- Ibland behöver man ändra på axlarnas skalor för att diagrammet skall bli presentabelt. Ställ pekaren vid respektive axel, högerklicka och välj **Formatera axel**. Under fliken **Skala** finns en del möjligheter att justera bilden.
- Stöddlinjer är överflödiga. Ställ pekaren på en av skallinjerna, högerklicka och välj **Radera**.
- Det ser bättre ut utan bakgrundsfärg och kantlinjer över och till höger. Ställ pekaren i diagrammet, högerklicka och välj **Formatera rityta**.
- Eventuellt kan man dessutom göra huvudrubriken något större och eventuellt i fetstil. Sedan är det bara att flytta över det till t ex Word.

2.3 Sammanfattande mått

Ofta behöver vi förutom frekvenstabell och diagram några sammanfattande mått på vårt datamaterial. Detta uppnås enklast genom att

1. **Verktyg—Data Analysis—Descriptive statistics.**
2. Markera de variabler som du vill ha sammanfattande statistik för. Ange dessutom om etiketterna är med.
3. Ange var du vill ha sammanfattningen lokaliserad. Tänk på att den består av ett stort antal rader.
4. Ange vad du vill skall vara med. I beskrivande syfte räcker det med att bocka för **Summary statistics**.

Sammanfattningen är *väl* tilltagen för de flesta användare. De för oss viktigaste måtten är

- *Lägesmått*: medelvärde (mean), median och typvärde (mode).
- *Spridningsmått*: standardavvikelse (standard deviation), varians (sample variance), standardfel för medelvärde (standard error) och variationsvidd (range).

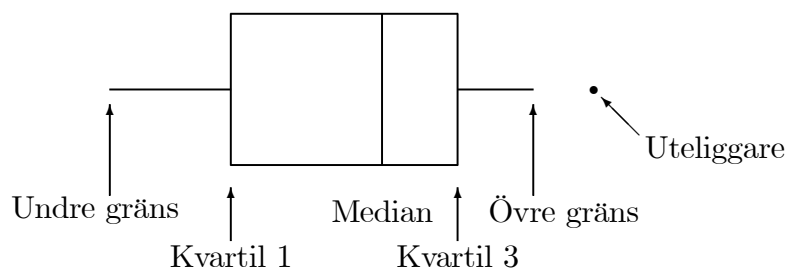
För övrigt kan man t ex använda *minimum* och *maximum* då man skall konstruera klassgränser för frekvenstabell. Några av de andra värdena kan också komma till användning då man vill anpassa Excel efter sina egna behov.

2.3.1 Boxplot

Ett utmärkt alternativ, eller komplement, till histogram är sk *boxplots* (låda-gram). Dessa faller under diagramkategorin *Explorativ Data Analys* (EDA) och finns tyvärr inte tillgängligt i Excel. Man kan dock förbereda värden för en boxplot genom att gå den hårda vägen, dvs genom att använda formler i Excel. För en nybörjare är detta besvärligt varför färdiga rutiner för detta finns att hämta i arbetsboken **Statistik i Excel**. Klicka fram arbetsbladet **Boxplot**. Här finns en rutin för att konstruera en boxplot för en variabel (univariat data) och en rutin för att skapa en boxplot med två boxar för en variabel uppdelad med avseende på någon kategorivariabel (0–1 variabel).

Univariata fallet

1. Mata in värden för den aktuella variabeln i kolumn A från A6 och nedåt
2. I kolumn B avgörs vilka (om några) observationer som är uteliggare. I B6 finns formeln för detta. Kopiera denna formel till de celler i B-kolumnen som är associerade med en observation i A-kolumnen.
3. Förutom uteliggarna (som är markerade med U i B-kolumnen) finns all relevant information under **Värden (univariata fallet)** i cellerna E7 till E12. Deras innebörd förklaras av följande boxplot.



Bivariata fallet

1. Mata in observationsvärdena för huvudvariabeln i kolumn H från H6 och nedåt och motsvarande värden för kategorivariabeln i kolumn I från I6 och nedåt. För att programmet skall fungera måste kategorivärdena vara 0 och 1. Om kategorierna är skrivna i klartext måste dessa kodas om, vilket uppnås genom att infoga en ny kolumn och använda OM-funktionen.
2. Nu skall observationerna delas upp med avseende på kategorivariabeln och sedan skall samma procedur som i det univariata fallet ovan utföras. Formler för detta finns i cellerna J6 till M6. Kopiera dessa formler nedåt i varje kolumn så att alla celler som associeras med observationer kommer med.
3. Förutom uteliggarna (som är markerade med U i L- respektive M-kolumnen) finns all relevant information under **Värden (bivariata fallet)** i cellerna E19 till E24 respektive E29 till E34. Innebörden av dessa värden är samma som i det univariata fallet och förklaras av boxploten ovan.

Nu har vi all information som behövs för att konstruera en boxplot. Då återstår bara problemet att få till själva diagrammet. Det finns ritprogram där det enklaste "programmet" är penna, papper och linjal.

3 Statistisk inferens

Statistisk inferens innebär statistisk *hypotesprövning* och/eller *konfidensintervall*. Det finns i Excel färdiga rutiner för att utföra hypotesprövning för två stickprov men inga rutiner för univariat data. Är man lite flitig går det dock utmärkt att skapa egna rutiner genom att kombinera statistiska, matematiska och logiska funktioner med formelkonceptet. De tillgängliga funktionerna finner man genom **Infoga-Funktion** (Alt-i,f) eller genom att trycka på funktionsikonen f_x . Detta är dock en marig uppgift som nybörjare i statistik (och eventuellt även i Excel) varför ett exempel på sådana rutiner finns att hämta på

driverhänvisning, fil:statistisk inferens i excel.

där jag har skapat inferensrutiner för den här kursen (även för tvåstickprovs-varianterna).

3.1 Arbetsboken “Statistik i Excel”

Förutom **Boxplot** finns fem kalkylblad i arbetsboken; **Ett stickprov**, **Oberoende stickprov**, **Parvisa observationer**, **Chi-2 metoden** och **Enkel linjär regression**. Den sistnämnda behandlas i ett senare avsnitt. Även om innehållet i dessa blad skiljer sig åt är upplägget av dem ungefär samma. I de viktigaste bladen, **Ett stickprov** och **Oberoende stickprov**, är framställningen uppdelad utifrån

1. om det är *rådata* eller *sammanfattning* av datamaterialet.
2. om det är en *kvantitativ* variabel eller en kodad kvalitativ, sk *0-1* variabel.

I **Parvisa observationer** och **Chi-2 metoden** förutsätts rådata.

Hur skall då data föras in i bladet? I celler där texten refererar till något statistiskt värde skall motsvarande värde hamna i cellen närmast till höger. Detta sker antingen genom att en formel använder information från andra celler i bladet eller att användaren matar in informationen i cellen. Om texten avslutas med ett frågetecken förväntas användaren bistå med information. Några av informationscellerna är självförklarande medan andra tarvar ytterligare förklaring.

- *Signifikansnivå*, α , skall anges som ett tal mellan 0 och 1.
- *Hypotetiskt värde* är värdet på populationsparametern (μ , $\mu_2 - \mu_1$, etc) under H_0 .
- *Mothypotesen* H_1 kan antingen vara ensidig, t ex $H_1 : \mu < \mu_0$, $H_1 : \mu > \mu_0$ eller tvåsidig, t ex $H_1 : \mu \neq \mu_0$. Vilken typ man vill ha anger man i cellen som associeras med *Sidor*($<$, $>$, $><$).
- *Konfidensgrad* bestäms automatiskt som $1 - \alpha$.
- Vid två oberoende stickprov skall man dessutom svara på huruvida populationsstandardavvikelserna förutsätts vara samma. Detta gör man genom att ange j eller n efter *Samma varians* (j/n).

Hypotestest och konfidensintervall för kvantitativa variabler använder t-fördelningen, dvs

- vid små stickprov förutsätts normalfördelat material.

- inga fördelningskrav vid stora stickprov ($n \geq 30$).
- båda stickproven stora vid jämförelse utan förutsättning om lika populationsstandardavvikelser.

För 0–1 variablerna förutsätts stora stickprov varför normalfördelningen används. Chi-2 metoden är lite speciell varför den får ett eget litet avsnitt.

3.1.1 Chi-2 metoden

Chi-2 metoden har framförallt två användningsområden för oss.

- Kvalitativ variabel med fler än två variabelvärden. (Med enbart två värden kodas den om till en 0–1 variabel.) Vi testar huruvida en viss sannolikhetsfördelning är rimlig. Exempelvis om en tärning är juste eller inte via

$$\begin{aligned} H_0 &: p_1 = p_2 = \dots = p_6 = \frac{1}{6} \\ H_1 &: \text{Någon annan fördelning.} \end{aligned}$$

- Undersök oberoendet mellan två kvalitativa variabler.

$$\begin{aligned} H_0 &: \text{Variablerna är oberoende} \\ H_1 &: \text{Någon form av beroende mellan variablerna} \end{aligned}$$

I båda fallen jämförs *observerade frekvenser* med *förväntade frekvenser*. Dessa matas i det första fallet in som två rader med vänster hörn förslagsvis i D5 respektive D15, och i det andra fallet som två korstabeller med övre vänster hörn förslagsvis i D5 respektive D15. Den statistiska jämförelsen görs via χ^2 -fördelningen med antal frihetsgrader i det första fallet lika med Antal kolumner-1, och i det andra fallet lika med (Antal kolumner-1)(Antal rader-1). Den statistiska funktionen som används i de här testen ligger i cellen till höger om rutan **p-värde** och är på formen

$$= \text{CHI2TEST}(\text{Observerade värden}; \text{Förväntade värden})$$

Här måste man berätta för Excel var datamaterialet finns. Det enklaste är att markera rutan med funktionen (B5) och ställa markören i läge och sedan markera det aktuella området. Har man två rader och tre kolumner för sina frekvenser blir det

$$= \text{CHI2TEST}(D5:F6; D15:F16)$$

4 Linjär regression

När vi undersöker samband mellan två eller flera numeriska variabler vill vi normalt göra mer än att bara rita ett spridningsdiagram, se avsnitt 2.2.3. Man skall dock alltid börja med ett sådant diagram för att få en idé om vilken typ av samband det rör sig om. Här förutsätter vi att sambandet är linjärt och skiljer på huruvida det endast är en förklarande variabel eller om det är flera.

4.1 Enkel linjär regression

Här förutsätter vi att vi endast har en förklarande variabel i modellen och första uppgiften blir därför att infoga en rät linje i spridningsdiagrammet. Vi förutsätter här att vi redan har konstruerat ett snyggt diagram enligt stegen i avsnitt 2.2.3. Regressionslinjen fås enklast genom att i diagrammet högerklicka på någon av datapunkterna, välja **Infoga trendlinje** och sedan välja **Linjär**. Detta är dock enbart i beskrivande syfte. Vi vill antagligen även kunna använda sambandet till att dra slutsatser för den bakomliggande populationen. För detta krävs en ordentlig regressionsanalys. Detta går att utföra i Excel men ger inte all den information som är önskvärd. Med hjälp av formler kan man dock skapa en rutin för enkel linjär regressionsanalys som ger en Minitabliknande utskrift. Denna finns i arbetsbladet *Enkel linjär regression* i arbetsboken *Statistik i Excel*.

4.1.1 Arbetsbladet “Enkel linjär regression”

I arbetsbladet *Enkel linjär regression* finns en Minitabutskrift i Exceltappning. Användaren har till uppgift att bistå med data där man skall placera y -värdena i kolumn A från A2 och nedåt samt x -värdena i kolumn B från B2 och nedåt. Sedan gäller det att få Excel till att utföra regressionen. Detta sker genom formeln REGR som är en sk matrisformel. Resultatet av en matrisformel kan vara ett stort antal värden vilket innebär att vi måste markera flera celler för utdataområde. Stegen blir

1. Mata in x - och y -värden. Skriv om så önskas in variabelnamnen i A1 respektive B1.
2. Markera området M1 till N5, dvs totalt 10 celler.
3. Välj **Infoga–Funktion** (Alt-i,f) (eller direkt via funktionsikonen, f_x). Välj **Statistik** och bläddra sedan ned till **REGR**. I denna dialogruta skall anges

- (a) **Kända y :** Markera y -värdena i A-kolumnen.
- (b) **Kända x :** Markera x -värdena i B-kolumnen.
- (c) **Konst:** Sant
- (d) **Statistik:** Sant
- (e) Avsluta nu *inte* med RETUR utan ange att detta är en matrisformel genom att avsluta med CTRL+SHIFT+RETUR.

Denna procedur leder till att cellerna M1 till N5 innehar följande regressionsinformation

b	a
s_b	s_a
R^2	s_e
F	fg
SSR	SSE

som sedan används för att skapa utskriften. Om storleken på datamaterialet förändras måste **REGR**-proceduren upprepas.

4.2 Multipel linjär regression

Nu förutsätter vi att vi har mer än en förklarande variabel i modellen. Det blir då svårare att åskådliggöra sambandet grafiskt och man får istället förlita sig till sina analytiska kunskaper. Även om den "riktiga" regressionsfunktionen i Excel inte är optimal är den tillräcklig för våra syften här. Stegen för att utföra en regressionsanalys är

1. Se till att ha datamaterialet väl samlat.
2. Välj **Verktyg–Data Analysis–Regression**.
3. I rutan **Regression** skall man nu ange
 - (a) **Input Y-range:** Markera (med etiketter) de celler som innehåller y -värdena (den beroende variabeln).
 - (b) **Input X-range:** Markera (med etiketter) de celler som innehåller x -värdena (de oberoende, eller förklarande, variablerna).
 - (c) Bocka för rutan **Labels**.
 - (d) Välj var utskriften skall placeras. Tänk på att det blir rätt många celler.