

**Något om forskningsmetodik och
om elementära statistiska storheter**

med övningsuppgifter, facit och hemtentamen

ISBN 978-91-639-6675-0

Eija I. Korhonen, 2005

Förord

När jag utöver mina bildkurser på Mångkulturella finska folkhögskolan, skulle undervisa YH-eleverna i forskningsmetodik och elementär statistik, visade det sig att skolan inte hade försett eleverna med några som helst läroböcker. YH-eleverna studerade på utbildningen "Behandlingspedagog med etnokulturell inriktning" och trots att de hade forskningsmetodik och statistik på schemat, saknade de litteratur i ämnet. Denna skrift som jag kallar för "*Något om forskningsmetodik och elementära statistiska storheter*" skrev jag, utan ersättning från skolan, för eleverna. Förutom textdelen, innehåller häftet övningsuppgifter och facit, samt en hemtentamen.

Medan jag vikarierade i forskningsmetodik och statistik, kompletterade jag denna skrift med korta utdrag ur vissa välskrivna böcker. Då de sidor som bifogades var viktiga för studenternas kunskapsinhämtning, och då de dessutom var pedagogiskt välformulerade, hade det för mig varit ett onödigt företag att försöka uttrycka det som redovisades på de aktuella sidorna bättre. Endast någon enstaka sida finns med i det här kompendiet, men för att den intresserade läsaren skall kunna förkovra sig, återges nedan en lista på de, för kunskapsinhämtningen, så viktiga sidorna, men det bästa är dock att läsa de nedan nämnda böcker i sin helhet :

Byström, J. (1998). Grundkurs i statistik. Stockholm: Natur och kultur.

- *Sidorna 22, 23 och 24*
- *Sidan 32*
- *Sidorna 48 och 49*

Carlsson, M., Rönér Douhan, G. (1992). Statistik – en introduktion. Runa Förlag.

Sidan 9 i denna skrift är en kopia av

- *sidorna 14 och 15*

Sidorna 30 till 43 i denna skrift är kopior av

- *Sidorna 31-39*

Till en uppgift, på sidan 16, har jag i detta kompendium, använt mig av en artikel i Göteborgs-Posten, 28 november 2005.

Sidorna 45 -49 i detta kompendium är exempel på empiri hämtad från min undersökning i samarbete med överläkare Anders Ringdahl på SU och min handledare Maud Eriksson-Mangold bland hörselskadade (*sidorna 1-5*).

Frågorna på sidan 50 i denna skrift baseras på frågor hämtade ur

Patel, R., & Davidsson, B. (1991). Forskningsmetodikens grunder. Lund: Studentlitteratur.

Också bilagan, sidan 52, är hämtad från ovan bok.

*

Då detta kompendium är avsedd som en introduktion till forskningsmetodik och elementär statistik, presenteras här långt ifrån alla undersökningsmetoder, och endast elementära statistiska storheter. Läsaren hänvisas till forskarutbildningskurser och utförlig speciallitteratur i ämnet!

Varför forskningsmetodik?

Idag översvämmas vi av information och kunskap i form av marknadsundersökningar, rapporter, utredningar, prognoser, analyser, forskningsrapporter, med mera. Nya råd och rön inom en uppsjö av olika ämnesområden når oss dagligen via dagstidningar, TV, radio, tidskrifter, internet, etc., och på våra arbetsplatser, skolor, sjukhus, mm., matas vi ständigt med ny information. Därutöver slinker reklamslag in i våra liv där det finns det minsta utrymme; både genom vår brevlåda och via media.

En stor del av vetenskaplig kunskap når allmänheten via massmedia. En dag står det i tidningen att vi måste äta fiberrik mat som är snålt på fettinnehåll om vi vill stanna friska och krya, men redan följande vecka kan rubrikerna ropa ut att ”läkare varnar för fettsnål kost” och ”fiberrik föda inte bra för alla!” Förskräckta läser vi de senaste forskningsrönen som handlar om att ”för mycket fibrer och för lite fett” gör småbarn sjuka, och medan vi följer rubrikerna, upptäcker vi att den ena larmrapporten, t.ex. om cancerframkallande ämnen, överröstar den andra. Specialisterna förser oss med minutiösa kunskaper om tillvarons minsta beståndsdelar, men trots en uppsjö av information, håller vi på att ”tappa fattningen” – eller kanske just på grund av det stora informationsflödet?

Specialiseringen i vårt postmoderna samhälle innebär i mångt och mycket att vår helhetsbild av tillvaron fragmenteras. Föreställ dig, till exempel *bilden av människan*, som idag har styckats upp i mikrobiologiska strukturer, DNA och molekyler. Men vem är hon ... människan? DNA säger oss föga om människans tillvaro, hennes värderingar, önskningar, förhoppningar eller drömmar.

Idag gäller reduktionismen och det kritiskt rationella vetenskapsidealet. Biologismen är åter på frammarsch och en del forskare letar i våra gener även efter förklaringar till mänskligt beteende. Våra gener bestämmer visserligen vår biologi, men vår biologi bestämmer inte helt och fullt över vår fria vilja, men statistiken, matematiken och experimentell design anses idag vara de viktigaste metodologiska verktygen för att producera pålitlig vetenskaplig kunskap.

Men det är inte bara därför som grundläggande kunskaper i statistik är viktiga. Statistiken används i mängder av undersökningar, utvärderingar och rapporter för att presentera för oss kunskap och information i komprimerad och överskådlig form. Till och med reklamen hänvisar till ”vetenskapliga bevis”. Företagen marknadsför allehanda produkter under förevändningar såsom ”7 av 9 blev nöjda, lyckliga, glada och fick ett underbart sexliv” tack vare att de har använt just det företagets produkter, men bör vi tro på allt som sägs vara vetenskapligt bevisat? Vilken sorts kunskap skall vi lita på? Kan vetenskapen lösa alla mänsklighetens problem? Eller tillkommer det rent av problem tack vare vetenskapen?

Frågor föder nya frågor och för att kunna orientera sig i informationsflödet och bedöma uppgifters pålitlighet är det bra att veta något om forskningsmetodik och något om elementära statistiska storheter. Det går, tyvärr, också ljuga med siffror och diagram! Därför skall vi titta lite närmare på de möjligheter statistiken erbjuder.

Idag är kunskapen dessvärre ofta kortlivad. Det gäller att ”hänga med” och vara delaktig i vad som händer i samhället, på arbetsplatsen, i skolan. Vi behöver ha ”öga för” vilken sorts information vi bör ta del av, och vilken sorts information vi kan leva förutan. Det gäller att överblicka, gallra, tolka och kritiskt granska det stora informationsflödet med hänsyn till

informationens reliabilitet, validitet och relevans. Häri ligger ett viktigt skäl varför det är på sin plats med elementära kunskaper i forskningsmetodik och statistik.

Enligt den allmänna uppfattningen samlar vetenskapen in fakta, men bilden är missvisande. I lika hög utsträckning *tillverkar* vetenskapsmännen och kvinnorna kunskaper. Vetenskap är en skapande process, så låt oss kika på hur tillverkningen av kunskap kan se ut och vilka är ingredienserna i processen. Låt oss lära oss något om metodologin och om hur man undviker fällorna med falsk marknadsföring med stöd i statistik! Det bästa sättet att lära sig är att börja med grunderna i statistik.

Grundläggande statistik

Problemställning och population

Statistik och diagram används ofta i syfte att presentera data i lättöverskådlig, komprimerad form. Statistik används också för att analysera insamlad information, men innan vi börjar undersöka ett fenomen eller en företeelse, och skaffa oss så mycket information om det som möjligt, måste vi precisera både vår **problemställning** och vilken **målgrupp** eller **population** vår problemställning gäller.

Med ett problem eller problemområde, i det här sammanhanget, syftar vi helt enkelt på något som vi vill skaffa oss kunskaper om. Därför är det viktigt att vi också letar upp den all den kunskap som redan finns i ämnet för att få en överblick över ämnesområdet innan vi börjar vår undersökning. Kunskapen är till en del kumulativ, och även om stora framsteg också kan ske plötsligt genom banbrytande undersökningar som till och med kan kullkasta tidigare kunskaper, skall man inte glömma att bakom stora framsteg alltid finns mycket nedlagt tid och möda. Ju mer kunskap vi har om vårt ämnesområde, desto lättare blir det för oss att precisera vår problemställning, dvs. avgränsa frågeställningen och välja de, för ämnet lämpligaste undersökningsmetoderna.

Det är, med andra ord, många saker som måste tas i beaktande redan från start. Några viktiga frågor är: *Vad avser vi att undersöka?* Vad, vilken företeelse, och vilken population (grupp/enhet) är målet för vårt intresse? Vad har tidigare gjorts inom ämnesområdet? Vilken typ av undersökning är det som vi vill genomföra, dvs. vilka metoder ämnar vi använda oss utav, och är dessa metoder de mest lämpade för ändamålet? *Varför?* Osv.

Genom att ställa oss själva frågan *varför* klargör vi syftet med vår undersökning. Ett sådant förtydligande hjälper oss också att precisera om det är etiskt försvarbart att vi genomför den planerade undersökningen. Vad säger forskningsetiken inom ämnesområdet? Vare sig vår studie är ett experiment, en intervju eller en enkätstudie, uppstår frågan: Tar tillräckligt med hänsyn till de personer som vi kommer att involvera i vår undersökning? En intervju kan till exempel uppfattas som alltför närgående eller alltför lång och tröttsam av deltagarna.

I alla sorters undersökningar måste vi fundera om det är rimligt att utsätta människor – våra tänkta intervjuobjekt, experimentdeltagare eller enkätbesvarare – för de frågor vi vill ställa, eller att utsätta dem för det besvär som vår undersökning medför. Här kommer kraven på *absolut frivillighet, fullständig information om undersökningen och dess syfte*,

anonymitetskravet och konfidentiell behandling av de erhållna svaren och personuppgifterna in i bilden.

Det är ytterst viktigt att vi ger korrekt, fullständig och uppriktig information till de presumtiva deltagarna. Det är inte etiskt försvarbart att t.ex. av rädsla för att någon skulle backa ut ur vår studie, låta bli att yttra hela sanningen om undersökningen och dess syfte till deltagarna. Vi måste vara uppriktiga även om det skulle innebära att vi mister försökspersoner!

Alla sorters deltagande i undersökningar skall bygga på frivillighet som baseras på fullständiga och korrekta uppgifter om undersökningen och dess syfte. Om vi t.ex. lovar att de inblandade personerna får vara anonyma, får vi inte ens koda eventuella enkäter för att kunna skicka påminnelser (!) om dem, och har vi fått personuppgifter eller andra uppgifter om deltagarna som kan binda individer – till exempel till enkätsvarsblanketter – måste vi efter avslutat undersökning avkoda (!) personuppgifterna.

Det är mycket viktigt att behandla all information vi får konfidentiellt och respektera deltagarnas rätt till utförlig information och fullständig anonymitet! Därutöver måste vi också ta ställning till om vår undersökning är ekonomiskt försvarbar och om tiden räcker till. (Läs vidare i kurslitteratur).

Tre huvudtyper av undersökningar med hänsyn till tidigare kunskapen inom problemområdet

Saknas det tidigare kunskaper inom problemområdet eller om den företeelse som vi är intresserade av, och vi vill lära oss så mycket som möjligt om problemområdet, gör vi kanske en *explorativ pilot studie eller explorativ undersökning*. Syftet med en explorativ undersökning är att få så allsidig kunskap som möjligt om företeelsen/problemställningen i fråga. En pilot studie syftar på en första studie som undersöker problemområdet i avsikt att ge en fingervisning om våra antaganden om ämnesområdet/verkligheten. Om våra antagande bekräftas, kommer de sedan ytterligare att prövas mot verkligheten med hjälp av flera och större studier.

Finns det redan mycket kunskaper inom ett problemområde, vill vi kanske fördjupa oss i en väl avgränsad del av problemet och göra en detaljerad *deskriptiv undersökning* om företeelsen eller problemet ifråga.

Inom problemområden där man redan har stora kunskaper, har man utvecklat modeller eller teorier om verkligheten. Modellerna utvecklas till teorier, och föranleder nya antaganden, så kallade *hypoteser* om verkligheten. Hypoteserna kan uttrycka olika sorters samband mellan olika faktorer, till exempel "om ... så"; "om en person har ett stabilt socialt nätverk, så kan detta speglas i graden av individens självkänsla".

Att hypoteser om verkligheten är rimliga och relevanta, måste testas. Kunskapen bör grundas *empiriskt*, dvs. kunskap bör *bekräftas genom observationer av verkligheten*. När så sker har vi att göra med *hypotesprövande undersökningar* vars resultat antingen *falsifierar vår hypotes*, dvs. observationer av verkligheten motsäger hypotesen som då förkastas som felaktig, eller *verifierar vår hypotes* – med andra ord; observationer av verkligheten ger stöd åt hypotesen. Verifieras vår hypotes innebär det att man sedan kan arbeta vidare med den och använda den som stöd för modeller i teoribildning som förklarar verkligheten.

Vilken typ av undersökning vi än ämnar företa oss, och oavsett vad vårt problemområde handlar om, är det viktigt att vi inhämtar kunskaper om problemområdet innan vi börjar vår undersökning. Genom att lära oss mycket så mycket som möjligt om ämnesområdet i fråga, får vi en bra överblick och har lättare att avgränsa problemet. Frågeställningen blir hanterbar: Helt enkelt, vi preciserar problemet. I Runa Patels och Bo Davidssons bok om *Forskningsmetodikens grunder* står det mycket matnyttigt om tillvägagångssätten. Läs boken!

Forskning, teoribildning, modeller

Deduktion och induktion

När ett antal hypoteser inom ett problemområde har *verifierats*, dvs. hypoteserna har empiriskt prövats mot verkligheten, och forskarnas antaganden har grundats på observationer av verkligheten utvecklar forskarna en *modell* eller en *teori*.

En teori bygger sålunda på ett antal verifierade hypoteser som ger en sammanhängande bild. En modell, däremot, har inte samma anspråk som en teori att slutgiltigt förklara verkligheten. En modell fungerar mer som en länk mellan teori och verklighet, och ibland som ett led i teoribildningen, då forskarna prövar hypoteser. Ett sådant tillvägagångssätt vid generering av teorier kallas för den *hypotetiskt deduktiva metoden*, eller *bevisningens väg*, därför att den förklarar verkligheten utifrån befintlig kunskap och med hjälp av empiriskt prövbara hypoteser. *Induktion* eller *upptäckandets väg*, däremot, inhämtar först kunskap från verkligheten för att sedan dra slutsatser och bygga teorier från den kunskap som man har samlat in.

En invändning mot den induktiva teoriproduktionen är att man aldrig kan vara säker på kunskapens generaliserbarhet. Ett exempel som jag ofta hörde under min utbildning var att om man observerar svanar och upptäcker enbart vita svanar, så drar man kanske den felaktiga slutsatsen att ”alla svanar är vita”. Det kan alltid dyka upp ett undantag. Problemet med den deduktiva metoden för teorigenerering, i sin tur, kan till exempel vara att man går miste om kunskap om verklighetens beskaffenhet tack vare att man är styrd av antaganden eller teorier redan i utgångsläget. Vi bär alla färgade glasögon. Ingen går fri från förutfattade meningar. Dessutom har vi alla en historia i vårt bagage och inte ens forskaren kan ställa sig objektivt utanför sin historia eller utanför själva forskningsprocessen.

Kvantitativa och kvalitativa metoder

Kvantitativa och kvalitativa metoder syftar på hur vi samlar in, bearbetar, beskriver, analyserar och rapporterar informationen. De statistiska metoderna hör till kvantitativa metoder. Statistik syftade från början på *kunskaper om staten*, och den *beskrivande* eller *deskriptiva statistiken* svarar bland annat på frågor som: Var? När? Vilka? – eller beskriver samband.

Kvalitativa undersökningar och verbala analyser har avsikten att *tolka* och *förstå* ett fenomen och/eller dess underliggande mönster. Men också i kvantitativa undersökningar förekommer verbala analyser baserade på statistiska beräkningar. Vi talar då om statistisk *inferens* eller *slutledning*. Det är heller inte uteslutet att man i kvalitativa analyser använder sig av deskriptiv statistik eller tabeller. Man kan t.ex. med hjälp av en korstabell berätta något om *demografiska variabler* (demografi: befolkningsstatistik, demografiska variabler kan t.ex. vara utbildning, ålder och yrke, etc.).

Utöver de rent metodologiska skillnaderna mellan kvantitativa metoder som grundar sig på statistiska beräkningar och kvalitativa metoder vilka avser att verbalt tolka och förstå undersökningsresultaten, finns det en viktig skillnad mellan dessa metoder beträffande resultatens *generaliserbarhet*. Endast kvantitativa resultat kan (om viktiga kriterier är uppfyllda, t.ex. att sampelurvalet är tillräckligt stort och representativt för den grupp eller population som man avser undersöka) generaliseras till att gälla hela den population som man söker kunskaper om. Ur de kvalitativa studierna kan man endast dra generella slutsatser i förhållande till teorier och inte i förhållande till populationen! Detta är en viktig skillnad att hålla i minnet.

Skillnaden mellan deskriptiv statistik och statistisk inferens

Statistik används alltså på två olika sätt, dels för att *beskriva*, och dels för att *analysera*.

Beskrivande statistik avser i första hand att genom *tabeller och grafiska framställningar* ge en överskådlig, komprimerad presentation av det material som för tillfället studeras. I den beskrivande statistiken använder man sig till exempel av *lägesmått* och *spridningsmått*.

Till skillnad från den beskrivande statistiken, kan man med hjälp av analyserande statistiskt – *statistisk inferens* (dock med vissa säkerhetsmarginaler) dra slutsatser från materialet.

Begrepp – variabel - mätning

När man talar om *begrepp* inom vetenskaperna syftar man på en *bestämning*, en *definition* eller en *rad inbördes relaterade bestämningar/definitioner*, som *betecknar ett objekt, en företeelse eller en grupp, eller grupper av objekt/företeelser, osv.* Ett exempel är det väldefinierade begreppet en *triangel*. Det är allmänt accepterat och överenskommet att en triangel är en tresidig, rätlinjig geometrisk figur vars vinkelsumma alltid är 180 grader. Ingen ifrågasätter den definitionen av begreppet ”triangel”, men det är sällan som våra begrepp är så entydiga.

Inom samhällsvetenskaplig forskning kan begreppen ibland vara otydliga och svårfångade. Ju mer ”flytande” eller ”otydligt” ett begrepp är, desto mer noggrant måste vi förklara vad vi menar med begreppet. För att sedan också kunna forska på vårt begrepp, måste vi *operationellt definiera* det. Operationalisering betyder att vi måste med bästa möjliga *reliabilitet* (”tillförlitlighet”) och *validitet* (”med stöd i fakta” - mer om reliabilitet och validitet

senare) kunna beskriva begreppet på ett sådant sätt att vi också kan undersöka eller ”mäta” det som vi syftar på med begreppet. När man samlar in kvantitativ data kallar man det för *mätning*.

Begreppet måste alltså få en operationell definition. Skall vi till exempel undersöka någonting svårfångat, låt oss säga intelligens, måste vi precisera vad vi menar med intelligens, och hur vi avser att mäta just denna egenskap. Som ni förstår, är intelligens inte ett enhetligt begrepp. Handlar intelligens om faktakunskaper eller problemlösningsförmåga? Om emotionell intelligens eller moral? Betyder det att man är snabbtänkt och kvick, eller har tempot ingen betydelse för definitionen av intelligens? Man kan vara bra på något och okunnig om mycket annat, eller förstå sina känslor, men sakna förmågan att kommunicera dem, osv. Intelligens kan också tänkas variera över tid och mognadsgrad, osv. Hur undersöker man bäst eller mäter någonting sådant?

När våra begrepp är definierade och operationaliserade, måste vi fundera på vilket tillvägagångssätt bäst lämpar sig för att undersöka det vi avser att undersöka. Vi kanske gör en kvalitativ undersökning och intervjuar människor kring ett visst begrepp, en frågeställning, en attityd (osv.), eller riggar upp en välplanerad och välkontrollerad experimentell situation, eller sammanställer enkäter med frågor med svarsalternativ i intervallskala? Oavsett vilket, skall metoden som vi väljer väl tjäna sitt syfte så att vår undersökning blir meningsfull och pålitlig. Trots allt, ber vi människor om en tjänst när de skall delta i vår undersökning. Vi får inte besvära folk i onödan genom att vara dåligt förberedda! Såsom tidigare påpekats, måste vi vara på det klara med relevansen av våra metoder och val av vårt tillvägagångssätt, och medvetna om att vi måste noggrant följa etiska forskningsregler. Då detta kompendium är avsedd som en introduktion till forskningsmetodik och elementär statistik, presenteras här inte många undersökningsmetoder. Läsaren hänvisas till forskarutbildningskurser och utförlig speciallitteratur i ämnet!

När vi kommit så långt att vi har gjort observationer gällande vårt begrepp, en viss företeelse, eller om vissa aspekter eller egenskaper (osv.) och dessa observationer antar varierade värden, talar vi om *en variabel*. En variabels *variationsområde* eller *variationsvidd*, dvs. *range*, syftar på det totala antalet värden som en variabel kan anta. Exempelvis kan man tänka sig att man i statistikprovet kan få från 0 poäng upp till 10 poäng, vilket innebär att provresultaten har en variationsvid eller range 0-10.

Om man på ett meningsfullt sätt kan ordna variabelvärdena (observationerna) i en serie och säga att t.ex. en person har ett större värde än en annan, låt oss säga på statistikprovet, kallas variabeln *kvantitativ*. Kön däremot, som vanligtvis betecknar endera en man eller en kvinna, är en *kvalitativ* variabel. Obs. Förväxla inte kvalitativa variabler med kvalitativa metoder!

När en kvantitativ variabel endast kan anta heltalsvärden, 0, 1, 2, osv., kallas den *diskret* (eller *diskontinuerlig*). En kvantitativ variabel som kan anta samtliga reella talvärden kallas *kontinuerlig*. Exempel på kontinuerliga variabler är reaktionssnabbhet, längd, vikt och ålder.

Låt oss tänka att statistikprovresultatet endast kan anta heltalsvärden från 0 till 10. Statistikprovresultatet är sålunda en diskret variabel. För vissa syften är det dock praktiskt att tänka oss variabeln som kontinuerlig. Om det bakomliggande syftet med vårt statistikprov är t.ex. att mäta ”förmågan att förstå och hantera siffror” och om vi dessutom på ett meningsfullt sätt kan rangordna provresultaten efter antalet erhållna poäng, så att en person kan ha högre

poäng i ”förmågan att förstå och hantera siffror” än en annan, kan vi behandla vår variabel som *approximativt kontinuerlig*.

Population – sampling – stickprov

När vi ämnar göra en undersökning av något slag, måste vi bestämma oss för vilken grupp eller *population* vi söker information om. Alla som ingår i en population har en på förhand bestämd egenskap gemensam. En population kan t.ex. vara alla svenskfödda sexåringar i Sverige eller alla treåringar i Göteborg, men det kan också vara alla EU länder, alla tidningsartiklar om ett visst ämne, alla hästar på Åby travbana i oktober i år, osv.

Är populationen liten, kan vi göra observationer eller fråga alla (individer, hästar, eller vad de nu mån vara) beträffande det som vi ämnar undersöka. En sådan undersökning kallas för *totalundersökning*, men vi kan sällan tillfråga/observera så stora grupper. Därför måste vi på ett lämpligt sätt välja ut ett hanterbart antal personer, hästar, etc. som får representera hela populationen, dvs. ett *stickprov*.

För att vi skall kunna dra generella slutsatser från stickprovet måste stickprovet vara tillräckligt stort och representativt för hela populationen. Det krävs att vi plockar ut stickprovet genom ett *obundet slumpmässigt urval (OSU)* så att ingen individ (häst eller vad det nu gäller) har varken större eller mindre möjlighet att bli utplockat till att vara med i vår undersökning. Ett slumpmässigt urval av stickprov kallas för *sampling*. Vi kan t.ex. numrera alla individer i populationen och med hjälp av en slumpstalstabelle lottat fram undersökningsdeltagarna. Då har alla samma möjlighet att vara med i undersökningen.

Ett annat sätt är att systematiskt välja ut var 10:e eller var 20:e person i en förteckning över personer i populationen. Detta kallas för *systematisk urval*.

Ett tredje sätt är att i princip göra ”flera” slumpmässiga urval ur samma population: Om vi, t.ex. vet, att populationen består av 60 % män och 40 % kvinnor, kan vi slumpmässigt genom lottdragning få fram det antal kvinnor till undersökningen som motsvarar 40 % av stickprovets storlek, och också genom lottdragning få fram det antal män som motsvarar 60 % av stickprovets storlek till undersökningen. På så sätt ”härmar” vi populationens egenskaper genom att ha 60 % män och 40 % kvinnor också i stickprovet. Ett sådant urval kallas för *stratifierat urval*.

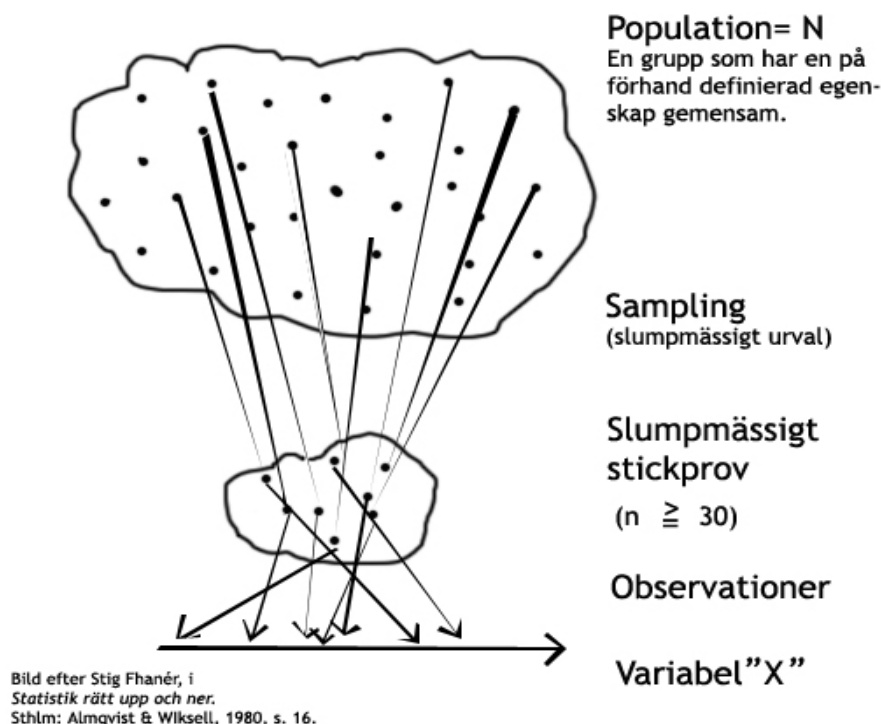


Bild 1. Population, samling, slumpmässigt stickprov, observationer av variabel X, illustration efter Stig Fhanér (1980).

Olika skalor

I en kvantitativ studie studerar vi olika egenskaper som varierar mellan deltagarna eller observationerna i studien. Det finns fyra skalnivåer. De är *nominalskala*, *ordinalskala*, *intervallskala* och *kvotskala*.

Nominalskala

Nominalskalans mätvärden ger en klassindelning i kategorier som oftast är av kvalitativ natur, t.ex. man eller kvinna. Man kan också tänka sig som exempel olika yrkeskategorier som arbetar på ett sjukhus; läkare, sjuksköterskor, undersköterskor och sjukvårdsbiträden. Siffervärdena på nominalskalan får alltså ingen numerisk innebörd. Om nominalskalan endast kan anta två kvalitéer, ex, man eller kvinna, kallas den dikotom eller binär skala.

Ordinalskala

Ordinalskalans mätvärden kan rangordnas men rangordningen säger ingenting om avståndet mellan mätvärden. Därför kan man inte jämföra t.ex. andra platsen med tredje platsen, då avståndet mellan mätvärden kan variera. Det kan, t.ex., vara fler poäng som skiljer placeringarna 1:a och 2:a, än vad som skiljer åt placeringarna 2:a och 3:a i en tävling. Ordinalskalan ger *ordning*.

Intervallskala

Intervallskalans mätvärden kan rangordnas. Dessutom kan man studera skillnader mellan de olika mätvärdena eftersom avståndet

dem emellan är det samma (man kan tala om ekvidistans).
Exempel på intervallskala är temperatur.

Kvotskala

Också kvotskalans mätvärden kan rangordnas. Dessutom kan man studera skillnader mellan olika mätvärden eftersom avståndet mellan mätvärden är det samma igenom hela skalan. Därtill har kvotskalan en absolut nollpunkt. Ett vanligt exempel på kvotskala är antalet barn i en familj eftersom en familj kan ha 0 barn. Observera att temperatur inte är av kvotskalekaraktär även om det finns en nollpunkt. Temperaturen 0 grader betyder inte avsaknad av temperatur – det finns också minusgrader!

Reliabilitet och validitet

När man talar om mätningens kvalitet, talar man om olika grader av validitet och reliabilitet. Det finns olika sorters validitet, men i princip syftar man med *hög validitet på att mätinstrumenten verkligen mäter det som man avser att mäta.*

Reliabilitet innebär graden av mätnoggrannhet, dvs. den säkerhet med vilken den aktuella variabeln mäts. Man brukar beräkna och ange reliabiliteten med Cronbach's reliabilitets koefficient α .

Kontrollfrågor: Har jag förstått?

Förstår jag innebörden av och skillnaderna mellan explorativa, deskriptiva och hypotesprövande undersökningar?

Förstår jag skillnaden mellan induktion och deduktion?

Vad menas med en hypotes?

Vad innebär det att en hypotes är verifierad? Eller att en hypotes är falsifierad?

(PS. En *korroborerad utsaga av verkligheten* är en tillfälligt antagen utsaga tills verkligheten motsäger den.)

Vad är en population?

Hur kan jag erhålla ett representativt urval ur en population och vad kallas ett sådant urval?

Vilka olika slumpmässiga urvalsprocesser känner jag till?

Vilkens sorts undersökningsresultat kan man generalisera till att gälla hela populationen?

Varför kan man inte generalisera alla undersökningsresultat till hela populationer?

Finns det etiska aspekter som är viktiga att tänka på om man ämnar genomföra en undersökning?

Vad menas med en variabel?

Vad är demografiska variabler för något?

Vad menas med en diskret variabel?

Vad menas med hög validitet?

Vad menas med mätningens reliabilitet?

Statistik används på två olika sätt. Vilka två sätt? Och vad innebär de?

Vad menas med en approximativt kontinuerlig variabel?

Vad är skillnaden mellan ordinalskala och intervallskala=

Vad karaktäriserar kvotskalan?

Vad är en nominalskala?

Rådata och fyrfältstabeller (korstabeller)

Ett behändigt och lättöverskådligt sätt att presentera data är att presentera den i så kallade fyrfältstabeller eller korstabeller. Sådana tabeller kan dessutom ha fler än fyra celler, men då förlorar man något av tabellens lättöverskådlighet. Korstabeller används oftast när man vill presentera resultatet på någon fråga i förhållande till en kvalitativ bakgrundsvariabel, t.ex. inställningen till statistik i relation till kön, utbildning, eller något annat.

För en korstabell måste vi ha mätt två variabler för varje person som ingår i undersökningen. *Rådatan* består sålunda av parvisa mätvärden. Nedan finns ett exempel på en fyrfältstabell för variablerna kön och inställning till statistik som besvarats med ”ja, jag tycker om statistik”, dvs. *positiv inställning* till statistik, och ”nej, jag tycker inte om statistik”, dvs. *negativ inställning*, vilket innebär att vår data är av ordinal karaktär. Observera hur man skriver rubrik och titlarna för tabellens kolumner och rader! Glöm heller inte att nämna sampelstorleken!

Tabell z: Fyrfältstabell över variablerna kön och inställning till statistik med absoluta frekvenser, $n=100$.

		Inställning till statistik		
		Positivt inställd	Negativt inställd	
Kön	Man	35	15	50
	Kvinna	40	10	50
		75	25	$n = 100$

När vi läser tabellen får vi information om hur könsfördelningen ser ut beträffande vår fråga, i ett stickprov på 100 personer varav 50 var män och 50 var kvinnor. Vid fyrfältstabeller är det viktigt att observera vad tabellen egentligen säger. Ett bra sätt att avgöra vad tabellen säger är att läsa av det totala antalet tillfrågade, dvs. den *absoluta frekvensen*, eller den totala andelen svarande uttryckt i procent, dvs. den *relativa frekvensen*. Ligger dessa summor radvis eller kolumnvis? Och vad är rubriken på raden, respektive kolumnen?

Övning: Rita en fyrfältstabell över följande bakgrundsdata i en pilotstudie på $n=49$.
 Antalet ensamboende kvinnor = 17, antalet sammanboende kvinnor = 8
 Antalet ensamboende män = 12, antalet sammanboende män = 12

	kvinnor	män	
ensamboende			
sambo			
			$n =$

Statistisk beskrivning av ett stickprov med en kvantitativ, diskret variabel

Låt oss säga att ni har gjort en stickprovsundersökning med avseende på en kvantitativ, diskret variabel och att ni vill på ett överskådligt sätt beskriva resultatet. Er variabel kan t.ex. vara antalet barn per familj. Eftersom "barn" enbart kan anta heltalsvärden är variabeln diskret till sin karaktär (numera mäter man dock ofta antalet barn med intervallskala och anger "procentuella delar av barn" ☺).

I vårt exempel har ni undersökt 50 familjer som valts ut slumpmässigt med hjälp av en slumpstalstabell. Det innebär att stickprovsstorleken är 50 familjer, vilket skall framgå av texten. Stickprovsstorleken betecknar vi med ett litet n . Populationen (som betecknas med ett stort N) kan t.ex. tänkas vara alla barnfamiljer, mantalsskrivna i storGöteborg.

För att statistiskt beskriva antalet familjer och barn, gör vi en frekvenstabell, där frekvensen, betecknad med f , anger antalet familjer som har 0, 1, 2, osv. barn. *Frekvensen visar alltså hur vanligt förekommande något är.* Frekvenstabellen i det här fallet beskriver hur familjerna i stickprovet fördelar sig över de olika variabelvärdena, helt enkelt, hur många barn familjerna har – eller med andra ord – vi beskriver stickprovets frekvensfördelning.

I kolumnen till vänster i vår Tabell x nedan, skriver vi vår X variabel, antalet barn. Nästa kolumn betecknar den *absoluta frekvensen*, betecknad med f . Frekvensen anger antalet familjer med ett visst antal barn (0, 1, 2, osv.)

Tabell x: Antal barn per familj

Antal barn (x)	Frekvens (f)	Kumulativ f	Relativ frekvens (f/n)	Frekvens i procent	kum.%
0	15	15	0.30	30	
1	10	25	0.20	20	
2	13	38	0.26	26	
3	6	44	0.12	12	
4	3	47	0.06	6	
5	3	50	0.06	6	
n=50			1.0	100%	100%

Räknar vi ihop siffrorna i kolumnen för den absoluta frekvensen får vi summan 50. Detta tal representerar sålunda stickprovsstorleken. Den kumulativa frekvensen får vi genom att summera, rad för rad, talen i den absoluta frekvensen. I kolumnen till höger om den kumulativa frekvensen har vi räknat ut den relativa frekvensen. Parentesen (f/n) anger hur den relativa frekvensen räknas ut. Vi delar, helt enkelt, frekvensen (f) med stickprovsstorleken, här ($n=50$). Den översta summan får vi alltså genom att dela frekvensen (barnlösa familjer) $15/50 = 0.30$. Nästa relativa frekvens får vi genom att dela frekvensen 10 (familjer med ett barn) med stickprovsstorleken, osv. *Summan av den relativa frekvensen blir alltid 1.*

Om vårt stickprov är minst 50, kan vi räkna ut frekvensen i procent. Då multiplicerar vi den relativa frekvensen med 100 för att få fram procenttalet.

Uppgift: Du kan själv räkna ut den kumulativa procenten och fylla i i tabellen längst ut till höger. *Summan av frekvensen uttryckt i procent blir alltid 100.*

En grafisk beskrivning av frekvensen familjer med x antal barn

Eftersom det i vårt exempel är frågan om en diskret variabel, skall vi välja ett stolpdiaqram för att grafiskt illustrera frekvensen familjer med X antal barn. Observera, att medan man skriver tabellrubriker ovanför av tabellerna, skriver man diagramrubriken, här *Diagram 1: Antalet barn per familj* under diagrammet. Det sistnämnda gäller också för bilder, dvs. bildrubriken skrivs alltid under bilderna.

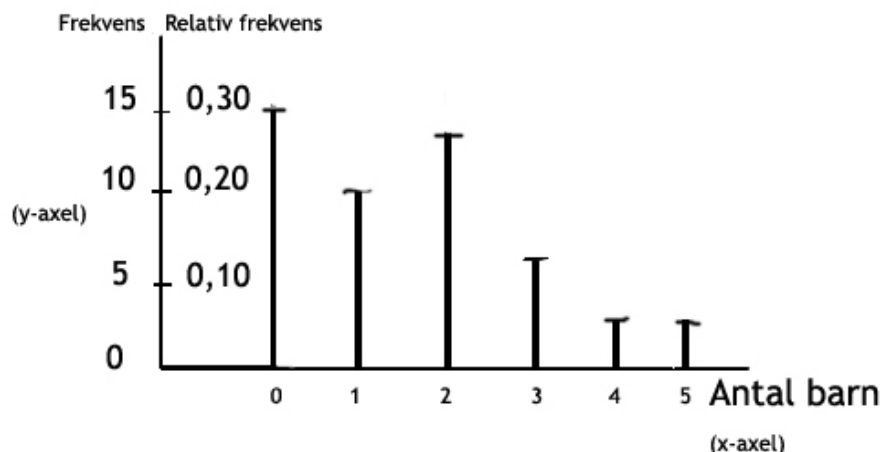


Diagram 1: *Antalet barn per familj*

Stolparnas längd i diagrammet anger frekvensen som vi kan läsa av på den lodräta y-axeln. På den vågräta x-axeln har antalet barn markerats. Informationen blir lättöverskådlig i diagramform, till exempel 15 familjer (titta på den lodräta y-axeln) har inga barn alls (värdet 0 på x-axeln som anger antalet barn).

Övning A.

I ett stickprov i årskurs två i gymnasieskolan, frågade man efter betyget i matematik. Följande data erhöles: 42 elever hade fått betyget 5; 30 elever hade betyget 3; 48 elever hade betyget 4; 37 elever hade betyget 2; 2 elever hade inte fått något betyg och 20 elever hade betyget 1. Beskriv undersökningsresultatet i en frekvenstabell med absolut frekvens och frekvens uttryckt i procenttal, samt rita ett diagram över materialet där (absolut) frekvens framgår. (Var noga med när du läser av data i uppgiften!)

Övning B.

Vi antar att vi frågar ett representativt urval av 15-åringar hur ofta de går på bio under ett år. Vi erhåller följande svar:

2 gånger per år	5 gånger per år	7 gånger per år
6 gånger	7 gånger	
1 gång per år	3 gånger	4 gånger
2 gånger	2 gånger	
3 gånger	4 gånger	1 gång per år
8 gånger	2 gånger	

Fyll i frekvenstabellen nedan (glöm inte rubriken ovanför tabellen!):

X	f	kumulativ f	relativ f
---	---	-------------	-----------

Observera: Då stickprovet är < 50 (mindre än 50), skall procent inte redovisas!

En grafisk beskrivning av den kumulativa frekvensen

I vårt exempel, undersökningen om antalet barn per familj, fanns en frekvenstabell som återges igen nedan. I frekvenstabellen kan vi nu även läsa den kumulativa frekvensen uttryckt i procent i kolumnen ytterst till höger.

Tabell x: Antal barn per familj

Antal barn (x)	Frekvens (f)	Kumulativ f	Relativ frekvens (f/n)	Frekvens i procent	kum.frekvens i %
0	15	15	0.30	30	30
1	10	25	0.20	20	50
2	13	38	0.26	26	76
3	6	44	0.12	12	88
4	3	47	0.06	6	94
5	3	50	0.06	6	100
n=50			1.0	100%	100%

Ett bra sätt att åskådliggöra kumulativ frekvens (som till sin karaktär är stigande), är att räkna om den till procent om stickprovet är tillräckligt stort ($n > 50$, $n = 50$) och rita en *trappstegskurva*. Nedan i Diagram 2, ser du ett trappstegdiagram över materialet.

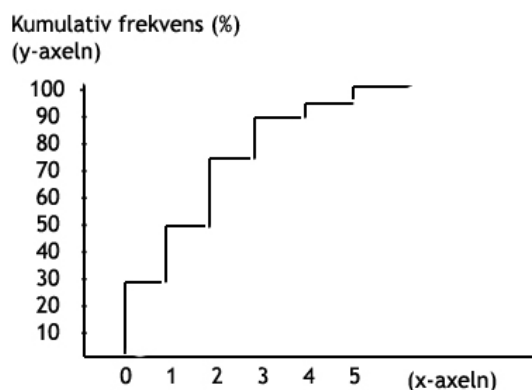


Diagram 2: Antalet barn per familj uttryckt i kumulativ frekvens (%).

FACIT – övningsuppgifterna A och B

Övning A.

I ett stickprov i årskurs två i gymnasieskolan, frågade man efter betyget i matematik. Följande data erhöles: 42 elever hade fått betyget 5; 30 elever hade betyget 3; 48 elever hade betyget 4; 37 elever hade betyget 2; 2 elever hade inte fått något betyg och 20 elever hade betyget 1. Beskriv undersökningsresultatet i en frekvenstabell med absolut frekvens och frekvens uttryckt i procenttal, samt rita ett diagram över materialet där (absolut) frekvens framgår. (Var noga med när du läser av data i uppgiften!)

Tabell x: Undersökningsgruppens
fördelning på matematikbetyg
i gymnasieskolan

Betyg	f	%
0	2	1
1	20	11
2	37	21
3	30	17
4	48	27
5	42	23

n = 179 (100%)

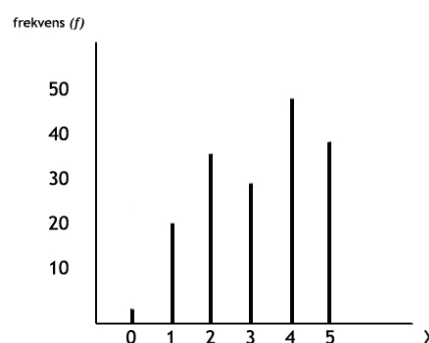


Diagram X: Undersökningsgruppens fördelning
på matematikbetyg i gymnasieskolan.

Övning B: frekvenstabell. Facit

Vi antar att vi frågar ett representativt urval av 15-åringar hur ofta de går på bio under ett år. Vi erhåller följande svar: 2 gånger per år, 5 gånger per år, 7 gånger per år, 6 gånger, 7 gånger, 1 gång per år, 3 gånger, 4 gånger, 2 gånger, 2 gånger, 3 gånger, 4 gånger, 1 gång per år, 8 gånger, 2 gånger

Fyll i frekvenstabellen nedan (glöm inte rubriken ovanför tabellen!):

Tabell B: Frekvensen av femtonåringarnas biobesök under ett år.

X	f	kumulativ f	relativ f
1	2	2	0,13
2	4	6	0,27
3	2	8	0,13
4	2	10	0,13
5	1	11	0,07
6	1	12	0,07
7	2	14	0,13
8	1	15	0,07
		n=15	1

Övning C:

Övning. Fyll i diagrammet i Göteborgs Postens artikel om stormen Gudrun från 2005 ! Jag har lämnat uppgifter till hjälp i det urklippta diagrammet! :)

Siffrorna som du behöver finner du längst ner, efter artikeln.

GÖTEBORGS-POSTEN MÅNDAG 28 NOVEMBER 2005

VÄSTSVENRIKES | 9

Elva döda i Gudruns spår

VÄXJÖ: Stormen Gudruns härjningar för snart ett år sedan har hittills kostat elva människor livet. De flesta har hamnat under rotvältor som slagit igen som fällor.

Redan i april gick Arbetsmiljöverkets generaldirektör Kenth Pettersson ut hårt och krävde att säkerhetskraven skulle följas vid arbete i den stormfällida skogen. Pettersson hotade till och med att stoppa arbetet. Då hade fyra personer omkommit inom loppet av någon vecka och Arbetsmiljöverket såg en farlig utveckling.

Men verkets värsta farhågor infriades inte.

Håkan Rosengren, biträdande tillsynsdirektör vid Arbetsmiljöverkets Växjödistrikt, förnekar att de krav som verket ställde

i våras var någon form av kompromiss med tanke på de stora värden som låg i de stormfällida skogarna.

– Vi kompromissar inte om arbetstäckens liv, säger han till TT.

– Den massiva information som vi gick ut med gjorde nytta. Folk blev mer försiktiga.

Vid sidan om dödsfallen har ändå hundratals skadat sig mer eller mindre allvarligt. Arbetsmiljöverket har fått in 128 anmälningar om skador, men mörkertalet är mycket stort.

– Vi får bara in anmälningar på sådana fall där det råder ett arbetsgivar-arbetsstämmandeförhållande.

För att få reda på exakt hur många som skadats i skogen skulle man behöva gå igenom journalerna vid samtliga vårdcentraler och akutmottagningar i de drabbade länen – ett sisyfosarbete.

– Men vi skulle gärna se att någon forskare tog sig an detta. Man skulle säkert kunna utläsa många intressanta saker ur resultaten, säger Håkan Rosengren.

Fler maskiner

Att stormskogen ändå inte krävt fler dödsoffer beror, enligt Rosengren, på den ökade mekaniseringen.

I dag sitter skogsarbetarna ofta skyddade i stora skogsmaskiner. Efter den stora stor-

Dödsolyckor i skogen

Antal dödade skogsarbetare i Sverige.



* T o m 25/11 -05 (varav 11 stormrelaterade)

Källa: Arbetsmiljöverket

svenska grafikbyrå/TT

FAKTA:

Väldsamaste ovädret på 35 år

Stormen Gudrun drog in över södra och västra Sverige på kvällen den 8 januari i år. Ovädret var det väldsamaste på 35 år.

Sju personer miste livet under ovädrets härjningar. Ytterligare minst elva personer har omkommit i röjningsarbetet efter stormen.

Stormen skapade omkring 180 000 kalhyggen. Totalt föllades cirka 75 miljoner kubikmeter skog.

Runt 730 000 elabonnenter blev strömlösa.

En kvarts miljon hushåll blev utan fast telefon.

Närmare 3 000 mil elnät skadades.

Försäkringsbolagens kostnader blev 4,2 miljarder kronor, vilket gör stormen Gudrun till den dyraste enskilda skadehändelsen i Sverige. (TT)

men 1969, som ändå orsakade mindre skador än Gudrun, omkom ett 40-tal skogsarbetare.

Vid sidan om dödsfallen och de fysiska skadorna får man också räkna in de psykiska problemen som drabbat enskilda skogsägare. Skogsstyrelsen gjorde inom ramen för sitt stormanalysprogram djupintervjuer med ett antal skogsägare som fått sina livsverk förstörda. De vittnade om sorg och vrede, men också om framtidstro.

Enda vinnarna på stormen är älgarna och möjligen de som äger timmerbilar. Det var helt enkelt för farligt för älgägarna att ge sig ut i de småländska skogarna.

LENNART HANSSON

11

1969	40
1975	20
1995	13
2000	7
2001	3
2002	10
2003	3
2004	1
2005	5

Lite om skalor och diagram

När en *kvalitativ variabel* är i fråga, används *nominal skala* eller *ordinal skala*. Kvalitativ variabel är inte att förväxla med kvalitativ metod som kanske bäst illustreras av intervjuer som analyseras. Skall en kvalitativ variabel illustreras används ofta *ytdiagram*, *cirkeldiagram*, *rektangeldiagram* eller i de fall variabeln är dikotom och *inte* består av talvärden, kanske med ett *stapeldiagram*. En *dikotom, kvalitativ variabel* kategoriserar något i två kategorier, t.ex. en "man" eller en "kvinna", eller antar antingen "ja" och "nej" alternativet.

I ett *linjediagram* visar man ofta värden på en variabel och hur de förändras över tid. Ett *ytdiagram* liknar linjediagrammet, men den har överlappande ytor eller staplade ytor. Ett *cirkeldiagram* kallas också "tårtdiagrammet". Cirkeldiagram rekommenderas när summan av alla värden utgör en helhet, till exempel det totala antalet elever i klassen (=100%). Cirkeln delas i lämpliga "tårtbitar" beroende på hur stor andel av den totala "tårtan/cirkeln" resultatet illustrerar. Cirkeldiagrammet används ofta när det inte finns någon tydlig ordning mellan "delarna" eller "tårtbitarna", i vilka man ofta också anger procentsatser eller annan mängdinformation. *Rektangeldiagram* liknar cirkeldiagrammet, förutom till formen.

Stolpdiagram används ungefär i samma syfte som stapeldiagram (histogram), men stolpdiagram visar frekvensen för olika variabel värden. För kumulativ frekvens används inte sällan en *summapolygon*.

När en kvantitativ, diskret variabel skall illustreras eller när frekvenser skall illustreras som i fallet ovan med kumulativa värden, används med fördel en *trappstegskurva*.

Kvantitativa kontinuerliga variabler illustreras ofta med ett *histogram* (också kallad stapeldiagram) och stora mängder sifferdata åskådliggörs bäst genom att man delar upp datan i klasser, t.ex. alla barn födda i januari, alla barn födda i februari, osv. Vi skall nedan titta lite på klassindelning.

Klassindelad material

Approximativt kontinuerlig variabel

Nedan exempel som återberättas fritt har hämtats ur Stig Fhanérs bok *Statistik – rätt upp och ned*. Stockholm: Almqvist och Wiksell (1980), s. 18-20.

Låt oss anta att vi har ett stickprov på 150 sjuåringar och vi testar deras intelligensnivå. Testet kan anta poäng från 78 till 131. Eftersom testet endast kan anta heltalsvärden är den egentligen en diskret variabel, men intelligens kan betraktas som en kontinuerlig "egenskap". Därför är vår variabel en *approximativt kontinuerlig variabel*.

Det är alltså frågan om den underliggande egenskapens karaktär som avgör om variabeln kan betraktas som approximativt kontinuerlig variabel, trots att mätvärden endast antar heltalsvärden!

Testet har stort variationsvidd. Med sina 78 och 131 poäng kan testet anta 54 olika värden och för att få bättre överskådlighet över resultaten, kommer vi att dela resultatvärden i ett antal klasser. Det går till så att vi bestämmer lika stora poängintervaller inom vilka vi för samman individer till klasser. I det här fallet bestämmer vi klasserna till att bestå av ca 9 poängenheter

vardera. Då blir poängvärden från 78 till 86 en klass, 87 till 95 en klass, osv. I frekvenstabellen (Tabell 2) nedan har vi beräknat frekvensen individer i varje klass.

Tabell 2: *Resultaten på intelligenstestet bland 150 barn på 7 år*

Klassintervall	frekvens (f)
78 – 87	2
87 – 95	35
96 – 104	74
105 – 113	24
114 – 122	7
123 – 131	8

Om vi hade skrivit materialet utan att först ha delat in det i klasser hade ovan Tabell 2 blivit oerhört mycket längre! Nu har vi presenterat materialet i komprimerad form. *Vi vinner överskådlighet, men vi förlorar något av informationens noggrannhet!*

Som ni förstår, går det inte ur Tabell 2 läsa vilket var det lägsta, respektive det högsta poängvärdet som något av barnen fick på testet. Vi kan endast läsa att två barn hamnade i klassen med de lägsta poängvärdena och åtta barn i klassen med de högsta poängvärdena. Det tål dessutom att funderas på om testet beträffande barnens intelligens verkligen är pålitligt och *säger något om den egenskap som man avser att mäta*, i det här fallet barnens intelligens. Man talar då om en test *validitet*.

Mittvärdet i en klass, så kallad *klassmitten*, får man genom att man tar det högre värdet i klassen och adderar med det lägre värdet i samma klass och dividerar summan därefter med två.

Klassmitten för den första klassen är sålunda 82. Det räknas ut så här: $(86 + 78)/2 = 82$.

Klassmitten beräknas som ett slags typiskt värde för klassen. Klassmitten anges då man ritar ett diagram över fördelningen.

Övning: Vilket blir klassmitten för de övriga klasserna? Räkna!

Här är första delen av statistikkompendiet slut. Andra delen finns i fotograferad form nedan tills bättre utskrift blir tillgängligt. I den fortsätter vi att tala om klassindelad material innan vi talar om elementära statistiska storheter, beräknar medelvärden, variationsmått och nämner korrelationer.