
3. MODELLER

Många använder modellbegreppet synonymt med teoribegreppet. Här använder vi modeller som hjälpmedel när det gäller att omvandla teoretiska påståenden till prövbara hypoteser och när det gäller att lösa specifika problem inom något verklighetsområde. En modell kan således vara identisk med en teori, men består oftast av delar från flera olika teorier. De modeller vi arbetar med här byggs ofta upp dels av teorier om det studerade verklighetsområdet och dels av teorier om de mätinstrument som används för att göra observationer. På så sätt kan vi säga att en modell beskriver vad vi bör observera enligt en eller flera teorier. I det här kapitlet diskuterar vi modellbegreppet något närmare och undersöker hur modeller kan användas för att beskriva händelser och fenomen samt som hjälpmedel för att analysera observationer. Det typiska för händelser som studeras inom statistiken är att "slumpen" har en viss betydelse. Ett exempel är ett företag som skall göra en försäljningsprognos på en av sina produkter. Man vet då inte på förhand exakt hur stor försäljningen kommer att bli, utan "slumpen" kommer att till viss del bestämma den försålda kvantiteten. För oss är det alltså särskilt intressant att studera modeller av situationer då "slumpen" inverkar på resultatet. Sådana modeller beskrivs i avsnitt 3.3. Innan vi tittar närmare på modeller behöver vi bekanta oss med några begrepp som ofta används i statistiska modeller, nämligen population, urval och variabel. Detta gör vi i kapitlets första avsnitt.

3.1 Population, urval och variabel

För att mer precist kunna diskutera syftet med statistiska undersökningar och de grundläggande teorierna om statistiska metoder måste vi införa begreppen population, urval och variabel.

Med *population* menar vi en mängd individer eller objekt (både konkreta och abstrakta) som har en eller flera gemensamma egenskaper. Låt oss betrakta några exempel.

Exempel 3.1

Mängden människor som är svenska medborgare ett visst bestämt datum utgör en population. Individernas gemensamma egenskap är att de är svenska medborgare det givna datumet. Individerna har också en mängd andra egenskaper, t. ex. ålder, kön och längd, som kan vara av intresse att studera. Dessa egenskaper kan dock vara olika för olika individer i populationen. ■

Exempel 3.2

Mängden av alla bostadshus i en viss kommun ett visst datum utgör en population. En individ i denna population är sålunda ett bostadshus. Egenskaper som inte nödvändigtvis är lika för alla individer i denna population är t. ex. bostadsyta och antal människor som bor i huset. ■

Exempel 3.3

Mängden av alla tärningskast som kan göras med en viss tärning utgör en population. Här är ett kast med tärningen en individ i populationen. Individerna i denna population är abstrakta. Antalet prickar som är på ovansidan tärningen är en egenskap som kan variera mellan olika kast, dvs mellan olika individer i populationen. ■

Observera att antalet individer i en population kan variera högst avsevärt. En population med ett ändligt antal individer brukar kallas en ändlig population till skillnad från en oändlig population som har oändligt många individer. Populationerna i exemplen 3.1 och 3.2 är alltså ändliga. Populationen i exempel 3.3 däremot är oändlig eftersom vi, åtminstone teoretiskt, kan göra hur många kast som helst med en tärning.

Det är viktigt att den population man vill studera är noga definierad så att det inte uppstår tveksamheter om en individ tillhör populationen eller ej. T. ex. är populationen i exempel 3.2 ofullständigt definierad. Vad menar vi med ett bostadshus? Räknas ett radhus som ett bostadshus eller som flera? Räknas ett hyreshus som ett bostadshus? Det är syftet med den analys som vi avser göra som får avgöra hur populationen skall definieras. Om vi t. ex. endast är intresserade av fristående villor är det förstås ingen idé att även låta radhus ingå i populationen.

De individer som vi faktiskt observerar är nästan alltid en del, och oftast en mycket liten del, av hela populationen. Om alla individerna i en population observeras, säger vi att vi gör en *totalundersökning* av populationen. I annat fall gör vi en *urvalsundersökning* och de utvalda individerna kallas *urval*.

Det är i och för sig möjligt att mäta längden på alla svenska män i åldern 20 till 25 år för att beräkna deras medellängd. Det vore dock betydligt lättare om det räckte med att mäta längden av individerna i ett relativt litet urval och använda dessa observationer för att dra slutsatser om medellängden hos individerna i hela populationen.

De egenskaper som kan vara olika hos olika individer i en population kallas *variabler*. Våra observationer är således observationer på en eller flera variabler. Då vi observerar en variabel som kan representeras med tal, t. ex. längd, pris och ålder, kallas variabeln *kvantitativ*. Ibland har man även anledning att studera variabler som inte svarar mot tal, t. ex. kön, ögonfärg och partisympati. En sådan variabel kallas *kvalitativ*.

Kvantitativa variabler kan vara kontinuerliga eller diskreta. Kontinuerliga kvantitativa variabler är sådana som, eventuellt inom vissa gränser, kan anta vilket värde som helst. Längden av en människa kan t. ex. vara 174 cm, 182.65 cm och 176.92378504 cm. Däremot kan inte kroppslängden vara negativ. Inom vissa gränser finns det sålunda inget tal som inte kan vara längden av en människa. Variabeln ”kroppslängd” är därför en kontinuerlig kvantitativ variabel. Diskreta kvantitativa variabler är sådana som bara kan anta vissa värden. Antal prickar som kommer upp då man kastar en tärning kan bara anta värdena 1, 2, 3, 4, 5 och 6 och är därför en diskret kvantitativ variabel. Eftersom det endast är kvantitativa variabler som kan vara kontinuerliga eller diskreta kommer vi fortsättningsvis att underförstått mena kontinuerliga kvantitativa respektive diskreta kvantitativa variabler då vi skriver kontinuerliga respektive diskreta variabler.

Övning 3.1

Försök att mer precist definiera följande populationer och avgör om de är ändliga eller oändliga.

a) Tomatplantorna i ett växthus då skörden hos varje planta under en viss tid mäts.

- b) Fembarnsfamiljer då antal flickor i familjen räknas.
- c) Kommunerna i riket då antalet invånare i kommunerna räknas.
- d) Alla 18-åringar i Sverige då man undersöker vilken utbildning hon/han har.

Övning 3.2

Avgör om variablerna i övning 3.1 är kvalitativa eller kvantitativa och i förekommande fall om de är kontinuerliga eller diskreta.

I en del sammanhang är det praktiskt att tilldela olika tal till olika kvalitativa egenskaper. T. ex. kan observationer på sympatier med ett visst politiskt parti tilldelas talet "1" medan övriga observationer på den variabeln, dvs. sympatier med något annat politiskt parti, tilldelas talet "0". Genom att addera alla observationer får vi då antal observerade sympatisörer med det studerade partiet. Denna kodifiering till tal av kvalitativa egenskaper innebär dock inte att egenskapen har blivit kvantitativ. Valet av tal för att representera sympatisörer respektive icke-sympatisörer är helt godtyckligt. Vi kan lika gärna tilldela sympatisörerna talet "0" och icke-sympatisörerna talet "1". Summan av alla observationerna blir i det fallet antal observerade icke-sympatisörer, som ju tillsammans med en uppgift om antal observationer, ger samma information som antalet sympatisörer. När vi observerar en kvantitativ variabel, som t. ex. kroppslängd, har vi ofta flera mätskalor att välja mellan, t. ex. centimeter, meter och tum. När mätskalan väl är bestämd är också de olika individernas mätvärden entydigt bestämda.

3.2 Något om modeller och deras användning

I det här avsnittet skall vi introducera modellbegreppet och diskutera hur man kan arbeta med modeller. Förmodligen har du tidigare i ditt liv arbetat med modeller. Tänk bara på din barndom, då du säkert hade dockor eller leksaker som du kallade bilar, båtar och flygplan. Att leksaker i form av bilar, båtar och flygplan är modeller accepterar vi nog utan att fundera så mycket, men är dockor modeller? Eftersom vi kan påstå att en leksaksbil verkligen är en modell, måste vi ha en uppfattning om vad en modell är och

vilka egenskaper en modell har. Låt oss tänka efter vad en leksaksbil är och vilka egenskaper den har. En leksaksbil är i vissa avseenden en kopia av en riktig bil, men det finns en rad viktiga skillnader mellan riktiga bilar och leksaksbilar. En leksaksbil är t. ex. mycket mindre och den har oftast ett enklare utförande än den riktiga bilen. Med enklare utföranden menar vi sådana saker som att dörrarna och bagageluckan inte kan öppnas, ratten inte kan vridas osv. Det gör ingenting att dörrarna och bagageluckan inte kan öppnas, för ingen skall sätta sig i leksaksbilen och ingen skall ha något bagage i bagageutrymmet. Allmänt är det så att leksaksbilen är förenklad på alla punkter där det inte gör något att den är förenklad. Om vi t. ex. skall köra leksaksbilen är det viktigt att den har en motor, men om vi inte skall köra den kan den lika gärna vara utan motor. Sammanfattningsvis kan vi säga att leksaksbilen är en förenkling av den riktiga bilen och att leksaksbilen är "naturlig" på just de punkter vi vill att den ska "bete sig" som en riktig bil.

Våra leksaksmodeller vi hittills har betraktat har alla varit förenklade avbildningar av någonting som finns i verkligheten och vårt användningsområde av modellen har bestämt vad som kan förenklas. Vi skall nu betrakta några andra situationer som inte är lek och finna att man även där har användning av förenklade avbildningar av verkligheten. Låt oss först betrakta flygplanskonstruktören som vill studera ett flygplans aerodynamiska egenskaper, dvs. se om flygplanet kan flyga. Konstruktionen kan förstås bygga ett riktigt flygplan, sätta sig i det, starta motorn och se vad som händer. Med en stor portion tur kommer flygplanet att lyfta och flyga och med ännu mera tur kan flygplanet också landas och fortfarande vara oskadd. Har konstruktionen otur störtar planet. Det är dock tveksamt om det finns några flygplanskonstruktörer som är nöjda med sin arbetssituation om de måste testa sina konstruktioner på det sättet. Finns det då någon alternativ arbetsmetod? Ja, tack och lov. Flygplanskonstruktionen kan bygga en modell av sin konstruktion, som till det yttre ser precis likadan ut som det riktiga flygplanet och som har precis samma viktfördelning. Sedan kan konstruktionen placera modellen i en vindtunnel och se vad som händer. Metoden har några uppenbara fördelar:

- det är mycket billigare att bygga en modell än att bygga ett riktigt flygplan
- konstruktionen behöver inte riskera sitt liv då konstruktionen ska testas

- det är lätt att göra förändringar i modellen för att prova olika alternativ i konstruktionen

En modell betyder alltså för oss en förenklad beskrivning av något verkligt och beskrivningen måste vara relevant för det vi ska använda modellen för. Vi har också möjligheten att ersätta verkligheten med symboler. Om vi t. ex. inte behöver ha en funktionsduglig ratt i en modellbil kan vi måla en ratt på instrumentbrädan. Den avmålade rattan är då en symbol för en verklig ratt.

Vi kommer att flitigt använda modeller med symboler i den här framställningen. Därför skall vi diskutera dessa begrepp ytterligare något. Antag att vi sitter i en bil och åker med konstant hastighet. Ju längre tid vi åker, desto längre sträcka färdas vi. Tydligt finns det ett samband mellan den tid och det avstånd vi färdas. Från fysiken vet vi att sambandet är

$$s = v \cdot t$$

där s betyder längden på den tillryggalagda sträckan, v hastigheten och t den tid vi har färdats. Denna "formel" använder symbolerna s , v och t för någonting som finns i verkligheten, nämligen tillryggalagd sträcka, färdhastighet och färdtid. Vidare är den en beskrivning av något verkligt, nämligen att om vi färdas med hastigheten v under tiden t , så hinner vi precis sträckan $s = v \cdot t$. "Formeln" är alltså en modell som beskriver tillryggalagd sträcka då vi färdas med en viss hastighet under en viss tid.

Med den terminologi vi introducerade i föregående avsnitt är storheterna sträcka, hastighet och tid variabler. Mer precist är variablerna kvantitativa och kontinuerliga. I modellen symboliseras variablerna med bokstäver, nämligen s , v respektive t . Vi kan alltså hantera symbolerna s , v , och t som om de vore vanliga tal. I modellen finns också en vanlig matematisk operation, nämligen multiplikation. Modeller som med hjälp av matematiska operationer beskriver hur variabler, symboliserade av bokstäver, beror av varandra, kallas matematiska modeller.

Varför har vi nu bemödat oss att diskutera matematiska modeller? Det finns flera anledningar till det. Som vi såg i de första exemplen på modeller, är det enklare att arbeta med modeller än med verkligheten själv. Detsamma gäller naturligtvis även för matematiska modeller. De matematiska relationerna "uttrycker" verklighetens relationer i en förenklad, kompakt och

användbar form. De matematiska relationerna är användbara för numeriska beräkningar. Med den matematiska modellen ovan är det möjligt att snabbt, enkelt och billigt att räkna ut hur lång sträcka vi hinner färdas om vi åker med en viss hastighet under en viss given tid. Alternativet att ta reda på det är att utföra experimentet i verkligheten, dvs. springa, cykla, åka bil etc. med den givna hastigheten under precis den givna tiden och sedan mäta sträckan hur långt vi hann. För att utföra längdmätningen behöver vi dessutom någon måttstock. Detta är betydligt krångligare än att utföra beräkningen med hjälp av modellen.

Övning 3.3

Diskutera om en fotomodell egentligen är en modell och i så fall vad den är en modell av.

Övning 3.4

Diskutera på vilket sätt modellen " $s = v \cdot t$ " är en förenkling av verkligheten.

3.3 Modeller av slumpmässiga försök

Begreppet *försök* innebär allmänt att utföra något slags handling och sedan invänta resultatet av den utförda handlingen. Ett exempel är att kasta ett mynt eller en tärning och notera resultatet av kastet, dvs. vilken sida av myntet eller tärningen som kommer upp.

För en del försök kan vi, innan försöket utförs, exakt bestämma vad försöket kommer att resultera i. Antag t. ex. att vi har en tärning vars sex sidor alla är märkta med en prick. Om vi kastar tärningen och sedan noterar antalet prickar på den sida som kommer upp, vet vi redan från början att försöket kommer att ge resultatet "en prick" därför att försöket alltid kommer att ge det resultatet. På motsvarande sätt kommer försöket "släpp en sten från en meters höjd" alltid att resultera i att stenen faller till marken. Ett försök som vi i förväg kan bestämma vad det kommer att resultera i kallas ett *deterministiskt försök*. För ett deterministiskt försök är det alltså alltid möjligt att förutsäga resultatet innan försöket utförs, bara man känner betingelserna för försöket.

Ofta kan man inte kontrollera försöket (eller betingelserna) så väl att det blir ett deterministiskt försök. Vid varje upprepning av försöket erhåller man ett resultat som kan varieras från gång till gång på ett sätt som man inte kan förutse innan försöket utförs (t. ex. upprepade kast med en normal tärning). Man säger att slumpen har inverkan på försöksresultatet och man säger att man utför ett *slumpmässigt försök*. I det här avsnittet skall vi diskutera modeller av slumpmässiga försök.

De olika resultat som ett försök kan resultera i kallas försökets olika *utfall*. Ett försök kan aldrig ge två olika utfall samtidigt. Bara ett utfall är möjligt vid varje upprepning av försöket. Om man kastar en vanlig tärning kan endast ett av utfallen "en prick", "två prickar", ... , "sex prickar" inträffa. Försöket kan inte resultera i två prickar och fem prickar samtidigt, så att "två prickar och fem prickar" är inte något utfall. Däremot kan utfallet "två prickar" inträffa vid ett kast och utfallet "fem prickar" inträffa i nästföljande kast.

Betrakta det slumpmässiga försöket att kasta ett mynt och se vilken sida av myntet som kommer upp. Vi vet inte vilket utfall det blir, men vi kan tala om vad det kan bli, dvs. vi kan räkna upp alla möjliga utfall. När vi kastar ett mynt blir det antingen krona eller klave. Om vi symboliserar utfallet krona med KR och klave med KL så är

KR KL

en uppräkningslista av alla utfall som kan inträffa när vi utför försöket.

En uppräkningslista av vad som kan inträffa vid ett slumpmässigt försök kallas för *utfallsrum*. Tydligt är utfallsrummet en viktig bit av en modell av ett slumpmässigt försök.

I fortsättningen kommer vi att låta Ω vara en symbol för utfallsrum och vi kommer att skriva utfallsrummet med mängdsymboler. I försöket med myntkastet får vi alltså

$$\Omega = \{KR, KL\}.$$

Övning 3.5

Skriv upp utfallsrummet för försöket att kasta en tärning. Hur många utfall finns det i utfallsrummet?

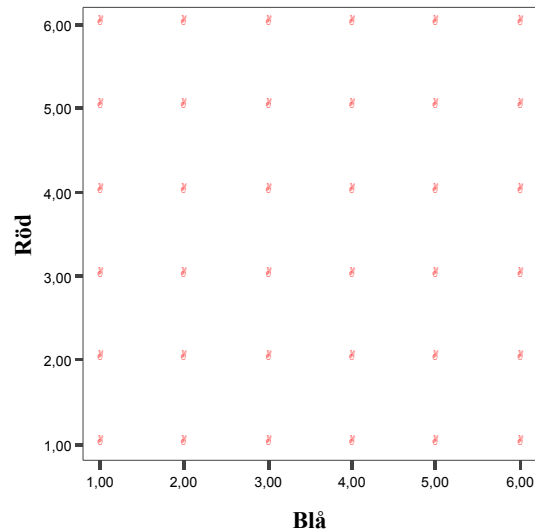


Figure 3.1: Utfallsrummet för kast med två tärningar kan illustreras med en tvådimensionell punktmängd.

Exempel 3.4

Antag att vi utför det slumpmässiga försöket att kasta en blå och en röd tärning och noterar antalet prickar på tärningarna. För att utreda hur utfallsrummet till detta försök ser ut måste vi ta reda på vad som kan inträffa. Ett utfall som kan inträffa är att den blå tärningen visar en prick och den röda tärningen visar tre prickar. Vi kan låta ett talpar, $(1,3)$, symbolisera det utfallet, så att det första talet visar antal prickar på den blå tärningen och det andra talet visar antal prickar på den röda tärningen. På detta sätt kan vi räkna upp alla utfall, från det att båda tärningarna visar en prick, $(1,1)$, till att båda tärningarna visar sex prickar, $(6,6)$. Detta ger oss utfallsrummet

$$\Omega = \{(1,1), (1,2), \dots, (6,6)\}$$

som alltså har $6 \cdot 6 = 36$ utfall. I det här fallet kan utfallsrummet åskådliggöras med en punktmängd i ett koordinatsystem, se figur 3.1. ■

Alla utfallsrum vi har diskuterat hittills har ett ändligt antal utfall. Exempelvis har försöket att singla en slant två utfall, försöket att kasta två tärningar har 36 utfall osv. Vi säger att utfallsrummet är *ändligt* om det finns ett ändligt antal utfall. Ett utfallsrum som har oändligt många utfall kallas *oändligt* utfallsrum.

Exempel 3.5

Om vi utför försöket att kasta ett mynt tills vi får klave för första gången, kan det inträffa att vi får klave redan i det första kastet, (KL), eller vid det andra kastet, (KR, KL), eller vid det tredje kastet, (KR, KR, KL), osv. Även om det inte är troligt att vi måste kasta myntet hundra gånger innan vi får klave för första gången, är det faktiskt möjligt, och utfallet med 99 stycken KR följt av ett KL måste tas med i utfallsrummet. Det är ännu mindre troligt att vi får klave första gången vid det tusende kastet, men även det är möjligt och motsvarande utfall måste alltså noteras i utfallsrummet. Tydligen får vi

$$\Omega = \{(KL), (KR, KL), (KR, KR, KL), \dots\}$$

dvs. det finns ett oändligt antal utfall i utfallsrummet. ■

Exempel 3.6

Betrakta försöket att tända en glödlampa och observera hur länge lampan lyser innan den går sönder. Antag att brinntiden ligger någonstans mellan 1200 och 1800 timmar. Utfallsrummet är då ett intervall av tal

$$\Omega = \{[1200, 1800]\}.$$

Trots att utfallsrummet i det här fallet är begränsat, är det oändligt, eftersom det finns oändligt många reella tal i intervallet $[1200, 1800]$. Om vi däremot bara hade möjlighet att mäta tiden i hela timmar skulle vi observera utfall från ett ändligt utfallsrum med exakt 601 utfall:

$$\Omega = \{1200, 1201, 1202, \dots, 1799, 1800\} \quad \blacksquare$$

I exempel 3.5 är utfallsrummet oändligt. Observera att vi kan skriva utfallen som en följd: $(KL), (KR, KL), (KR, KR, KL), \dots$. Vi säger då att vi har

ett uppräknligt antal utfall. Ett utfallsrum som har ett ändligt eller ett uppräknligt oändligt antal utfall kallas *diskret*. Ett utfallsrum som består av ett intervall av reella tal är ett *kontinuerligt* utfallsrum. Att mäta brinntiden för en glödlampa, som i exempel 3.6, är alltså ett slumpmässigt försök med ett kontinuerligt utfallsrum.

Övning 3.6

- a) Skriv upp utfallsrummet för det slumpmässiga försöket "Du spelar två tennismatcher mot din lärare och skriver upp resultatet av varje match". Är utfallsrummet ändligt eller oändligt?
- b) Skriv upp utfallsrummet för det slumpmässiga försöket "Du spelar ett set i tennis mot din lärare och skriver upp gameställningen". (Tiebreak tillämpas ej.) Är utfallsrummet ändligt eller oändligt?

Övning 3.7

Nedan beskrivs några slumpmässiga försök. Avgör om försökens utfallsrum är ändliga eller oändliga respektive diskreta eller kontinuerliga.

- a) En familj väljs slumpmässigt ut och antal barn noteras.
- b) En fembarnsfamilj väljs slumpmässigt ut och antal pojkar noteras.
- c) En fembarnsfamilj väljs slumpmässigt ut och antal barn noteras.
- d) En person väljs slumpmässigt ut och personens längd noteras.
- e) En person väljs slumpmässigt ut och personens ålder noteras (flera svar möjliga).
- f) En person väljs slumpmässigt ut och personens kön noteras.
- g) Snödjupet i Stockholm den 1 januari nästa år.

Vi har hittills endast diskuterat utfallsrum, dvs. gjort en beskrivning av vad som är möjligt att inträffa då ett slumpmässigt försök genomförs. En modell av ett slumpmässigt försök måste även innehålla något mer, nämligen en uppgift om hur troligt det är att de olika utfallen inträffar. En uppgift om hur troligt det är att ett utfall inträffar kallas sannolikheten att utfallet

inträffar. En fullständig modell över det slumpmässiga försöket att kasta ett mynt och se vilken sida som kommer upp kan se ut så här:

$$\left\{ \begin{array}{l} \Omega = \{KR, KL\} \\ \text{"Det är lika troligt att det blir KR som KL"} \end{array} \right. .$$

Övning 3.8

Gör en fullständig modell av det slumpmässiga försöket i övning 3.5.

Övning 3.9

Gör en fullständig modell av de slumpmässiga försöken i övning 3.6.

Vi har nu sett att det oftast är möjligt, och ibland till och med lätt, att skriva upp utfallsrummet för ett slumpmässigt försök. Däremot är det ofta betydligt svårare att ange hur troligt det är att de olika utfallen inträffar, dvs. att ange sannolikheter. I nästa två avsnitt diskuterar sannolikhetsbegreppet lite närmare och något om de problem som är förknippade med sannolikhetsbegreppet.

3.4 Sannolikhet av ett utfall och en händelse

I föregående avsnitt såg vi att sannolikheter spelar en viktig roll i modeller av slumpmässiga försök. Sannolikheten för ett utfall är ett tal som uttrycker hur troligt det är att utfallet inträffar. För detta ändamål använder vi ett antal symboler. Vi låter e symbolisera ett utfall och $P(e)$ symbolisera talet som anger hur troligt det är att utfallet e inträffar. Exempelvis så låter vi $P(KL)$ vara det tal som anger hur troligt det är att vi får utfallet KL då vi utför det slumpmässiga försöket "kast med ett mynt". Talet $P(e)$ kallar vi för sannolikheten att e inträffar.

I fortsättningen kommer vi inte bara att prata om utfall utan även om *händelser*. En händelse är en delmängd av utfallsrummet. En händelse kan t. ex. vara att vi får ett udda antal prickar då vi kastar en tärning. Den händelsen består då av utfallen "en prick", "tre prickar" och "fem prickar". Om A symboliserar händelsen "udda antal prickar", så symboliserar $P(A)$

sannolikheten för händelsen att få ett udda antal prickar då vi utför det slumpmässiga försöket att kasta en tärning.

3.5 Tolkningar av sannolikheter

3.5.1 Den klassiska sannolikhetsdefinitionen

Frågar vi någon hur stor sannolikheten är att få klave när ett välbalanserat mynt kastas, svarar förmodligen de flesta 50 % eller $1/2$. Varför tycker vi att det är ett rimligt svar? Ett argument är att eftersom det totala antalet utfall är två, så måste sannolikheten för ett av utfallen vara $1/2$. Skriver vi detta resonemang i symboler, får vi

$$P(e) = \frac{1}{\text{antal utfall i } \Omega}.$$

Detta kallas för den *klassiska sannolikhetsdefinitionen*. Den klassiska sannolikhetsdefinitionen för en händelse blir analogt

$$P(A) = \frac{\text{antal utfall som gör att } A \text{ inträffar}}{\text{totala antalet utfall i } \Omega}.$$

Exempel 3.7

Antag att vi vill beräkna sannolikheten att få ett udda antal prickar när vi kastar en tärning. Vi vet att det finns tre utfall, "en prick", "tre prickar" och "fem prickar", som resulterar i händelsen "udda antal prickar". Vi vet också att det finns sex utfall i utfallsrummet. Detta ger oss att

$$P(\text{udda antal prickar}) = \frac{3}{6} = 0,5. \quad \blacksquare$$

Att beräkna sannolikheten att en händelse inträffar enligt den klassiska sannolikhetsdefinitionen handlar alltså mycket om att räkna antal utfall som leder till en viss händelse och att räkna totala antalet utfall i utfallsrummet. Sådana räkningar brukar kallas för kombinatorik. Vi kommer här inte att närmare gå in på den kombinatoriska teorin.

Innan vi avslutar det här avsnittet skall vi peka på en svårighet vi råkar ut för om vi ohämmat tillämpar den klassiska sannolikhetsdefinitionen. Betrakta det slumpmässiga försöket att kasta två tärningar. Låt A vara händelsen att pricksumman blir 5. Det finns då minst två sätt att beräkna antalet utfall som leder till händelsen A . Ett sätt är att låta utfallsrummet vara $\Omega = \{(1, 1), (1, 2), \dots, (6, 6)\}$ som i exempel 3.4. Utfallen som ger pricksumman 5 är då $(1, 4), (2, 3), (3, 2)$ och $(4, 1)$, dvs. fyra utfall ger A . Eftersom totala antalet utfall är 36 får vi $P(A) = 4/36 = 1/9$. Vi kan också hävda att utfallsrummet är $\Omega = \{2 \text{ prickar}, 3 \text{ prickar}, \dots, 12 \text{ prickar}\}$, dvs 11 möjliga utfall och ett av dessa är det som ger händelsen A . Alltså är $P(A) = 1/11$.

Den klassiska sannolikhetsdefinitionen ger tydligen olika svar på sannolikheten beroende på hur vi definierar utfallsrummet och de utfall som ingår i det. För att klara av detta problem måste vi kräva att utfallen måste väljas så att de är "lika möjliga", dvs. att de är lika sannolika. Problemet är nu att vi då måste använda sannolikheter för att definiera sannolikheter, dvs. vi får ett cirkelresonemang. Ett försök att rädda den klassiska sannolikhetsdefinitionen är genom att använda den så kallade *indifferensprincipen*. Den säger att om det inte finns något skäl att förvänta något speciellt utfall kan vi säga att utfallen är lika möjliga och den klassiska sannolikhetsdefinitionen kan användas. Efterhand som vi eventuellt får kunskaper om försöket kan det inträffa att det inte längre finns skäl att anse att utfallen är lika möjliga, vilket ånyo innebär att vi inte kan använda den klassiska sannolikhetsdefinitionen.

Övning 3.10

Utfallsrummet för försöket "välj slumpmässigt en trebarnsfamilj och observera könen hos barnen" kan bestå av utfallen "bara flickor", "bara pojkar", "två flickor och en pojke" samt "en flicka och två pojkar". En sådan uppdelning är emellertid inte lämplig då den klassiska sannolikhetsdefinitionen skall tillämpas eftersom dessa utfall inte är lika möjliga. Tar vi däremot hänsyn till den ordning barnen föds i, kan utfallen $FFF, FFP, \dots, PPP$ anses vara lika möjliga. Detta förutsätter att det är lika möjligt att det föds en pojke som att det föds en flicka, vilket dock inte är fallet i verkligheten. Beräkna nu med hjälp av den klassiska sannolikhetsdefinitionen sannolikheten för följande händelser: Familjen har

a) minst två pojkar, b) högst två pojkar, c) exakt en pojke samt d) exakt två pojkar.

Övning 3.11

Ett slumpmässigt försök består i kast med två tärningar. Sök sannolikheterna för följande händelser, om varje utfall tilldelas samma sannolikhet

- a) antalet prickar på båda tärningarna är jämt.
- b) antalet prickar på åtminstone en av tärningarna är jämt.
- c) summan av antalet prickar på tärningarna är jämt.

Övning 3.12

Skriv upp en modell av det slumpmässiga försöket ”kasta en tärning och notera antalet prickar som kommer upp”. Ange sannolikheterna för utfallen med hjälp av den klassiska sannolikhetsdefinitionen. Vad är sannolikheten för utfallet ”sex prickar”?

3.5.2 Relativa frekvensens stabilitet

Vid det här laget börjar vi bli väl förtrogna med det slumpmässiga försöket att singla slant. Vi betraktar försöket än en gång för att demonstrera den sannolikhetsdefinition vi skall diskutera i detta avsnitt. Om vi gör en lång serie försök och noterar varje utfall kommer vi att få en till synes helt regellös följd av utfall. Exempelvis kan vi få

KR, KR, KL, KR, KL, KL, KR, KL, KL, KL...

Den relativa frekvensen för utfallet KL vid n stycken kast med ett och samma mynt definieras som

$$R_n(\text{KL}) = \frac{\text{antal gånger KL inträffar}}{n}.$$

Om vi nu analyserar den relativa frekvensen för utfallet KL kommer vi snart underfund med att det finns en hel del information att hämta. För ovanstående följd av försöksresultat får vi följande följd av relativa frekvenser

n	1	2	3	4	5	6	7	8	9	10
$R_n(KL)$	0,00	0,00	0,33	0,25	0,40	0,50	0,43	0,50	0,56	0,60

I början har följden av relativa frekvenser kraftiga variationer, men dessa blir mindre ju längre kastserien blir. Den relativa frekvensen stabiliserar sig kring ett visst tal då antalet kast ökar.

Om vi gör en ny kastserie kommer förloppet att bli något annorlunda, men det tal som den relativa frekvensen stabiliserar sig kring är detsamma. Denna stabilitetsegenskap hos den relativa frekvensen brukar benämnas *relativa frekvensens stabilitet*. Det tal som den relativa frekvensen $R_n(KL)$ stabiliserar sig kring tolkar vi som sannolikheten $P(KL)$. Sannolikheten för en händelse tolkas alltså i det här fallet som hur ofta händelsen i medeltal inträffar i ett mycket stort antal oberoende och identiska upprepningar av ett slumpmässigt försök.

Exempel 3.8

Vi har anledning att tro att det är lika troligt att en person är född den 10:e som den 20:e dagen i en månad. Däremot är det mindre troligt att en person är född den 29:e som vanligen inte finns i februari, eller den 30:e som aldrig finns i februari eller 31:a som saknas i fem månader. Figur 3.2 visar relativa antalet personer födda någon av de första tio dagarna i en månad vid successiv genomgång av ett mycket stort register över personers födelsedata. Kurvan har samma allmänna förlopp som då vi kastar ett mynt, dvs den relativa frekvensen varierar starkt i början av sekvensen men stabiliserar sig efterhand kring ett visst tal. I det här exemplet stabiliserar sig den relativa frekvensen kring ett värde som är något större än $1/3$. ■

3.5.3 Sannolikhet som grad av tilltro

Betrakta det slumpmässiga försöket ”du spelar en tennismatch mot din lärare och noterar vem som vinner”. Hur stor är sannolikheten att du vinner? För att bestämma det kan vi inte använda den klassiska sannolikhetsdefinitionen, för vi kan knappast påstå att utfallen är lika möjliga. Vi kan heller inte bestämma sannolikheten genom att upprepa försöket många gånger och studera den relativa frekvensen. I och för sig kan du och din lärare spela

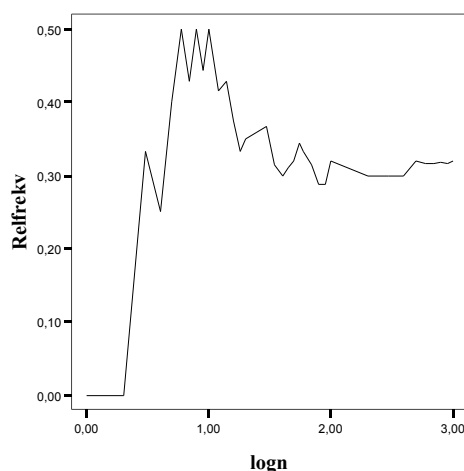


Figure 3.2: Relativa frekvensen personer i ett register födda 1-10 i en månad. Logskala används på den vågräta axeln så att 1 svarar mot $n = 10$, 2 svarar mot $n = 100$ och 3 svarar mot $n = 1000$.

många tennismatcher mot varandra, men matcherna kommer inte att ske under samma betingelser. Två olika matcher blir två olika slumpmässiga försök snarare än två upprepningar av samma försök.

I dagligt tal använder vi ofta sannolikheter för att uttrycka vår osäkerhet utan att vi baserar osäkerheten på relativa frekvensen av ett stort antal upprepningar av ett försök. Sannolikheter skall i det här fallet snarare tolkas som en personlig bedömning av vår osäkerhet om möjliga utfall. Denna bedömning kan baseras på vår personliga övertygelse eller vår kunskap om ett faktiskt sakförhållande, men inte på egenskaper hos en lång rad av upprepningar. Sådana sannolikheter kallar vi för subjektiva sannolikheter och uttrycker sålunda en grad av tilltro till att en viss händelse skall inträffa.

En tolkning av den subjektivt bestämda sannolikheten att du vinner en tennismatch mot din lärare är 0,8 är att säga att "Mot bakgrund av vad jag tror och vad jag vet, är jag beredd att ingå ett vad om att vinna tennismatchen. Jag är så säker att jag satsar 8 kr mot 2 kr". Vi säger då att din subjektivt bestämda sannolikhet att du vinner är $\frac{8}{8+2} = 0,8$. Alternativt kan vi säga

att den insats jag är beredd att ge (8 kr) är lika med sannolikheten (0,8) multiplicerad med vinsten jag kan vinna i vadet (8 + 2 kr).

En tolkning av den subjektivt bestämda sannolikheten $P(A)$ är alltså "Jag är beredd att ingå ett vad om att händelsen A inträffar (mot att händelsen A inte inträffar) med en insats som är proportionell mot $P(A)$ mot en insats som är proportionell mot $1 - P(A)$. Subjektiva sannolikheter är därför användbara för att uttrycka personliga övertygelser om vad som kommer att inträffa och vad som inte kommer att inträffa. Subjektiva sannolikheter bestäms alltså utifrån vår egen intuition, kunskap, erfarenhet och utbildning, men också av vår egen personlighet, våra åsikter, förhoppningar och farhågor.

Gränsen mellan sannolikheter som grad av tilltro och de andra två sannolikhetsbegreppen kan ibland te sig suddig. När t. ex. en prognos om vädret nästa dag rapporterar att sannolikheten för regn är 0,40 kan det betyda att i 40 % av fallen då väderleksläget har varit identiskt med det nuvarande så har det blivit regn nästa dag. Alternativt kan påståendet vara ett utslag av meteorologens personliga uppfattning och erfarenheter om liknande vädersituationer. När någon gör ett uttalande om en sannolikhet är det därför nödvändigt att undersöka vilken sorts tolkning man kan ge åt uttalandet.

Övning 3.13

Diskutera relationerna mellan subjektiva sannolikheter och odds på en travbana.

3.5.4 Kolmogorovs axiom

För att de sannolikheter vi definierar för olika händelser skall vara meningsfulla måste de uppfylla vissa rimlighetskrav. Vi vill t. ex. att sannolikheten för en händelse som alltid inträffar (inträffar med säkerhet) skall vara lika med ett. Detta krav, tillsammans med två andra krav som vi skall diskutera längre fram, utgör vad som brukar kallas Kolmogorovs axiom.

Låt först Ω inte bara betyda utfallsrummet, utan även händelsen att något av utfallen i utfallsrummet inträffar. Eftersom Ω är en uppräkningslista av alla

möjliga utfall måste händelsen Ω alltid inträffa. Om det finns n stycken utfall i Ω får vi enligt den klassiska sannolikhetsdefinitionen

$$P(\Omega) = \frac{n}{n} = 1$$

Man kan visa att samma egenskaper gäller för sannolikheter definierade som relativa frekvenser, se övning 3.14.

Det tal som står i täljaren i den klassiska sannolikhetsdefinitionen måste vara ett positivt heltal eller noll eftersom det talet anger att antal. Detsamma gäller för talet i nämnaren. Därav drar vi slutsatsen att sannolikheten för en händelse måste vara större än eller lika med noll. Skriver vi detta i symboler får vi

$$P(A) \geq 0$$

för en händelse A . I övning 3.14 visar vi att motsvarande resultat gäller även för sannolikheter definierade som relativa frekvenser.

Låt nu A och B vara två händelser som inte har något utfall gemensamt. Om A inträffar kan då B inte inträffa och vice versa. Sådana händelser kallas ömsesidigt uteslutande. Det antal utfall som gör att A eller B inträffar måste då vara lika med antalet utfall som gör att A inträffar plus det antal som gör att B inträffar. Sannolikheten att A eller B inträffar blir då enligt den klassiska sannolikhetsdefinitionen

$$\begin{aligned} P(A \text{ eller } B) &= \frac{\text{antal utfall som gör att } A \text{ eller } B \text{ inträffar}}{\text{totala antalet utfall i } \Omega} \\ &= (\text{antal utfall som gör att } A \text{ inträffar} \\ &\quad + \text{antal utfall som gör att } B \text{ inträffar}) / \text{totala antalet utfall i } \Omega \\ &= \frac{\text{antal utfall som gör att } A \text{ inträffar}}{\text{totala antalet utfall i } \Omega} \\ &\quad + \frac{\text{antal utfall som gör att } B \text{ inträffar}}{\text{totala antalet utfall i } \Omega} \\ &= P(A) + P(B) \end{aligned}$$

Även för dessa resultat finns motsvarigheter för sannolikheter definierade som relativa frekvenser vilket visas i övning 3.14.

Vi har därmed visat följande tre egenskaper hos sannolikheter enligt den klassiska definitionen och för sannolikheter definierade som relativa frekvenser:

$$\begin{aligned} P(\Omega) &= 1 \\ P(A) &\geq 0 \\ P(A \text{ eller } B) &= P(A) + P(B) \text{ då } A \text{ och } B \text{ är ömsesidigt uteslutande.} \end{aligned}$$

Dessa tre egenskaper kallas Kolmogorovs axiom. Det är inte helt självklart att sannolikheter definierade som grad av tilltro uppfyller Kolmogorovs axiom. Däremot är det möjligt att med vissa tilläggsantaganden, t. ex. på hur man får formulera sina bud i en vadslagning, kan få sannolikheter som grad av tilltro att uppfylla Kolmogorovs axiom.

I de följande kapitlen skall vi utgående från Kolmogorovs axiom visa räkneregler för sannolikheter och utföra olika sannolikhetsberäkningar. Vi måste därför kräva att den sannolikhetsdefinition vi använder uppfyller Kolmogorovs axiom. Räkneregler och beräkningarna kommer då att vara oberoende av hur sannolikheterna är definierade.

Övning 3.14

Övertyga dig om att sannolikheter definierade som relativa frekvenser uppfyller Kolmogorovs axiom.

3.6 Osannolika händelser och otroliga händelser

Om vi vet att sannolikheten att en händelse A inträffar är liten (men inte nödvändigtvis noll), kan vi tydligen dra slutsatsen att A sällan inträffar. Vid en försöksserie där ett försök utförs per dag och sannolikheten att A inträffar är 0,001, kommer A att inträffa i genomsnitt en gång vart tredje år. Är sannolikheten 0,000001 inträffar A i genomsnitt en gång på 3000 år osv.

I många tillämpningar av sannolikheteasteori spelar händelser med små sannolikheter, osannolika händelser, en viktig roll. Om man kan visa att en händelse är osannolik kan man med "praktisk säkerhet" räkna med att händelsen inte kommer att inträffa vid ett enda eller ett litet antal upprepningar av försöket. Observera dock det faktum att en osannolik händelse inte är det samma som en omöjlig händelse. En osannolik händelse kan mycket väl

inträffa, trots att sannolikheten att den skall inträffa är liten. I många slumpmässiga försök är det till och med så att det alltid är en osannolik händelse som inträffar. Betrakta t. ex. ett lotteri med 10000 lotter. Sannolikheten att en viss lott skall bli vinstlott är då 0,0001, dvs händelsen att just den lotten blir en vinstlott är tämligen osannolik. Å andra sidan är det alltid en lott med just den sannolikheten att vinna som faktiskt vinner lotteriet!

För att avgöra om det är troligare att en händelse A inträffar än en annan händelse B , kan vi jämföra sannolikheterna $P(A)$ och $P(B)$. Om t. ex. $P(A)$ är större än $P(B)$ säger vi att A är troligare än B . Alternativt kan vi undersöka kvoten

$$L = \frac{P(A)}{P(B)}.$$

Om L är större än ett så måste $P(A)$ vara större än $P(B)$, dvs A är troligare än B , och omvänt, om L är mindre än ett så måste $P(A)$ vara mindre än $P(B)$, dvs. B är troligare än A . För lika sannolika händelser är L lika med ett.

I exemplet med lotteriet har alla lotter samma sannolikhet att bli vinstlott. Därför är L lika med ett då vi jämför två olika lotter, dvs. alla händelser är lika troliga. Antag nu i stället att en person, Alexander Lucas, köper en lott och en annan person, Joakim von Anka, köper alla övriga 9999 lotter. Låt A vara händelsen att Alexander Lucas vinner och B händelsen att Joakim von Anka vinner. Då är

$$L = \frac{P(A)}{P(B)} = \frac{0,0001}{0,9999} \approx 0.0001.$$

Händelsen A är alltså inte bara osannolik, den är också mycket otrolig jämfört med händelsen B .

I sannolikhetslärans tillämpningar är det ofta intressant att jämföra sannolikheter och bedöma händelsernas troligheter (på engelska likelihood) i förhållande till varandra. Genom sådana jämförelser kan troliga händelser sorteras fram och otroliga händelser kan lämnas därhän.

3.7 Några exempel på modeller av slumpmässiga försök

3.7.1 Förekomsten av en gen

Antag att $p \cdot 100\%$ av individerna i en population har en viss gen och låt oss kalla genen A . Det innebär att om vi slumpmässigt väljer ut en individ ur populationen kommer den valda individen att antingen ha genen eller vara utan den. Utfallsrummet för detta försök blir därför

$$\Omega = \{A, \overline{A}\},$$

där A betyder att individen har genen och $\overline{A}$ betyder att individen inte har genen. Det finns alltså två möjliga utfall, A och $\overline{A}$, och det är de utfallen som räknas upp i utfallsrummet. Sannolikheterna för de två utfallen är

$$\begin{aligned} P(A) &= p \\ P(\overline{A}) &= 1 - p. \end{aligned}$$

Sannolikhetsmodellen, bestående av utfallsrum och sannolikheter, är densamma oavsett tolkning av sannolikheter. Däremot är tolkningen av modellen olika för den frekventistiska och den subjektiva ansatsen. Med frekventistisk tolkning av sannolikheter säger modellen att i ett stort antal identiska upprepningar av försöket att slumpmässigt välja en individ ur populationen och se om den valda individen har genen eller ej, kommer $p \cdot 100\%$ av *försöken* att resultera i att vi får en individ som har genen. Om vi däremot har en specifik individ framför oss kan vi inte göra något frekventistiskt sannolikhetsuttalande huruvida individen har genen eller inte. Sannolikheten är ju då antingen 1, vilket inträffar om individen har genen, eller 0, vilket inträffar om individen saknar genen.

Ett subjektivt sannolikhetsuttalande uttrycker vad vi själva tror beträffande om en specifik individ har genen eller ej. I detta fall kan en person hänvisa t.ex. till sin "magkänsla" och hävda att sannolikheten är nära 100 % att individen har genen, medan en annan person kan med motsvarande argument hävda motsatsen. Båda uttalandena är lika riktiga, eftersom de uttrycker personernas subjektiva uppfattningar, och uppfattningar kan ju naturligtvis vara helt olika.

3.7.2 Att mäta avstånd

Antag att avståndet mellan två punkter är μ meter. När avståndet skall mätas med hjälp av ett mätinstrument kan det inträffa att det uppmätta avståndet inte blir exakt μ utan någonting i närheten. Anledningen till att vi inte exakt observerar avståndet beror på att vi får ett visst så kallat mätfel. Mätfel uppkommer av många olika orsaker, t.ex. kan det finnas en osäkerhet i mätinstrumentet. Mätinstrumentet kan ge oss olika resultat beroende på den temperatur, luftfuktighet etc. som råder vid mättillfället.

I en frekventistisk ansats skulle man säga att det finns ett visst fixt avstånd, μ , mellan punkterna och att när vi mäter uppstår ett slumpmässigt mätfel. Låt oss beteckna mätfelet med ε . Det innebär att det mätresultat vi får är summan av det verkliga, sanna, avståndet och mätfelet. Vi kan skriva det som $\mu + \varepsilon$. En mycket enkel modell för mätfelet är

$$\Omega = \{\varepsilon = -0,1; \varepsilon = 0,0; \varepsilon = 0,1\}$$

$$P(\varepsilon = -0,1) = 0,3$$

$$P(\varepsilon = 0,0) = 0,4$$

$$P(\varepsilon = 0,1) = 0,3$$

I den här mycket förenklade modellen säger vi att mätfelets storlek är ett av tre möjliga värden, 0,1 meter, 0 meter (alltså inget mätfel) respektive 0,1 meter. Den frekventistiska tolkningen av sannolikhetsuttalandena är att om mätningen utförs ett stort antal gånger kommer det observerade mätvärdet att vara 0,1 meter kortare än det sanna värdet ($\varepsilon = -0,1$) i 30 % av mätningarna, det kommer att vara lika med det sanna värdet ($\varepsilon = 0$) i 40 % av mätningarna och det kommer att vara 0,1 meter längre än det sanna värdet ($\varepsilon = 0,1$) i 30 % av mätningarna. För en viss given mätning kan vi dock inte ange någon frekventistisk tolkning av mätfelet och dess storlek.

En subjektiv ansats ger uttalanden om vad modelleraren själv tror om avståndet. Det är därför möjligt att tänka på avståndet som slumpmässigt så att en modell skulle kunna vara

$$\Omega = \{\mu = 9,9; \mu = 10,0; \mu = 10,1\}$$

$$\begin{aligned}
P(\mu = 9, 9) &= 0,3 \\
P(\mu = 10, 0) &= 0,4 \\
P(\mu = 10, 1) &= 0,3
\end{aligned}$$

En tolkning av den här modellen är att den som skriver modellen tror att avståndet är 10 meter och 0,4 är ett mått på hur säker hon eller han är på det uttalandet. Som alternativ tror modelleraren att avståndet är 9,9 meter eller 10,1 meter och ett mått på hur säker hon eller han är på det är 0,3 för respektive avstånd. Ett sätt att tolka måtten på säkerhet är att modelleraren kan tänka sig ingå ett vad och där satsa 40 % av en tänkt vinst på att avståndet är 10 meter. Motsvarande insatser gäller för avstånden 9,9 och 10,1 meter.