
17. PRÖVNING AV HYPOTESER OM POPULATIONSPARAMETRAR

17.1 Grundläggande begrepp

Vi har tidigare behandlat hur observationer från ett urval kan användas för att uppskatta populationsparametrar. I detta kapitel ska vi introducera en annan typ av inferens där observationer från ett urval används för att avgöra om en hypoteser om populationsparametrar kan förkastas eller ej. Sådana avgöranden är ofta viktiga då beslut av olika slag skall fattas. Några exempel här är problemet att avgöra om ett parti av en viss produkt innehåller en viss utlovad kvalitet, om en tillverkningsprocess fungerar som planerat och om en viss behandling av en viss sjukdom har någon effekt på patienter.

De statistiska metoderna som används för att fatta denna typ av beslut baseras på två komplementära antaganden om populationen. Partiets produkter har antingen den utlovade kvaliteten eller så har de den inte. Antingen fungerar tillverkningsprocessen som planerat eller så gör den det inte. Antingen har behandlingen effekt på patienterna eller så har den det inte. Dessa antaganden kallas vanligen *hypoteser* och metoderna att fatta beslut av det slag vi har exemplifierat ovan, baserade på observationer från ett urval, kallas *hypotesprövning*. I detta kapitel diskuterar vi metoder för att pröva hypoteser om populationsparametrar.

Vi introducerar först ett antal begrepp och utgår då från ett exempel. Antag att de ekonomiska kalkyler ett företag gör för en viss produkt bygger på att tillverkningskostnaderna är 200 kr per enhet. Genom en rad osäkerhetsfaktorer kommer tillverkningskostnaden att variera för olika tillverkade enheter. Exempel på sådana faktorer är att tillverkningstiden kan variera, kvaliteten i de använda komponenterna varierar osv. Vad kalkylen bygger på är att tillverkningskostnaden i medeltal, sett över hela populationen av tillverkade enheter, är 200 kr. Antag också att man har använt sådana mar-

ginaler i kalkylerna att en viss avvikelse från 200 kr kan tillåtas, men om den sanna tillverkningskostnaden är 220 kr måste kalkylerna revideras.

Vår uppgift är nu att på basis av observationer på tillverkningskostnaden för ett urval av tillverkade enheter avgöra om (1) vi kan behålla de uppgjorda kalkylerna eller (2) vi måste revidera kalkylerna. Eftersom vi observerar endast ett urval av enheter, finns en viss osäkerhet beträffande den sanna tillverkningskostnaden. Det beslut vi slutligen fattar beror på flera frågeställningar:

1. Vilka kostnader är förknippade med ett felaktigt beslut? Ett beslut är felaktigt om det är annorlunda det beslut vi skulle ha fattat om vi inte hade haft någon osäkerhet. I vårt exempel är ett felaktigt beslut antingen att (a) besluta om att anpassa kalkylerna till tillverkningskostnaden 220 kr fastän medeltillverkningskostnaden är 200 kr, eller (b) besluta om att inte revidera kalkylerna fastän medeltillverkningskostnaden är 220 kr. Det är här viktigt att påpeka att den osäkerhet som eventuellt kan leda till felaktigt fattade beslut av de här slagen enbart beror på det faktum att vi baserar beslutet på ett urval. Kostnaderna för de olika felen kan uttryckas på flera olika sätt. I vårt exempel är det kanske lämpligt att uttrycka kostnaderna i ekonomiska termer.
2. Vilka sannolikheter för de två olika felen i punkt 1 är acceptabla? Detta övervägande är givetvis förknippat med hur svårartade vi bedömer att de olika felen är.
3. När ger ett urval stöd för ett visst beslut eller för en viss hypotes om populationen? Vi måste utveckla en regel så att vi kan avgöra vilket beslut vi skall fatta eller vilket antagande om populationen vi skall tro på. Detta beslut måste vara konsistent med de sannolikheter för felaktigt fattade beslut vi har bestämt oss för.

Vi formaliserar nu de begrepp vi har introducerat. De antaganden om populationen vi arbetar med kommer vi fortsättningsvis att kalla hypoteser. Hypoteserna i vårt exempel är alltså

Tillverkningskostnaden är i medeltal 200 kr

Tillverkningskostnaden är i medeltal 220 kr

Vi låter μ beteckna populationens medeltillverkningskostnad. Den första hy-

hypotesen kallas vanligen för nollhypotes (här betyder "noll" att inga förändringar från kalkylens antagande föreligger). Uttryckt i populationsparametern μ skriver vi nollhypotesen som

$$H_0 : \mu = \mu_0 = 200.$$

Den andra hypotesen kallas alternativhypotes. Analogt skriver vi alternativhypotesen som

$$H_A : \mu = \mu_A = 220.$$

Symbolerna H_0 och H_A representerar nollhypotesen respektive alternativhypotesen. μ_0 är värdet på parametern μ då nollhypotesen H_0 är sann och μ_A är värdet på μ då alternativhypotesen H_A är sann. Hypoteserna formuleras alltid så att H_0 och H_A står i motsatsställning. Då den ena är sann så är den andra falsk och vice versa.

Vi måste påpeka här att de två värdena på μ (200 och 220) betraktar vi som de enda möjliga värdena, ett av dem måste vara sant. I praktiken kan givetvis andra värden, som 215 och 230, vara möjliga. För att förenkla diskussionen vill vi till att börja med begränsa oss till endast två möjliga värden. Längre fram i diskussionen skall vi generalisera metoderna så att även fler än två värden kan betraktas.

Det beslut som skall fattas, det vill säga att välja mellan hypoteserna H_0 och H_A , sker så att man dels får stöd från observationsmaterialet och dels handlar i enlighet med den analys av sannolikheterna för felaktiga beslut man har gjort upp.

Observationerna på tillverkningskostnaden måste uttryckas så att de kan ge meningsfullt stöd åt en av hypoteserna. Företagsledningen vill antagligen göra om kalkylen om medelvärdet av observationerna, dvs medeltillverkningskostnaden för de observerade enheterna, är klart större än 200 kr. På samma sätt får företagsledningen ett visst stöd för antagandet att medeltillverkningskostnaden för enheterna är 200 kr om medelvärdet av observationerna ligger i närheten av 200 kr. Beslutsfattaren skall alltså acceptera H_0 om $\bar{x}$ är i närheten av 200 och förkasta H_0 om $\bar{x}$ är klart större än 200. Den variabel som vi använder för att avgöra ett test, i det här fallet $\bar{X}$, kallas testvariabel. Frågan är nu när vi kan betrakta variabelns värde som "nära"

ett visst hypotetiskt värde och när den är ”klart större”, dvs var gränsen går mellan ”nära 200” och ”klart större än 200”.

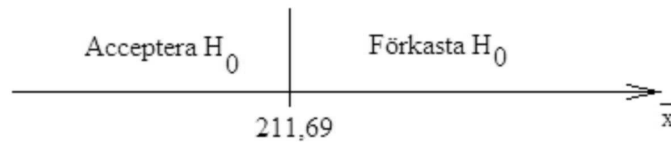
Att bestämma en sådan gräns kallas för att formulera en beslutsregel. Beslutsregeln i det här fallet är sålunda att förkasta H_0 om det observerade urvalsmedelvärdet $\bar{x}$ är större än ett visst tal, k , och acceptera H_0 om $\bar{x}$ är mindre än k . Talet k kallas kritiskt värde. När det kritiska värdet är bestämt så har vi en beslutsregel. Antag att det kritiska värdet i vårt exempel har bestämts till $k = 211,69$ kr. Det betyder att vi har beslutsregeln:

BESLUTSREGEL

Acceptera H_0 (dvs behåll kalkylen) om $\bar{X} \leq 211,69$

Förkasta H_0 (dvs revidera kalkylen) om $\bar{X} > 211,69$

Beslutsregeln identifierar två områden för möjliga värden på $\bar{X}$, det vill säga två områden i utfallsrummet för urvalsmedelvärdet. Dessa områden kallas acceptansområde och förkastelseområde och kan illustreras som i figur 17.1.



Figur 17.1 Utfallsrummet för $\bar{X}$ uppdelad i acceptansområde och förkastelseområde

Det observerade värdet på urvalsmedelvärdet $\bar{X}$ kommer att hamna antingen i acceptansområdet eller i förkastelseområdet och H_0 accepteras respektive förkastas beroende på i vilket område $\bar{X}$ hamnar. Då H_0 förkastas sägs resultatet vara statistiskt signifikant. Här har ordet ”signifikant” en speciell tolkning: alla testresultat är givetvis betydelsefulla, men det är bara då H_0 förkastas som testet säges vara statistiskt signifikant.

Tabell 17.1 Utfallen från ett hypotestest

	H_0 sann; $\mu = 200$	H_0 falsk; $\mu = 220$
Acceptera H_0 ; $\bar{X} \leq 211,69$	Korrekt beslut	Fel beslut; Typ II fel
Förkasta H_0 ; $\bar{X} > 211,69$	Fel beslut; Typ I fel	Korrekt beslut

Strukturen i beslutssituationen visas i tabell 17.1. Det finns två möjliga beslut: att behålla kalkylen eller att revidera den. Populationen har också två möjliga tillstånd: antingen är tillverkningskostnaden i medeltal 200 kr eller så är den 220 kr. För varje kombination av beslut och tillstånd hos populationen är det fattade beslutet antingen korrekt eller felaktigt. De två situationer då ett korrekt beslut fattas inträffar då (1) H_0 accepteras och H_0 är sann, dvs man beslutar sig för att inte revidera kalkylen och den är faktiskt baserad på korrekta antaganden, och (2) H_0 förkastas och H_0 är falsk, det vill säga man beslutar sig för att revidera kalkylen och den är faktiskt baserad på felaktiga premisser. På grund av slumpvariationen i observationerna kan emellertid två situationer uppkomma då ett felaktigt beslut fattas. Beroende på aktuell situation klassificeras ett felaktigt beslut som antingen ett typ I fel eller ett typ II fel. Typ I fel innebär att man förkastar H_0 fastän H_0 är sann (dvs att revidera kalkylen trots att den tidigare kalkylen egentligen bygger på riktiga antaganden). Typ II fel innebär att man accepterar H_0 trots att H_0 är falsk (dvs vi behåller kalkylen trots att den bygger på felaktiga antaganden).

För att utvärdera beslutsregeln beräknar man sannolikheterna att begå de två typerna av fel. Vi betecknar dessa sannolikheter med α respektive β där

$$\alpha = P(\text{Typ I fel}) = P(\text{förkasta } H_0 \mid H_0 \text{ är sann})$$

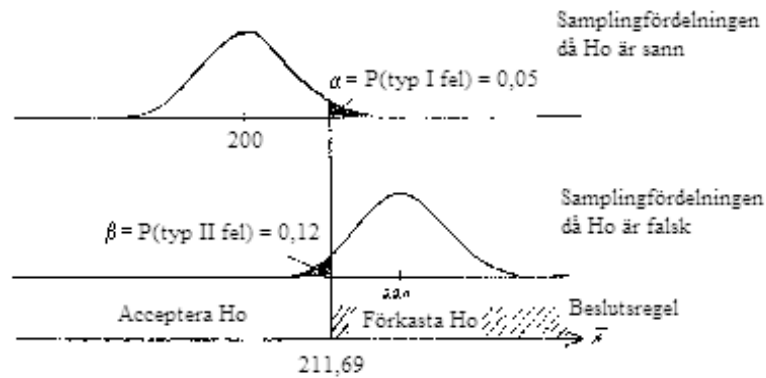
$$\beta = P(\text{Typ II fel}) = P(\text{acceptera } H_0 \mid H_0 \text{ är falsk})$$

Sannolikheten α kallas testets signifikansnivå.

Hur bra är valet 211,69 som kritiskt värde? Den frågan belyses genom att studera sannolikheten att begå ett typ I fel respektive ett typ II fel. Sannolikheten för ett typ I fel bestäms genom att studera samplingfördelningen för teststatistikan under förutsättning att H_0 är sann medan sannolikheten för typ II fel bestäms under förutsättningen att H_A är sann. Vi skall nu illustrera hur denna bestämning går till för vårt exempel med tillverkningskostnad. Antag därför att tillverkningskostnaden för komponenterna är

normalfördelad med variansen $\sigma^2 = 441$ och att urvalsstorleken är $n = 9$ komponenter.

Betrakta först fallet då H_0 är sann. Under denna förutsättning gäller alltså att varje observation dras från en population som är normalfördelad med väntevärdet $\mu = 200$ och variansen $\sigma^2 = 441$. Samplingfördelningen för urvalsmedelvärdet $\bar{X}$ baserat på $n = 9$ observationer är då normal med väntevärdet $\mu = 200$ och standardavvikelsen $\sigma = \sqrt{441/9} = \sqrt{49} = 7$. Den övre kurvan i figur 17.2 representerar samplingfördelningen för $\bar{X}$ under förutsättning att H_0 är sann.



Figur 17.2 Samplingfördelningar för $\bar{X}$ under H_0 och under H_A

Sannolikheten att begå ett typ I fel (representerat med den streckade ytan till höger om $k = 211,69$) är

$$\begin{aligned}\alpha &= P(\bar{X} > 211,69 \mid H_0 \text{ är sann}) = P\left(\frac{\bar{X} - 200}{7} > \frac{211,69 - 200}{7}\right) \\ &= P(Z > 1,67) = 1 - \Phi(1,67) \approx 0,05\end{aligned}$$

Under förutsättning att H_0 är falsk (och alltså H_A är sann) har $\bar{X}$ en annan samplingfördelning. Denna fördelning visas också i figur 17.2. Anledningen till att $\bar{X}$ har en annan samplingfördelning under H_A är att vi nu antar

att observationerna är dragna från en normalfördelning med väntevärdet $\mu = 220$ och variansen $\sigma^2 = 441$. Detta innebär att $\bar{X}$ är normalfördelad med väntevärdet 220 och standardavvikelsen $\sqrt{441/9} = \sqrt{49} = 7$. Sannolikheten att begå ett typ II fel (representerad med den streckade ytan till vänster om 211,69 i samplingfördelningen för $\bar{X}$ under H_A) är därför

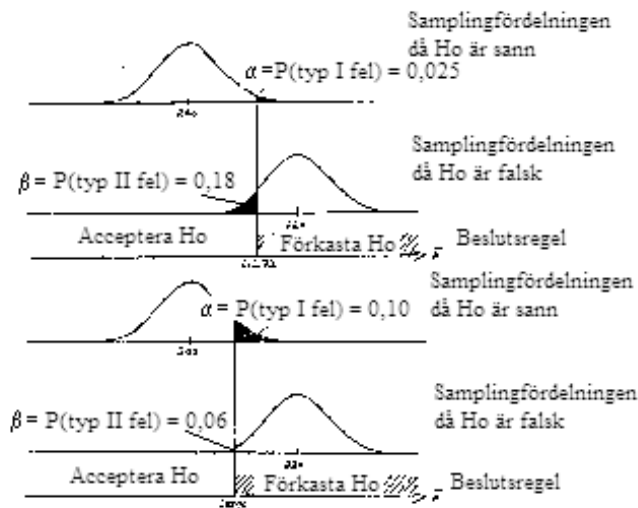
$$\begin{aligned}\beta &= P(\bar{X} \leq 211,69 \mid H_A \text{ är sann}) = P\left(\frac{\bar{X} - 220}{7} \leq \frac{211,69 - 220}{7}\right) \\ &= P(Z \leq -1,19) = \Phi(-1,19) \approx 0,12\end{aligned}$$

Testets styrka definieras som $1 - \beta$. I vårt exempel är således testets styrka $1 - 0,12 = 0,88$.

17.2 Att formulera hypoteser och välja beslutsregel

Beslutsregeln talar om för oss hur vi skall agera då olika utfall från urvalet inträffar. Detta sker genom att jämföra värdet på teststatistikan med det kritiska värdet. Som vi också har sett kan, på grund av slumpen, teststatistikan bli ovanligt stor eller ovanligt liten på ett sådant sätt att ett felaktigt beslut kommer att fattas. Sannolikheterna för att begå fel beror dels på hur beslutsregeln är utformad, dvs på valet av kritiskt värde, och dels på samplingfördelningen för teststatistikan. Speciellt vet vi att samplingfördelningen för urvalsmedelvärdet får en mindre varians då urvalsstorleken ökar. Därför påverkas felsannolikheterna även av urvalsstorleken. Detta medför att urvalsstorleken måste anpassas så att felsannolikheterna överensstämmer med de risker beslutsfattaren är villig att ta. Vi studerar nu hur det kritiska värdet och urvalsstorleken påverkar felsannolikheterna.

I exemplet med tillverkningskostnaderna hade vi det kritiska värdet $k = 211,69$. Betrakta nu två andra val av kritiska värden, $k = 208,96$ och $k = 213,72$. Hur kommer sannolikheterna α och β att förändras?



Figur 17.3 Effekten på felsannolikheterna då det kritiska värdet förändras

De två situationerna är illustrerade i figur 17.3. Jämför dessa situationer med figur 17.2 där fallet med $k = 211,69$ är illustrerad. Då vi ökar det kritiska värdet från 211,69 till 213,73 minskar sannolikheten att förkasta H_0 då H_0 är sann. Den nya sannolikheten att begå ett typ I fel är nu

$$\begin{aligned}\alpha &= P(\bar{X} > 213,73 \mid H_0 \text{ är sann}) = P\left(\frac{\bar{X} - 200}{7} > \frac{213,73 - 200}{7}\right) \\ &= P(Z > 1,96) = 1 - \Phi(1,96) \approx 0,025\end{aligned}$$

att jämföra med det tidigare $\alpha = 0.05$. Den nya beslutsregeln med $k = 213,72$ minskar alltså sannolikheten att begå ett typ I fel. Priset vi får betala för det är en ökad sannolikhet att acceptera H_0 då H_0 är falsk. Genom att öka det kritiska värdet ökar således sannolikheten att begå ett typ II fel, i vårt exempel från $\beta = 0,12$ till

$$\begin{aligned}\beta &= P(\bar{X} \leq 213,72 \mid H_A \text{ är sann}) = P\left(\frac{\bar{X} - 220}{7} \leq \frac{213,72 - 220}{7}\right) \\ &= P(Z \leq -0,90) = \Phi(-0,90) \approx 0,18\end{aligned}$$

Motsatt effekt fås om vi i stället minskar det kritiska värdet från 211,69 till 208,96. I detta fall ökar sannolikheten för typ I fel till $\alpha = 0,10$ medan sannolikheten för typ II fel minskar till $\beta = 0,06$.

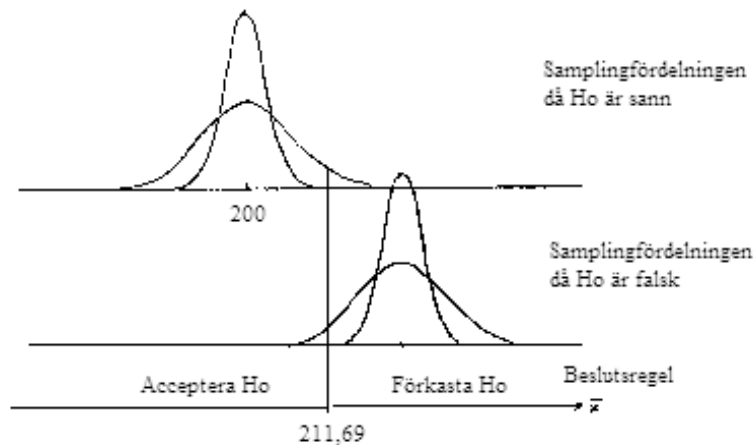
Detta exempel visar att för en viss urvalsstorlek kan man bara minska den ena feltypen på bekostnad av sannolikheten för den andra feltypen. Beslutsregeln måste utformas så att vi får en acceptabel balans mellan riskerna att begå de två typerna av fel.

Enda sättet att minska sannolikheterna för båda felen är att öka urvalsstorleken. En ökad urvalsstorlek kommer nämligen att minska variansen i samplingfördelningen för $\bar{X}$ såväl under H_0 som under H_A . Fördelningarna blir därmed koncentrerade kring sina väntevärden. Nettoeffekten blir mindre sannolikheter för felaktiga beslut. Om vi t ex ökar urvalsstorleken i vårt exempel från $n = 9$ till $n = 12$ blir standardavvikelsen i samplingfördelningen för $\bar{X}$ lika med $\sqrt{441/12} \approx 6,06$. Därmed får vi att

$$\begin{aligned}\alpha &= P(\bar{X} > 211,69 \mid H_0 \text{ är sann}) = P\left(\frac{\bar{X} - 200}{6,06} > \frac{211,69 - 200}{6,06}\right) \\ &= P(Z > 1,93) = 1 - \Phi(1,93) \approx 0,03\end{aligned}$$

$$\begin{aligned}\beta &= P(\bar{X} \leq 211,69 \mid H_A \text{ är sann}) = P\left(\frac{\bar{X} - 220}{6,06} \leq \frac{211,69 - 220}{6,06}\right) \\ &= P(Z \leq -1,37) = \Phi(-1,37) \approx 0,09\end{aligned}$$

I figur 17.4 finns dessa felsannolikheter illustrerade. Den slutsats vi kan dra är att om urvalsstorleken ökas så minskar sannolikheterna för både typ I fel och typ II fel. Eftersom varje observation är förknippad med en viss kostnad finns det en praktisk gräns för hur stora urvalen kan vara. Vanligen är kostnaden den begränsande faktorn för hur stora urval man tar.



Figur 17.4 Effekten på samplingfördelningen då urvalsstorleken ökar

Nollhypotesen och alternativhypotesen är varandras motsatser, så att då den ena är sann, är den andra falsk. Detta medger att vi bara behöver betrakta två typer av beslut, att acceptera H_0 eller att förkasta H_0 . Ett problem här är att avgöra vilken av hypoteserna vi skall kalla för H_0 . Grundprincipen för hur H_0 formuleras följer ur diskussionen i avsnitten 4.3 och 4.4 och logikens modus tollens regel. Det vill säga, om vi har förutsättningarna

(i) "Om hypotesen H_0 är sann så inträffar händelsen A "

(ii) "Händelsen A inträffar ej"

så är den logiska slutsatsen att

"Hypotesen H_0 är falsk"

Det intressanta här är att om vi ändrar förutsättningen (ii) till "Händelsen A inträffar" kan vi inte dra någon korrekt slutsats om hypotesen H_0 . Det är därför ofta mer intressant att kunna förkasta en nollhypotes än att inte förkasta den. Därav kommer också benämningen "statistiskt signifikant" i det fall då observationerna talar emot H_0 och vår beslutsregel säger oss att förkasta H_0 bör formuleras så att vår teori stärks om H_0 förkastas.

En tumregel för formulering av nollhypotesen är att den skall svara mot

ett antagande om ingen förändring, därav benämningen "noll". Om vi t ex prövar en ny behandling av en viss sjukdom, svarar H_0 mot ett antagande att den nya behandlingen inte har någon effekt. Om vi kan förkasta H_0 kan vi alltså dra slutsatsen att behandlingen har en effekt.

I exemplet med tillverkningskostnader formulerades hypoteserna som

$$H_0 : \mu = \mu_0 = 200$$

$$H_A : \mu = \mu_A = 220$$

Hypoteser som anger ett enda värde på en parameter brukar kallas enkla hypoteser. I exemplet med tillverkningskostnader är därför både H_0 och H_A enkla hypoteser. Man skulle emellertid kunna tänka sig att företagsledningen inte bara vill ha $\mu = 220$ kr som alternativ, utan även att $\mu = 201, 202, \dots, 220$, eller något annat större värde är tänkbara alternativ. Om företagsledningen är intresserad av alla medeltillverkningskostnader större än 200 kr vore hypoteserna

$$H_0 : \mu = \mu_0 = 200$$

$$H_A : \mu > \mu_0.$$

mer realistiska. Hypoteser som H_A kallas i detta fall för sammansatta hypoteser, eftersom de innehåller flera möjliga alternativa värden för μ . Ibland har man även anledning att studera sammansatta nollhypoteser. T ex kanske det finns anledning att studera $H_0 : \mu \leq \mu_0$. Sådana situationer är emellertid mer komplicerade att studera och vi lämnar dem därhän i denna framställning.

I vissa situationer är alternativhypotesen tvåsidig. Som nollhypotes har vi då i vårt exempel att medeltillverkningskostnaden μ är 200 kr, dvs $H_0 : \mu = \mu_0$. Alternativhypotesen är att medeltillverkningskostnaden är antingen lägre än 200 kr, dvs att $\mu < \mu_0$ eller att den är större än 200 kr, dvs att $\mu > \mu_0$. Alternativhypotesen är då den tvåsidiga sammansatta hypotesen $H_A : \mu \neq \mu_0$.

17.3 Prövning av hypoteser om populationsmedelvärdet i en normalfördelad population

Antag att vi har dragit ett urval om n individer från en normalfördelad population. Antag också att populationsvariansen σ^2 är känd, så att urvalsmedelvärdets varians, $\sigma_{\bar{X}}^2$ är känd. Det problem vi vill studera är hur hypoteser om populationsmedelvärdet μ prövas då vi på förhand har valt en signifikansnivå α , för testet. Vi visar först hur det kritiska värdet k bestäms vid en ensidigt sammansatt alternativhypotes:

$$H_0 : \mu = \mu_0$$

$$H_A : \mu > \mu_0$$

där μ_0 är ett godtyckligt tal. Vårt problem är alltså att bestämma ett tal k sådant att

$$P(\bar{X} > k \mid \mu = \mu_0) = \alpha.$$

Men,

$$\begin{aligned} P(\bar{X} > k \mid \mu = \mu_0) &= P\left(\frac{\bar{X} - \mu_0}{\sigma_{\bar{X}}} > \frac{k - \mu_0}{\sigma_{\bar{X}}}\right) = P\left(Z > \frac{k - \mu_0}{\sigma_{\bar{X}}}\right) \\ &= 1 - \Phi\left(\frac{k - \mu_0}{\sigma_{\bar{X}}}\right). \end{aligned}$$

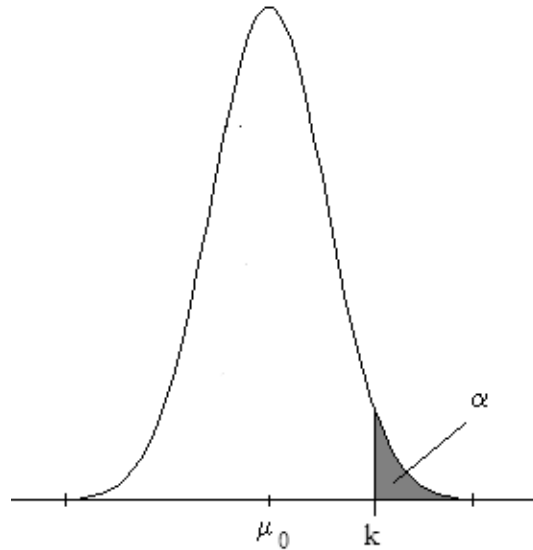
För t ex $\alpha = 0.05$ utnyttjar vi kunskapen att $\Phi(1,64) \approx 0,95$ så att

$$\frac{k - \mu_0}{\sigma_{\bar{X}}} \approx 1.64$$

eller

$$k \approx \mu_0 + 1,64 \cdot \sigma_{\bar{X}}.$$

Detta illustreras i figur 17.5.



Figur 17.5 Bestämning av kritiskt värde i ett ensidigt test

I det tvåsidiga testet av hypoteserna

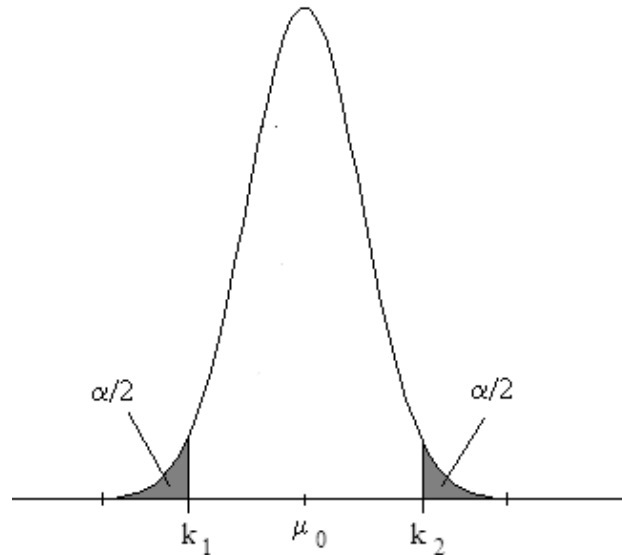
$$H_0 : \mu = \mu_0$$

$$H_A : \mu \neq \mu_0$$

förkastas H_0 både om $\bar{X}$ är "onormalt stor" och om $\bar{X}$ är "onormalt liten". Vi behöver alltså två kritiska värden, såsom illustreras i figur 17.6. Bestämningen av de kritiska värdena k_1 och k_2 sker på i princip samma sätt som när man bestämmer det kritiska värdet i det ensidiga fallet, men vi får tänka på att dela upp α i två hälfter, en för vardera svansen i samplingfördelningen för $\bar{X}$ under H_0 . De kritiska värdena skall alltså väljas så att

$$P(\bar{X} < k_1 \mid H_0 \text{ är sann}) = \alpha/2$$

$$P(\bar{X} > k_2 \mid H_0 \text{ är sann}) = \alpha/2$$



Figur 17.6 Bestämning av kritiska värden i tvåsidiga test

För t ex $\alpha = 0,05$ använder vi kunskapen att $\Phi(-1,96) \approx 0,025$ och $\Phi(1,96) \approx 0,975$ för att erhålla

$$k_1 = \mu_0 - 1,96 \cdot \sigma_{\bar{X}}$$

$$k_2 = \mu_0 + 1,96 \cdot \sigma_{\bar{X}}$$

När vi bestämmer kritiska värden använder vi oss av en standardisering av $\bar{X}$ så att vi kan använda $N(0,1)$ -fördelningen. I det tvåsidiga testet med signifikansnivån 0,05 bestämmer vi k_2 så att sannolikheten att $Z = (\bar{X} - \mu_0) / \sigma_{\bar{X}}$ är större än eller lika med $(k_2 - \mu_0) / \sigma_{\bar{X}}$ är 0,025. Vi får detta genom att sätta $(k_2 - \mu_0) / \sigma_{\bar{X}} = 1,96$. För att se efter om $\bar{X} > k_2$ kan vi alltså lika gärna se efter om $Z > 1,96$. Genom att använda den standardiserade variabeln Z får vi alltså beslutsregeln

Förkasta H_0 om $z < 1,96$ eller om $z > 1,96$

där $z = (\bar{x} - \mu_0) / \sigma_{\bar{X}}$ är det observerade värdet på den standardiserade testvariabeln. Detta är en betydligt enklare beslutsregel och används regelmässigt i datorprogram. I stället för att räkna ut kritiska värden i x -skalan använder vi oss alltså av standardiserade kritiska värden i z -skalan.

Antag nu att populationsvariansen σ^2 är okänd. Variansen för $\bar{X}$ kan då inte beräknas, utan måste ersättas med en uppskattning i våra formler. En lämplig uppskattning är S^2 , urvalsvariansen. Vi vet då att

$$T = \frac{\bar{X} - \mu_0}{S/\sqrt{n}}$$

är t -fördelad med $n - 1$ frihetsgrader. Vi visar nu hur man använder den kunskapen för att bestämma beslutsregel då vi prövar hypotesen

$$H_0 : \mu = \mu_0$$

$$H_A : \mu \neq \mu_0$$

Låt $t_{\alpha/2}(n - 1)$ vara $(\alpha/2) \cdot 100$ percentilen i t -fördelningen med $n - 1$ frihetsgrader, dvs det tal som är sådant att sannolikheten är $\alpha/2$ att en t -fördelad stokastisk variabel men $n - 1$ frihetsgrader är mindre än talet $t_{\alpha/2}(n - 1)$. På samma sätt som tidigare får vi då att

$$\begin{aligned} P(\bar{X} < k_1 \mid H_0 \text{ är sann}) &= P\left(\frac{\bar{X} - \mu_0}{S/\sqrt{n}} < \frac{k_1 - \mu_0}{S/\sqrt{n}}\right) \\ &= P\left(T < \frac{k_1 - \mu_0}{S/\sqrt{n}}\right) = \frac{\alpha}{2}. \end{aligned}$$

Därav följer att vi förkastar H_0 om det observerade värdet på T är tillräckligt litet,

$$t = t_{\alpha/2}(n - 1)$$

där

$$t = \frac{\bar{x} - \mu_0}{s/\sqrt{n}}$$

Låt $t_{1-\alpha/2}(n-1)$ vara $(1 - \alpha/2) \cdot 100$ percentilen i t -fördelningen med $(n-1)$ frihetsgrader, dvs det tal som är sådant att sannolikheten är $1 - \alpha/2$ att en t -fördelad stokastisk variabel med $(n-1)$ frihetsgrader är mindre än talet $t_{1-\alpha/2}(n-1)$. På motsvarande sätt som ovan får vi då att vi förkastar H_0 om det observerade värdet på T är tillräckligt stort,

$$t > t_{1-\alpha/2}(n-1)$$

Eftersom t -fördelningen är symmetrisk kring noll är $t_{\alpha/2}(n-1) = -t_{1-\alpha/2}(n-1)$.

På signifikansnivån α anser vi alltså att rimliga värden på t är värden i intervallet

$$t \pm t_{1-\alpha/2}(n-1)$$

Om t inte är mindre än $-t_{1-\alpha/2}(n-1)$ eller större än $t_{1-\alpha/2}(n-1)$ kan vi sålunda inte förkasta nollhypotesen, medan om $|t|$ är större än $t_{1-\alpha/2}(n-1)$ så förkastar vi nollhypotesen på signifikansnivån α .

Exempel 17.1

Antag att vi vill pröva hypotesen att herr A's restid till arbetet är 30 minuter. Standardavvikelsen i herr A's restid är okänd. Ett urval om 20 observationer dras. I urvalet uppskattas standardavvikelsen till 2. Vi skall nu pröva

$$H_0 : \mu = \mu_0 = 30$$

$$H_A : \mu \neq \mu_0$$

på signifikansnivån $\alpha = 0,05$. Eftersom

$$t_{1-0,05/2}(20-1) = t_{0,975}(19) = 2,093$$

följer det att rimliga värden på $t = \frac{\bar{x}-\mu_0}{s/\sqrt{n}} = \frac{\bar{x}-30}{2/\sqrt{20}}$ om hypotesen H_0 är sann, är alla värden i intervallet $-2,093 < t < 2,093$.

Om det observerade medelvärdet är $\bar{x} = 26,5$ får vi $t = \frac{26,5-30}{2/\sqrt{20}} = -7,83$ vilket alltså inte ligger i acceptansområdet, varför hypotesen om att herr A's restid är 30 minuter måste förkastas på signifikansnivån 5%. ■

17.4 Prövning av hypoteser om populationsmedelvärdet i en icke-normalfördelad population

Vi har hittills antagit att urvalet har skett från en normalfördelad population. Därmed har vi kunnat beräkna kritiska värden för testen genom att vi vet att

$$Z = \frac{\bar{X} - \mu}{\sigma/\sqrt{n}}$$

är normalfördelad med väntevärdet noll och variansen 1 i det fall då σ är känd, och

$$t = \frac{\bar{X} - \mu}{S/\sqrt{n}}$$

är t -fördelad med $n-1$ frihetsgrader i det fall då σ är okänd men uppskattas med S . Om urvalet inte sker från en population med okänd fördelning vet vi inte vilken samplingfördelning testvariabeln har. Däremot vet vi från centrala gränsvärdessatsen att fördelningen för Z blir mer och mer lik den standardiserade normalfördelningen då urvalsstorleken n ökar. Detta innebär att om vi har ett tillräckligt stort urval gör det inte så mycket att vi inte känner fördelningen för den population vi drar urvalet från. Vad som är ett "tillräckligt stort urval" beror bl a på hur bra precisionen måste vara i den enskilda tillämpningen och hur pass lik en normalfördelning populationens

verkliga fördelning är. Vanligen brukar man anse att $n = 30$ är ett tillräckligt stort urval.

Exempel 17.2

En cigaretttillverkare hävdar i sin reklam att de cigaretter de tillverkar innehåller mindre än 25 mg nikotin i genomsnitt. Ett urval om 36 cigaretter gav medelvärdet 24,5 mg och standardavvikelsen 3 mg. Vi prövar nu cigaretttillverkarens påstående genom att pröva hypotesen

$$H_0 : \mu = 25$$

$$H_A : \mu < 25$$

dvs vi prövar en enkelsidig hypotes. Vi väljer signifikansnivån $\alpha = 0,05$. Eftersom antalet observationer är stort använder vi oss av att

$$Z = \frac{\bar{X} - 25}{S/\sqrt{n}}$$

är approximativt normalfördelad och förkastar H_0 om det observerade z -värdet är mindre än -1,64. Vi får att

$$z = \frac{24,5 - 25}{3/\sqrt{36}} \approx -1,25$$

dvs vi kan inte förkasta H_0 . Slutsatsen blir därmed att det inte finns något empiriskt stöd för cigaretttillverkarens påstående. Det kan mycket väl vara så att det genomsnittliga nikotininnehållet är 25 mg och vi kan inte utesluta att den minskning till 24,5 mg vi ser i detta urval beror endast på slumpmässiga avvikelser. ■

17.5 Prövning av hypoteser om skillnader mellan populationsmedelvärden

Vi studerar nu hur man kan pröva hypoteser om skillnader mellan populationsmedelvärden i två speciella fall. Båda fallen kan oftast identifieras med problemet att studera om populationen har förändrat sig eller ej då den har varit utsatt för någon behandling eller annan påverkan, dvs vi vill se om behandlingen har haft någon effekt. De populationer vi då jämför är snarare samma population före och efter behandlingen. Låt μ_F vara populationsmedelvärdet före behandlingen och μ_E vara populationsmedelvärdet efter behandlingen. Den hypotes vi vill pröva är

$$H_0 : \mu_F = \mu_E$$

mot någon alternativhypotes.

I det första fallet vi studerar observerar vi samma individer både före och efter behandlingen. Låt t ex x_i vara observationen på den i -te individen före behandlingen och y_i vara observationen på samma individ efter behandlingen. Vi kan då bilda förändringen $d_i = x_i - y_i$ för den i -te individen. På samma sätt kan vi beräkna förändringar för alla individer, $i = 1, 2, \dots, n$. Om hypotesen är H_0 sann kommer $d_1, d_2, \dots, d_n$ att vara observationer på en stokastisk variabel, D , som har väntevärdet $\mu_D = 0$. Hypotesen H_0 kan därför omformuleras till

$$H_0 : \mu_D = 0$$

Att pröva en sådan hypotes sammanfaller nu med den metodik vi har studerat i de tidigare avsnitten i detta kapitel. Om t ex observationerna både före och efter behandlingen är normalfördelade blir differensen också normalfördelad. Berodande på om dess varians är känd eller ej baserar vi testet på z respektive t . Om observationerna inte är normalfördelade måste vi förlita oss på stora urvalsstorlekar och använda z .

Exempel 17.3

Vårt problem är att se om ett nytt viktnedskningsprogram har någon effekt. Antag att vi vet att vikten hos personerna i populationen är normalfördelad med väntevärdet μ_F före behandlingen enligt programmet och med medelvärdet μ_E efter behandlingen, med skillnaden $\mu_D = \mu_F - \mu_E$. Vi skall pröva hypotesen att programmet inte har någon effekt på vikten, dvs

$$H_0 : \mu_D = 0$$

mot

$$H_A : \mu_D > 0$$

Om hypotesen H_0 är sann och om observationerna är oberoende följer det att individernas viktsförändringar är oberoende observationer på en normalfördelad stokastisk variabel med väntevärde noll och varians σ_D^2 . Vanligen är σ_D^2 okänd och måste ersättas med en uppskattning. Antag att vi gör parvisa observationer på 9 individer och att följande observationer har erhållits

Individ	Vikt före	Vikt efter	Differens
1	84	71	13
2	72	67	5
3	81	81	0
4	80	71	9
5	85	78	7
6	71	64	7
7	81	72	9
8	75	75	0
9	68	68	0

I medeltal har alltså viktnedskningen varit 5,56 kg med standardavvikelsen 4,69. Om vi använder signifikansnivån $\alpha = 0,05$ skall vi alltså förkasta H_0 om det observerade z -värdet överstiger 1,64. I vårt exempel får vi att

$$z = \frac{5,56}{4,69/\sqrt{9}} \approx 3,56,$$

dvs vi kan dra slutsatsen att behandlingen har en statistiskt signifikant effekt på 5 % nivån. ■

I det andra fallet av prövning av hypoteser om skillnader mellan population-smedelvärden har vi inte parvisa jämförelser, utan vi har n stycken observationer från den ena populationen, säg $X_1, X_2, \dots, X_n$, med medelvärdet $\bar{X}$ och m stycken observationer från den andra, säg $Y_1, Y_2, \dots, Y_m$, med medelvärdet $\bar{Y}$. Vi skriver populationsmedelvärdena mer allmänt nu som μ_1 respektive μ_2 och formulerar den hypotes vi vill pröva som

$$H_0 : \mu_1 - \mu_2 = 0$$

Om observationerna är oberoende och normalfördelade med de kända varianserna σ_1^2 i den första populationen och σ_2^2 i den andra, så är skillnaden $\bar{X} - \bar{Y}$ normalfördelad med väntevärdet noll och variansen

$$\sigma_0^2 = \frac{\sigma_1^2}{n} + \frac{\sigma_2^2}{m}$$

Detta innebär att

$$Z = \frac{\bar{X} - \bar{Y}}{\sigma_0}$$

är normalfördelad med väntevärdet noll och variansen ett, dvs vi kan använda Z för att pröva hypotesen om lika populationsmedelvärden.

Om observationerna är oberoende och normalfördelade men med okända varianser måste vi förutsätta att varianserna är lika, säg σ^2 i båda populationerna. Vi baserar då en uppskattning av σ^2 på observationerna från båda populationerna genom

$$S_P^2 = \frac{(n-1) S_1^2 + (m-1) S_2^2}{n+m-2}$$

där S_1^2 och S_2^2 är urvalsvarianserna för observationerna från population 1 respektive 2. En uppskattning av variansen för $\bar{X} - \bar{Y}$ är därför $S_0^2 = S_P^2 \left(\frac{1}{n} + \frac{1}{m} \right)$, vilket medför att

$$t = \frac{\bar{X} - \bar{Y}}{S_0}$$

är t -fördelad med $n + m - 2$ frihetsgrader. Vi använder då fördelningen för denna testvariabel för att pröva hypotesen.

Om observationerna är oberoende men med okänd fördelning får vi lita till centrala gränsvärdessatsen, dvs att

$$Z = \frac{\bar{X} - \bar{Y}}{S_0}$$

är approximativt normalfördelad för stora urval, där S_0^2 ges av uttrycket

$$S_0^2 = \frac{S_1^2}{n} + \frac{S_2^2}{m}$$

Exempel 17.4

Vid en industri finns en maskin A som producerar vissa enheter. Man funderar över att köpa en ny maskin, B, som visserligen har en högre produktionshastighet när den fungerar som den ska, men den kan också ha ett större antal driftstopp. Man beslutar att köpa maskin B om dess genomsnittsproduktion är högre än maskin A:s genomsnittsproduktion. Båda maskinerna provas under 100 dygn och man erhåller

	Maskin A	Maskin B
Antal observationer	100	100
Medelproduktion	1305	1338
Standardavvikelse	25	100

Hypoteserna formuleras som

$$H_0 : \mu_A - \mu_B = 0$$

$$H_A : \mu_A - \mu_B < 0$$

Om maskin B genomsnittligt sett är bättre än A blir differensen $\mu_A - \mu_B$ negativ, vilket anges i mothypotesen, som på grund av frågeställningen är ensidig.

Det finns ingen anledning att anta att de båda variablerna är normalfördelade, men urvalen är stora. Vi beräknar därför en uppskattning av variansen för skillnaden till

$$s_0^2 = \frac{25^2}{100} + \frac{100^2}{100} = 106,25.$$

Då vi utför prövningen på signifikansnivån $\alpha = 0,05$ skall vi sålunda förkasta H_0 om z -värdet är mindre än 1,64. Med våra data får vi

$$z = \frac{1305 - 1338}{\sqrt{106,25}} \approx 3,20.$$

I detta exempel får vi alltså att $z \approx 3,20$ dvs H_0 förkastas och vi anser att maskin B är bättre. ■

17.6 Prövning av hypoteser om proportioner

Att pröva en hypotes om proportioner bygger på att binomialfördelningen kan approximeras med normalfördelningen. Antag t ex att andelen individer med en viss egenskap i en stor population är π . Låt X vara antalet individer med den egenskapen i ett urval om n individer. Då är X en binomialfördelad stokastisk variabel med parametrarna n och π . Vidare vet vi att $E[X] = n\pi$ och $V[X] = n\pi(1 - \pi)$. Låt proportionen av individer i urvalet med den studerade egenskapen vara P . Då är $P = X/n$ och det följer ur räknereglerna för väntevärden och varianser att

$$E[P] = E\left[\frac{X}{n}\right] = \frac{E[X]}{n} = \frac{n\pi}{n} = \pi$$

$$V[P] = V\left[\frac{X}{n}\right] = \frac{V[X]}{n^2} = \frac{n\pi(1-\pi)}{n^2} = \frac{\pi(1-\pi)}{n}$$

Om urvalet är stort kan fördelningen för

$$Z = \frac{P - \pi}{\sqrt{\pi(1-\pi)/n}}$$

approximeras med normalfördelningen med väntevärdet noll och variansen ett.

Exempel 17.5

För en stor population av kvinnor i en viss ålder gäller vid en viss tidpunkt att 64 % är körkortsinnehavare, vilket man fastställt genom en totalundersökning. Vid ett senare tillfälle vill man undersöka om andelen körkortsinnehavare i den aktuella åldern är oförändrad. Man drar därför ett urval om 200 individer, varav 160 visar sig vara körkortsinnehavare. Vi vill nu pröva på signifikansnivån 5 % om andelen körkortsinnehavare i populationen har ändrats.

Hypoteserna blir

$$H_0 : \pi = 0,64$$

$$H_A : \pi \neq 0,64$$

Nollhypotesen anger här att andelen körkortsinnehavare är oförändrad medan alternativhypotesen säger att en förändring har ägt rum. Alternativhypotesen är således tvåsidig. En förändring kan ju antingen bestå i att andelen har minskat eller att den har ökat. Båda svansarna i fördelningen för P blir därför aktuella som kritiskt område. Använder vi Z som testvariabel skall vi alltså förkasta H_0 om det observerade z -värdet är mindre än -1,96 eller om det är större än 1,96. Vi ser att den observerade proportionen är $p = 160/200 = 0,80$ och variansen för P om H_0 är sann är

$$\frac{\pi(1-\pi)}{n} = \frac{0,64 \cdot (1-0,64)}{200} = 0,001152$$

varför det observerade z -värdet blir

$$z = \frac{0,80 - 0,64}{\sqrt{0,001152}} \approx 4,71.$$

Eftersom $4,71 > 1,96$ förkastar vi, på signifikansnivån 5 %, hypotesen om ingen förändring. ■

17.7 p-värden

När vi prövar en hypotes försöker vi avgöra om det är rimligt eller orimligt att få de data vi har fått under förutsättning att nollhypotesen är sann. Den metod för att pröva en nollhypotes som vi har studerat bygger därför på att sätta upp ett förkastelseområde som är konstruerat så att det är liten sannolikhet att teststatistikan får ett värde i förkastelseområdet om nollhypotesen är sann. Som ett alternativ kan man beräkna sannolikheten att få ett datamaterial som är lika sannolikt eller mer orimligt än de data vi har fått under förutsättning att nollhypotesen är sann. Denna sannolikhet kallas p -värde.

Antag att vi har en modell som säger att våra observationer är normalfördelade med väntevärdet μ och med känd varians σ^2 och vi vill pröva hypotesen

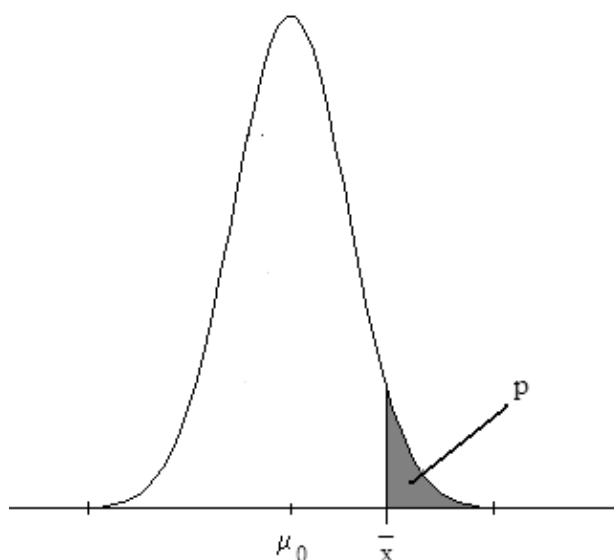
$$H_0 : \mu = \mu_0$$

$$H_A : \mu > \mu_0$$

Vi förkastar då H_0 om det observerade urvalsmedelvärdet $\bar{x}$ är tillräckligt stort (större än något kritiskt värde) eller ekvivalent, om $z = (\bar{x} - \mu_0) / \sigma / \sqrt{n}$ är tillräckligt stort. p -värdet är sannolikheten att den stokastiska variabeln $\bar{X}$ är större än det observerade urvalsmedelvärdet $\bar{x}$,

$$\begin{aligned} p &= P(\bar{X} > \bar{x}) = P\left(\frac{\bar{X} - \mu_0}{\sigma / \sqrt{n}} > \frac{\bar{x} - \mu_0}{\sigma / \sqrt{n}}\right) \\ &= 1 - \Phi\left(\frac{\bar{x} - \mu_0}{\sigma / \sqrt{n}}\right) \end{aligned}$$

Denna sannolikhet illustreras i figur 17.7.



Figur 17.7 p -värdet är sannolikheten att den stokastiska variabeln $\bar{X}$ får ett värde som är större än det observerade värdet $\bar{x}$

p -värdet är alltså sannolikheten att få ett värde på teststatistikan som är minst så ovanligt som det observerade värdet, under förutsättning att nollhypotesen är sann. p -värdet kan således inte tolkas som sannolikheten att nollhypotesen är sann. p -värdet är ett uttalande om data (och potentiella data) och ingenting annat. Ett p -värde som är 0,02 betyder alltså att sannolikheten att observera det värde vi fått på teststatistikan eller ett mer ovanligt värde är 0,02.

Då p -värdet är stort har vi alltså inte observerat något som är särskilt orimligt. Händelser med stor sannolikhet inträffar ju relativt ofta. Ett stort p -värde innebär alltså att data är konsistenta med nollhypotesen och det finns ingen anledning att förkasta nollhypotesen. Här skall vi dock komma ihåg att det finns många andra hypoteser som kan generera de data vi fått,

så att ett högt p -värde bevisar inte att nollhypotesen är sann. Det mesta vi kan säga då vi har ett högt p -värde är att nollhypotesen inte verkar vara falsk. Formellt säger vi att vi inte kan förkasta nollhypotesen.

Då p -värdet är litet innebär det att vi har observerat något som är mycket ovanligt om nollhypotesen är sann. Vi måste då avgöra om det är så att nollhypotesen är sann och vi har observerat något mycket ovanligt eller om det är så att nollhypotesen är falsk och det var fel att utgå ifrån den då vi beräknade p -värdet.

Det är i allmänhet inget större problem att beräkna ett p -värde och de flesta datorprogram som utför hypotestest räknar automatiskt ut p -värden. En viktig skillnad mellan testförfarandet med kritiskt område och att använda p -värden är att när vi använder ett kritiskt område bestämmer vi själva signifikansnivå och anger därmed vad som är signifikant avvikande från nollhypotesen. När vi rapporterar ett p -värde överlåter vi åt läsaren att avgöra om p -värdet är tillräckligt litet för att det svarar mot att data avviker signifikant från nollhypotesen.

Exempel 17.1 (fortsättning)

För att förenkla beräkningarna antar vi att standardavvikelsen σ är känd och lika med 2. Medelvärde av restiden från 20 observerade resor är $\bar{x} = 26,5$. Observera att detta är ett tvåsidigt test varför p -värdet blir sannolikheten att $\bar{X}$ avviker mer än $|26,5 - 30| = 3,5$ både uppåt och nedåt från det hypotetiska värdet $\mu_0 = 30$.

$$\begin{aligned} p &= P(\bar{X} \leq 26,5 \text{ eller } \bar{X} \geq 33,5) = 2P(\bar{X} \leq 26,5) \\ &= 2P\left(Z \leq \frac{26,5 - 30}{2/\sqrt{20}}\right) \approx 2\Phi(-7,83) \approx 0 \quad \blacksquare \end{aligned}$$

Övningar

17.1 Vid kontrollvägning av 12 st 2 kg förpackningar av mjöl var medelvikten i urvalet 1,95 kg. Man har vid längre kontrollserier kunnat konstatera en konstant standardavvikelse om $\sigma = 0,10$ kg. Man kan vidare betrakta vikterna som oberoende normalfördelade stokastiska variabler.

a) Är det erhållna resultatet förenligt med hypotesen att populationen av alla producerade förpackningar håller en genomsnittsvikt om 2 kg? Signifikansnivån väljs till 5 %.

b) Diskutera huruvida förutsättningarna kan anses realistiska.

17.2 Från flera tidigare undersökningar vet man att proportionen defekta enheter i en viss produktionsprocess legat nära 8 %. Efter att vissa ändringar i produktionsprocessen vidtagits vill man undersöka hur stor proportion defekta enheter processen nu ger. Bland enheter som tillverkats efter ändringarna utväljs 160 enheter slumpmässigt. 10 av dessa är defekta. Är det tänkbart att de vidtagna åtgärderna inte haft någon som helst effekt, dvs att processens genomsnittliga felfrekvens ligger kvar på 8 %? Risken att felaktigt påstå detta får vara högst 5 %.

17.3 Betrakta övning 16.5. Testa hypotesen att livslängden hos producerade transistorer är i medeltal 30 timmar.

17.4 Vid ett företag vill man pröva om "medellivslängden" hos ett skärverktyg är 3000 skär eller om det är mindre än 3000 skär. Om sex observationer på livslängden var 2970, 3020, 3005, 2900, 2940 och 2925, vilken slutsats bör man dra då signifikansnivån 0,05 används? Livslängden antas vara normalfördelad.

17.5 Chefen för bageriföretaget Rokak AB har som en av sina målsättningar att minst 40 % av kunderna på marknaden har kännedom om produktnamnet "Rokaks kex".

a) Chefen har en känsla av att den nuvarande marknadsföringen är dålig. Om han kan få belägg för att mindre än 40 % känner till "Rokaks kex" kommer han att vidta en "drastisk omorganisation" av marknadsavdelningen. En marknadsundersökning genomförs för att undersöka marknadens kännedom om "Rokaks kex". Det visar sig då att av 500 slumpmässigt utvalda individer på marknaden hade 180 kännedom om "Rokaks kex". Använd ett statistiskt test för att ge chefen ett beslutsunderlag. Redovisa och motivera förutsättningar och antaganden för det test du använder.

b) Strax efter undersökningen startas en marknadsföringskampanj för produkten. Därefter genomförs en ny marknadsundersökning. Antag att i den nya undersökningen har 195 av 500 slumpmässigt utvalda individer kännedom om produktnamnet. Använd ett statistiskt hypotestest för att testa om kampanjen har haft effekt och ökat andelen kunder som känner till produkten.

17.6 En tillverkare av glödlampor hävdar att hans produkter fungerar i medeltal 800 timmar. Ett slumpmässigt urval om 30 glödlampor togs ut för att pröva tillverkarens påstående. Följande data erhöles:

510	600	800	810	690	640	610	780	600	420	740	600	750	780
900	910	780	620	420	420	750	720	680	390	650	540	640	450
600	840												

Testa tillverkarens hypotes vid 99 % signifikansnivå.

17.7 Vid tillverkning av bildrör till TV-apparater vet man att dessa har en genomsnittlig livslängd på 1200 timmar med en standardavvikelse på 300 timmar. För att effektivisera tillverkningen infördes en ny produktionsteknik. Ett urval av bildrör från den nya produktionsmetoden visar sig ha en genomsnittlig livslängd på 1265 timmar. Ger den nya produktionsprocessen bildrör med längre livslängd än den gamla?

17.8 I en kommun med 65 000 röstberättigade fick X-partiet 20,0 % av rösterna vid 2006 års val. Partiet tror att resultatet av den kampanj man fört efter valet blir att partiet kommer att öka sin andel sympatisörer. För att ta reda på hur det förhåller sig beslutar partistyrelsen att låta göra en partisympatiundersökning. Man bestämmer att undersökningen skall baseras på 400 personer, slumpmässigt valda från röstlängden. Resultatet blev att av 400 personer "röstade" 92 med X-partiet. Har partiet en signifikant högre andel sympatisörer vid undersökningstillfället jämfört med valresultatet 2006 eller ligger skillnaden inom den statistiska felmarginalen? Signifikansnivån väljs till 0,05.

17.9 Tio idrottsutövare som normalt inte använder tobak deltog i ett experiment för att se om snus påtagligt ökar hjärtverksamheten. I experimentet mättes idrottarnas vilopulser. De fick sedan lägga in var sin portionspåse snus. Efter 20 minuter mättes pulsarna igen. Följande mätvärden erhöles:

Försöksperson	Puls före	Puls efter
1	81	105
2	81	91
3	68	87
4	61	86
5	67	82
6	74	78
7	75	87
8	64	94
9	70	93
10	60	90

Föreslå ett lämpligt test för att pröva om snusen ger en ökning av pulsen. Ange förutsättningarna för det test du använder. Använd signifikansnivån 5 %.

17.10 En psykolog gjorde följande experiment. Vid en väg mättes (med dolda instrument) bilarnas hastighet på två ställen med några kilometers mellanrum. Efter lite flyttande av mätinstrumenten bestämdes två punkter där bilarna höll ungefär samma hastighet. Därpå parkerades en polisbil på en parkeringsficka mellan de båda mätpunkterna. Hastigheten mättes sedan hos 12 bilar enligt nedanstående tabell.

Bil nr	Hastighet före poilsbilen	Hastighet efter polisbilen
1	83	61
2	88	84
3	106	95
4	131	121
5	84	83
6	87	79
7	129	92
8	97	88
9	92	69
10	91	90
11	124	115
12	99	100

Psykologens hypotes var att blotta åsynen av polisbilen skulle få bilisterna att sakta farten. Föreslå ett lämpligt statistiskt test. Ange förutsättningarna för testet. Genomför testet på signifikansnivån 0,05.

17.11 Bromsförmågan jämfördes för två olika biltyper av 2009 års modell. Från båda biltyperna valdes 64 bilar slumpmässigt ut. Vid testet uppmättes den bromssträcka (i meter) som fordrades för att stanna bilen när bromsningen satten in vid hastigheten 65 km/tim. Efter testet gjordes följande beräkningar:

	Biltyp 1	Biltyp 2
Urvalsstorlek	64	64
Medelvärde	29,50	27,25
Varians	6,37	5,44

Ger det redovisade bromstestet tillräckligt belägg för påståendet att det råder en skillnad mellan biltypernas genomsnittliga bromssträcka? Ställ upp lämpliga hypoteser för frågeställningen ovan. Gör nödvändiga antaganden och genomför ett statistiskt test på signifikansnivån 5 %. Vilken slutsats följer av resultaten?