

F1:A Ekonometri, introduktion

Kap 1-3 Repetition av Sannolikhetssteori och Statistisk inferens

Allmänt tänker vi oss ett experiment med ett antal möjliga utfall.

Mängden av alla möjliga utfall betecknas S och benämns utfallsrummet.

Låt A vara en delmängd av utfallsrummet, $A \subseteq S$.

A kallas en händelse (event)

Experimentet har utförts och A inträffade.

Ex 1) Experiment = Kast med tärning och vi noterar antalet prickar upptill

Utfallsrum:	$S = \{1, 2, 3, 4, 5, 6\}$	
Händelser:	$A = \{1, 2, 3\}$	"mindre än 4"
	$A' = \{4, 5, 6\}$	"större än 3"
	$B = \{1, 3, 5\}$	"udda tal"
	$B' = \{2, 4, 6\}$	"jämnt tal"
	$C = \{2\}$	"tvåa"
	$D = \{1, 2, 3, 4, 5, 6\}$	"något hände"

Ex 2) Experiment = Kast med tärning tills vi noterar en sexa

Utfallsrum:	$S = \{1, 2, 3, \dots\}$	
Händelser:	$A = \{1\}$	"sexa i 1a kastet"
	$B = \{1, 2, 3\}$	"högst 3 kast"
	$C = \{4, 5, 6, 7, \dots\}$	"minst 4 kast"

- En variabel vars värde bestäms av utfallet av ett slumpförsök (experiment) kallas *slumpvariabel* el. *stokastisk variabel*.

- Slumpvariabler betecknas oftast med stor bokstav X, Y, Z osv. och de värden som variabeln antar betecknas oftast med små bokstäver x, y, z osv.

- En slumpvariabel är antingen
 - *diskret*, dvs. den kan anta ett uppräknligt antal möjliga utfall.
 - *kontinuerlig*, dvs den kan anta värden från ett kontinuerligt (sammanhängande) intervall

- Man betecknar sannolikheten för en händelse $X = x$ med

$$P(X = x)$$

Oberoende händelser

Om sannolikheten för en händelse A inte påverkas av om en annan händelse har inträffat (eller inte) så säger vi att de är oberoende:

$$P(A | B) = P(A) \iff A \text{ och } B \text{ oberoende}$$

där symbolen $|$ betecknar betingning, dvs "givet att B har inträffat".

Om A och B är oberoende gäller att

$$P(A \cap B) = P(A)P(B)$$

om de är beroende gäller att

$$P(A | B) = \frac{P(A \cap B)}{P(B)}$$

Diskreta slumpvariabler

Låt X vara en diskret slumpvariabel som kan anta värden $x_1, x_2, \dots, x_n, \dots$. Då låter vi

$$P(X = x) = \begin{cases} f(x) & \text{för alla } x_1 \text{ i utfallsrummet} \\ 0 & \text{annars} \end{cases}$$

där $f(x)$ är en funktion av det tänkta utfallet utfallet som ger oss sannolikheten för att just detta ska inträffa.

Funktionen kallas för *sannolikhetsfunktionen* för X , eller *probability mass function* (pmf).

Obs! För att $f(x)$ ska vara en pmf måste följande gälla:

$$1 : 0 \leq f(x) \leq 1$$

$$2 : \sum_{x \in S} f(x) = 1$$

Ex 1) Beteckna händelsen "det regnar idag" med $X = 1$ och motsatsen med $X = 0$. Låt p beteckna sannolikheten för den första händelsen då kan vi skriva

$$f(x) = p^x (1-p)^{1-x}$$

Ex 2) Låt X vara en slumpvariabel med utfallsrum $1, 2, 3, \dots$ med pmf

$$f(x) = (1-p)^{x-1} p$$

där $0 < p < 1$ är en *parameter* som bestämmer hur funktionen ser ut. Ett experiment som vi kan tänka oss med denna beskrivning är "kasta en tärning tills vi observerar en sexa". Parametern p betecknar då sannolikheten för sexa, $p = 1/6$.

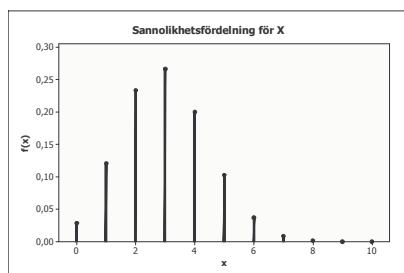
Sannolikheten att vi får en sexa i fjärde kastet beräknas

$$\begin{aligned} P(X = 4) &= f(4) = \left(1 - \frac{1}{6}\right)^{4-1} \cdot \frac{1}{6} \\ &= \frac{125}{1296} \approx 0.09645 \end{aligned}$$

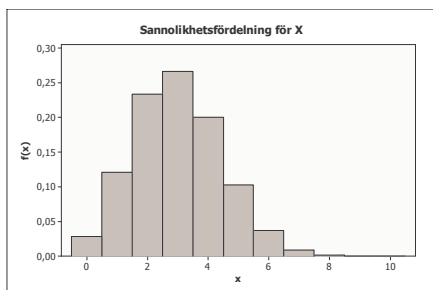
Sannolikheten att vi får vänta *högst* fyra kast beräknas som

$$\begin{aligned} P(X \leq 4) &= \sum_{k=1}^4 f(k) \\ &= f(1) + f(2) + f(3) + f(4) \\ &= 1 - \left(1 - \frac{1}{6}\right)^4 = \frac{671}{1296} \\ &\approx 0.51775 \end{aligned}$$

Diagram som åskådliggör sannolikhetsfunktionen för en diskret slumpvariabel



alternativt



Kontinuerliga slumpvariabler

Om ett experiment medger ett kontinuerligt utfallsrum får man en kontinuerlig slumpvariabel. Vi kan beräkna sannolikheten för händelser av typen $P(a \leq X \leq b)$ dvs att X hamnade mellan två värden med

$$P(a \leq X \leq b) = \int_a^b f(x) dx$$

Funktionen $f(x)$ benämns *täthetsfunktion* eller *probability density function* (pdf).

Obs! Händelser och sannolikheter av typen $P(X = x)$ är meningslösa för en kontinuerlig variabel då

$$P(X = a) = \int_a^a f(x) dx = F(a) - F(a) = 0$$

Formuleringen $f(x)$ används för att *beskriva* funktionen, det är inte sannolikheter som ramlar ut från den!

Obs! För att $f(x)$ ska vara en riktig pdf måste det gälla att

$$1 : f(x) \geq 0$$

$$2 : \int_{-\infty}^{\infty} f(x) dx = 1$$

Allmänt brukar man ofta ange en *kumulativ fördelningsfunktion* eller *cumulative distribution function* (cdf) enligt

$$P(X \leq x) = F(x) = \int_{-\infty}^x f(t) dt$$

Ex 1) Vi mäter tiden till ett maskinhaveri vid någon anläggning. Pdf'en kan kanske skrivas som

$$f(x) = \frac{1}{\theta} \exp\left(-\frac{x}{\theta}\right), \quad 0 < x < \infty$$

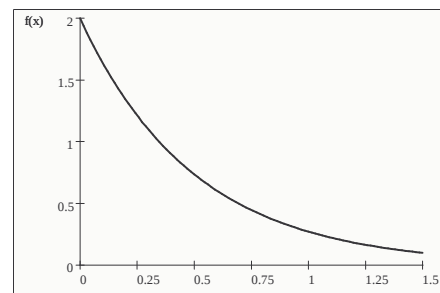
där θ är en parameter som bestämmer hur funktionen ser ut. Sannolikheten för att vi måste vänta högst x tidsenheter blir

$$F(x) = \int_0^x \frac{1}{\theta} \exp\left(-\frac{t}{\theta}\right) dt = 1 - e^{-\frac{x}{\theta}}$$

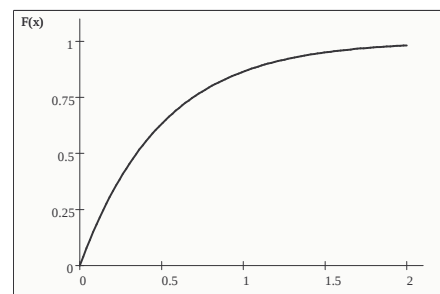
Sannolikheten att vi måste vänta mer än x tidsenheter blir

$$1 - F(x) = \int_x^{\infty} \frac{1}{\theta} \exp\left(-\frac{t}{\theta}\right) dt = e^{-\frac{x}{\theta}}$$

Pdf'en kan åskådliggöras enligt



och cdf'en enligt



Väntevärden / Expected value

För ett givet datamaterial med n observationer kan vi beräkna sådant som *genomsnittet* eller *stickprovsvariansen* enligt

$$\bar{x} = \frac{\sum x_i}{n}, \quad s^2 = \frac{\sum (x_i - \bar{x})^2}{n-1}$$

som mått på läge och spridning. Motsvarigheten för en slumpvariabel benämns *väntevärde* och *varians* och beräknas enligt

Diskret slumpvariabel:

$$E(X) = \sum_{x \in S} x f(x) = \mu$$

$$Var(X) = \sum_{x \in S} (x - \mu)^2 f(x) = \sigma^2$$

Kontinuerlig slumpvariabel:

$$E(X) = \int_{x \in S} x f(x) = \mu$$

$$Var(X) = \int_{x \in S} (x - \mu)^2 f(x) = \sigma^2$$

Allmänt kan man tänka sig väntevärden av valfria funktioner av en slumpvariabel:

$$E(h(X)) = \int_{x \in S} h(x) f(x)$$

Exempelvis

$$E(X^2) = \int_{x \in S} x^2 f(x) \quad \text{el} \quad E(e^X) = \int_{x \in S} e^x f(x)$$

Egenskaper för väntevärdesberäkningar (a och b är konstanter):

$$1 : \quad E(b) = b$$

$$2 : \quad E(aX + b) = aE(X) + b = a\mu + b$$

$$3 : \quad E(aX + bY) = aE(X) + bE(Y)$$

Om X och Y är *oberoende* variabler gäller att

$$4 : \quad E(XY) = E(X)E(Y) = \mu_X \mu_Y$$

För *varianser* gäller (a och b är konstanter):

$$1 : \operatorname{Var}(X) = E[(X - \mu)^2] = \sigma^2$$

$$2 : \operatorname{Var}(X) = E(X^2) - E(X)^2 = E(X^2) - \mu^2$$

$$3 : \operatorname{Var}(a) = 0$$

$$4 : \operatorname{Var}(aX + b) = a^2 \operatorname{Var}(X) = a^2 \sigma^2$$

Om X och Y är *oberoende* variabler gäller att

$$5 : \operatorname{Var}(X + Y) = \operatorname{Var}(X) + \operatorname{Var}(Y) = \sigma_X^2 + \sigma_Y^2$$

$$6 : \operatorname{Var}(X - Y) = \operatorname{Var}(X) + \operatorname{Var}(Y)$$

$$7 : \operatorname{Var}(aX + bY) = a^2 \operatorname{Var}(X) + b^2 \operatorname{Var}(Y)$$

Om två slumpvariabler är beroende gäller

$$1. \operatorname{Var}(X + Y) = \operatorname{Var}(X) + \operatorname{Var}(Y) + 2\operatorname{Cov}(X, Y) \\ = \sigma_X^2 + \sigma_Y^2 + 2\rho\sigma_X\sigma_Y$$

$$2. \operatorname{Var}(X - Y) = \operatorname{Var}(X) + \operatorname{Var}(Y) - 2\operatorname{Cov}(X, Y) \\ = \sigma_X^2 + \sigma_Y^2 - 2\rho\sigma_X\sigma_Y$$

och för en uppsättning av n stycken beroende slumpvariabler gäller att

$$3. \operatorname{Var}\left(\sum X_i\right) = \sum \operatorname{Var}(X_i) + 2 \sum_{i < j} \operatorname{Cov}(X_i, X_j)$$

Kovarians

Mått på hur två slumpvariabler samvarierar linjärt:

$$\begin{aligned} \operatorname{Cov}(X, Y) &= E[(X - \mu_X)(Y - \mu_Y)] \\ &= E(XY) - \mu_X\mu_Y \\ &= \sigma_{XY} \end{aligned}$$

Egenskaper (a, b, c och d konstanter):

$$1 : \operatorname{Cov}(X, Y) = 0 \iff X \text{ och } Y \text{ oberoende}$$

$$2 : \operatorname{Cov}(aX + b, cY + d) = ac \cdot \operatorname{Cov}(X, Y)$$

Korrelationskoefficienten

Mått på det linjära beroendet:

$$\rho = \frac{\operatorname{Cov}(X, Y)}{\sqrt{\operatorname{Var}(X)}\sqrt{\operatorname{Var}(Y)}} = \frac{\sigma_{XY}}{\sigma_X\sigma_Y}$$

Observera att måttet ρ är skaloberoende och att

$$-1 \leq \rho \leq 1$$

Normalfördelningen

Om X är normalfördelad skriver vi

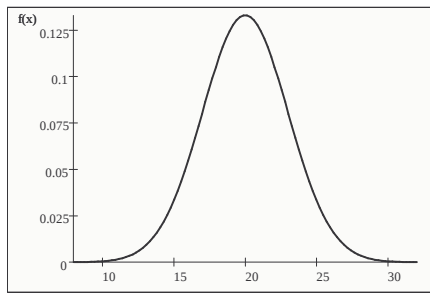
$$X \sim N(\mu, \sigma^2)$$

Förmodligen den viktigaste fördelningen.

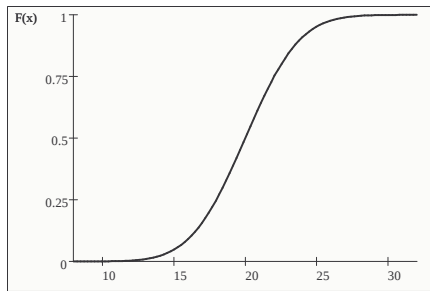
- Symmetrisk runt sitt väntevärde
- Utseendet bestäms av väntevärdet μ och variansen σ^2 .
- Linjärkombinationer av normalfördelade slumpvariabler är normalfördelade
- Centrala gränsvärdessatsen (CGS), linjärkombinationer av slumpvariabler går mot en normalfördelning när antalet variabler ökar, $n \rightarrow \infty$.

Finns diverse statistiska test för att avgöra om data kan sägas komma från en normalfördelning.

Täthetsfunktionen (pdf):



Fördelningsfunktionen (cdf):



Standardisering: Om $X \sim N(\mu, \sigma^2)$ så gäller att

$$Z = \frac{X - \mu}{\sigma} \sim N(\mu, \sigma^2)$$

Beräkning av sannolikheter görs med hjälp av tabeller för en standardiserad normalfördelning.

Ex) Antag $X \sim N(120, 20^2)$. Beräkna $P(X \leq 130)$

$$\begin{aligned} P(X \leq 110) &= P\left(Z \leq \frac{110 - 120}{20}\right) = P(Z \leq -0.5) \\ &= [\text{enl. tabell}] = 0.3085 \end{aligned}$$

χ^2 -fördelningen

- Om V är en χ^2 -fördelad slumpvariabel med r frihetsgrader skriver man

$$V \sim \chi^2(r)$$

- Frihetsgradstalet r är en parameter som bestämmer utseendet för fördelningen.

- Om $X \sim N(\mu, \sigma^2)$ så gäller att

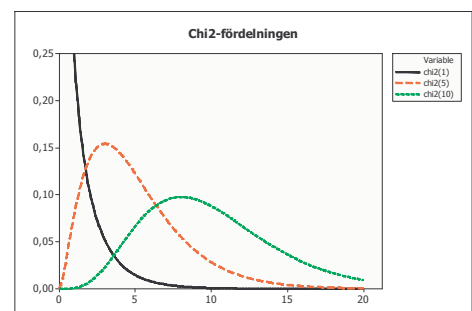
$$\frac{(X - \mu)^2}{\sigma^2} \sim \chi^2(1)$$

- Fördelningen är skev

- Fördelning för stickprovsvarianser

$$\frac{(n-1)S^2}{\sigma^2} \sim \chi^2(n-1)$$

- När antalet frihetsgrader ökar närmar den sig normalfördelningen.



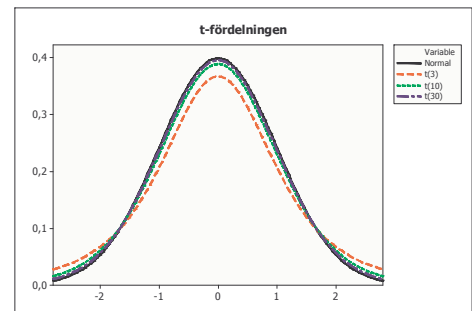
t-fördelningen

- Symmetrisk runt sitt medelvärde
- "Plattare" än en normalfördelning
- Utseendet bestäms av antalet frihetsgrader, r (parameter).
- När antalet frihetsgrader ökar, $r \rightarrow \infty$, närmar den sig en normalfördelning.
- Om $Z \sim N(0, 1)$ och $U \sim \chi^2(r)$ och om de är oberoende gäller att

$$T = \frac{Z}{\sqrt{U/r}} \sim t(r)$$

dvs en t-fördelning med r frihetsgrader. Ur detta kan man visa att

$$\frac{\bar{X} - \mu}{S/\sqrt{n}} \sim t(n-1)$$



F-fördelningen

- Om $U \sim \chi^2(r_1)$ och $V \sim \chi^2(r_2)$ och om de är oberoende gäller att

$$F = \frac{U / r_1}{V / r_2} \sim F(r_1, r_2)$$

dvs kvoten mellan två oberoende χ^2 -fördelade variabler (justerat för antalet frihetsgrader) är F -fördelad.

- Skev fördelning (utdragen till höger)
- När antalet frihetsgrader ökar närmar den sig normalfördelningen
- Koppling mellan t -fördelningen och F -fördelningen

$$[t(r)]^2 = F(1, r)$$

- Det gäller att

$$F(r_1, r_2) = \frac{1}{F(r_2, r_1)}$$

Punktskattningar

- Ofta har man ganska bra uppfattning om vad som gäller i en given situation, tex att observationer är (approximativt) normalfördelade, att man har oberoende observationer etc.
- Man kan ofta komma fram till att experimentet kan beskrivas / approximeras med någon fördelning / modell.
- Det man typiskt *inte* vet är parametervärdena som ingår i modellen.
- Dessa måste *skattas* eller *estimeras* med ledning av ett stickprov.
- En *statistika* är en funktion av stickprovet

$$h(x_1, x_2, \dots, x_n)$$

Exempel på statistikor)

$$\bar{X} = \frac{1}{n} \sum X_i$$

$$S^2 = \frac{1}{n-1} \sum (X_i - \bar{X})^2$$

även regressionkoefficienter i en linjär regressionsmodell

$$\hat{\beta}_1 = \frac{\sum (X_i - \bar{X})(Y_i - \bar{Y})}{\sum (X_i - \bar{X})^2}$$

$$\hat{\beta}_0 = \bar{Y} - \hat{\beta}_1 \bar{X}$$

- En statistika används typiskt för att skatta en parameter, vi kallar då statistikan för en *estimator* eller *skattning* för parametern.

- Dessa estimatorer $\hat{\theta}$ är slumpvariabler

Egenskaper hos estimatorer vid små eller moderata stickprovsstorlekar

- Väntevärdesriktighet
- Minimal varians
- Effektivitet
- Linjäritet
- BLUE
- Minimum mean square errors (MS) estimator

Väntevärdesriktighet

Om väntevärdet för en estimator $\hat{\theta}$ är lika med den parameter θ vi försöker skatta är den väntevärdesriktig (*unbiased*).

Bias = systematiskt fel (error)

$$Bias(\hat{\theta}) = E(\hat{\theta}) - \theta$$

"Träffar man i genomsnitt rätt? Eller måste man justera siktet?"

Minimal varians

Om variansen för en estimator $\hat{\theta}_1$ har mindre varians jämfört med vilken annan möjlig estimator $\hat{\theta}_2$ för samma parameter, säger vi att den har minimal varians

$$Var(\hat{\theta}_1) \leq Var(\hat{\theta}_2)$$

"Hur ser träffbilden ut? Samlat eller utspritt?"

Effektivitet

Låt $\hat{\theta}_1$ och $\hat{\theta}_2$ vara två väntevärdesriktiga estimatorer för θ . Om det gäller att

$$Var(\hat{\theta}_1) < Var(\hat{\theta}_2)$$

så är $\hat{\theta}_1$ en effektiv estimator jämfört med $\hat{\theta}_2$.

"Träffbilden är mer samlad för denna än för denna estimator"

Linjäritet

Om en estimator $\hat{\theta}$ är en linjärkombination av stickprovets observerade värden säger vi att $\hat{\theta}_1$ är en linjär estimator för θ .

$\bar{X}$ en linjärkombination av de observerade värdena

$$\sum \frac{1}{n} X_i = \bar{X}$$

BLUE

Om $\hat{\theta}_1$ är en linjär, väntevärdesriktig estimator för θ med minimal varians i klassen av linjära och väntevärdesriktiga estimatorer så är den en

Best linear unbiased estimator - BLUE estimator.

Minimum MSE estimator

$$MSE(\hat{\theta}) = E(\hat{\theta} - \theta)^2 = Var(\hat{\theta}) + Bias(\hat{\theta})^2$$

Ibland kan man kosta på sig en liten systematisk bias om vinsten i varians är stor, eller tvärtom

Övningsuppgifter

1. Diagnostiskt prov,
2. Uppgifter ur Kleinbaum et al., Kap 3: 1-6, 8, 11-13, 16, 18, 20, 23

F1:B Ekonometri, introduktion

Kap 1-3 Repetition statistisk teori

Egenskaper hos estimatorer vid stora stickprovsstorlekar

- Asymptotisk väntevärdesriktighet
- Konsistens
- Asymptotisk effektivitet
- Asymptotisk normalfördelning

Asymptotisk väntevärdesriktighet

Om bias minskar i takt med stickprovsstorleken och man kan skriva

$$\lim_{n \rightarrow \infty} Bias(\hat{\theta}) = 0$$

så är den asymptotiskt väntevärdesriktig.

Ex)

$$\hat{\sigma}^2 = \frac{\sum (X_i - \bar{X})^2}{n} \quad E(\hat{\sigma}^2) = \sigma^2 \left(\frac{n-1}{n} \right)$$

Konsistens

Om sannolikheten för att komma godtyckligt nära det sanna värdet ökar och går mot ett så är estimatören konsistent

$$\lim_{n \rightarrow \infty} P(|\hat{\theta} - \theta| < \varepsilon) = 1, \quad \forall \varepsilon > 0$$

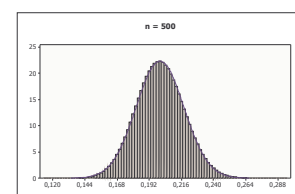
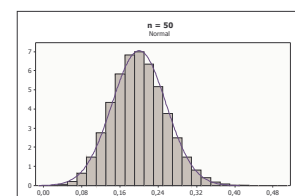
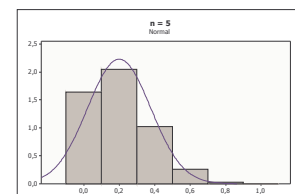
Ett tillräckligt villkor är att både bias och varians för estimatören går mot noll när $n \rightarrow \infty$.

Asymptotisk effektivitet

Om den asymptotiska variansen för en konsistent estimator är mindre än den asymptotiska variansen för alla andra konsistenta estimatorer så är den asymptotiskt effektiv.

Asymptotiskt normalfördelat

En estimator $\hat{\theta}$ sägs vara asymptotiskt normalfördelat om dess fördelning närmar sig en normalfördelning när $n \rightarrow \infty$



Hypotestest

- Börjar med en *nollhypotes* H_0 , ett påstående som ska prövas.
- En alternativ hypotes el. *mothypotes* H_1 eller H_A , ofta men inte alltid den logiska motsatsen till H_0 .
- Mothypotesen är ofta (men inte alltid) det vi vill bevisa.
- Om vi förkastar H_0 accepterar vi mothypotesen och vi säger att resultatet är *signifikant*.
- Enkla hypoteser, endast ett möjligt värde under hypotesen.
- Sammansatta hypoteser, fler än ett möjligt värde under hypotesen.

- Typiskt prövas om en given parameter har ett visst värde
- Vanligt med mer komplexa test.

Ex)

$$H_0 : p = 2/3 \quad \text{mot} \quad H_1 : p = 1/3$$

$$H_0 : \mu \leq 0 \quad \text{mot} \quad H_1 : \mu > 0$$

$$\begin{aligned} H_0 : \mu_1 = \mu_2 &\iff H_0 : \mu_1 - \mu_2 = 0 \\ \text{mot } H_1 : \mu_1 - \mu_2 &\neq 0 \end{aligned}$$

$$\begin{aligned} H_0 : &\text{"modell passar till data"} \\ \text{mot } H_1 : &\text{"modell passar inte till data"} \end{aligned}$$

Testfunktion eller teststatistika

Ofta baserat på just den estimator som man skattat parametern med, tex

$$\bar{X} \quad \text{estimator för } \mu$$

Vi känner till hur denna eller en transformation av den är fördelad under nollhypotesen, tex

$$\frac{\bar{X} - \mu_0}{\sigma/\sqrt{n}} = Z \sim N(0, 1)$$

där μ_0 är värdet för μ som vi antar vara det sanna värdet under H_0 , σ är en känd varians och n är stickprovsstorleken.

Två vägar för klassisk inferens:

1. Signifikantest
2. Konfidsensintervall

Konfidsensintervall

Man konstruerar ett intervall på sådant sätt att sannolikheten att observera ett värde på testfunktionen i det intervallet är lika med ett förutbestämt tal, ofta 95%.

Antag att vi testar

$$H_0 : \mu = 100 \quad \text{mot} \quad H_1 : \mu \neq 100$$

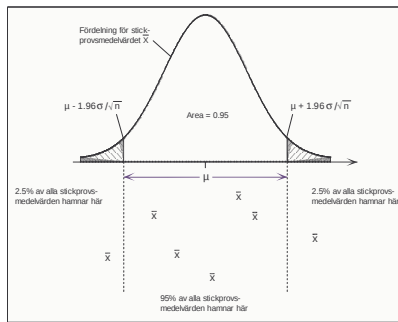
Om vi kan anta att $X_i \sim N(\mu, \sigma^2)$ så kan vi använda att

$$\bar{X} \sim N(\mu, \sigma^2)$$

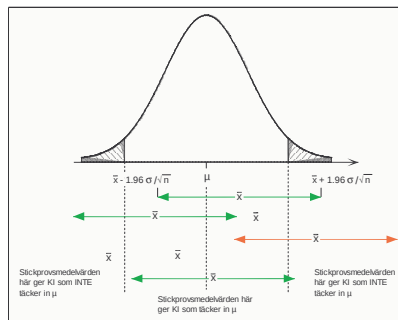
och bilda intervallet

$$\begin{aligned} 1 - \alpha &= \Pr\left(-z_{\alpha/2} \leq \frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \leq z_{\alpha/2}\right) \\ &= \Pr\left(\bar{X} - z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \leq \mu \leq \bar{X} + z_{\alpha/2} \frac{\sigma}{\sqrt{n}}\right) \end{aligned}$$

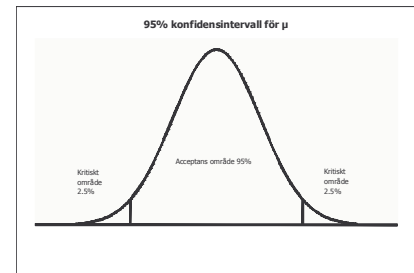
där $z_{\alpha/2}$ är tabellvärden.



Vi bildar sedan utifrån observerat medelvärde $\bar{x}$ ett KI



Kopplingen mellan signifikanstest och konfidensintervall



Intervallet innehåller alla de värden μ_0 som skulle leda till att vi accepterar nollhypotesen $H_0 : \mu = \mu_0$.

Signifikanstest

Hypotesprövning / Signifikanstest

1. Förutsättningar: Ett stickprov av storlek $n = 36$ från en normalfördelning med känd varians $\sigma^2 = 9$.
2. Nollhypotes: $H_0 : \mu = 5$
3. Moithypotes: $H_1 : \mu \neq 5$
4. Signifikansnivå / Felrisk: $\alpha = 5\%$
5. Teststatistika:

$$Z = \frac{\bar{X} - \mu_0}{\sigma/\sqrt{n}} = \frac{\bar{X} - 5}{3/6}$$

6. Fördelning: Under de givna förutsättningarna och under antagandet att nollhypotesen är sann, är teststatistikan $Z \sim N(0, 1)$. Om speciella antaganden behövs, anges dessa, tex "enligt CGS så är Z approximativt $N(0, 1)$ ".

7. Kritiskt värde/område: Definiera under vilka omständigheter vi skulle förkasta H_0 ; tex om vi observerar $|z_{obs}| > 1.960 = z_{0.025}$. Detta är ekvivalent med att vi observerar

$$\bar{x} > \frac{3}{6} \cdot 1.960 + 5 = 5.98$$

eller

$$\bar{x} < -\frac{3}{6} \cdot 1.960 + 5 = 4.02$$

Alternativ 1.

8. Beräkningar: Vi beräknar medelvärdet till 4.01.
9. Slutsats: Vi förkastar H_0 ty $\bar{x} = 4.01 < 4.02$. Det observerade medelvärdet är *signifikant skilt* från 5 på 5%-nivån. Data stöder inte påståendet att $\mu = 5$.

Alternativ 2.

8. Beräkningar: Vi beräknar medelvärdet till $\bar{x} = 4.35$.
9. Slutsats: Vi kan *inte* förkasta H_0 ty $4.02 < \bar{x} < 5.98$. Det observerade medelvärdet är *inte signifikant skilt* från 5 på 5%-nivån. Det finns inte tillräckligt stöd i data för att för att förkasta påståendet att $\mu = 5$.

Feltyper

För ett test finns alltid fyra möjligheter varav två leder till felaktiga slutsatser:

	Förkasta H_0	Förkasta ej H_0
H_0 sann	Typ I fel	ok
H_0 falsk	ok	Typ II fel

Sannolikheten för Typ I fel betecknas α , vilket alltså är felrisken. Sannolikheten För Typ II fel betecknas β .

	Förkasta H_0	Förkasta ej H_0
H_0 sann	α	$1 - \alpha$
H_0 falsk	$1 - \beta$	β

där $(1 - \beta)$ är testets *styrka* (eng. power).

Observera att det är *betingade* sannolikheter, tex

$$\Pr(\text{Förkasta } H_0 \mid H_0 \text{ sann}) = \alpha$$

p-värde eller *p*-value

Ex) Vi tänker oss ett enkelsidigt test $H_0 : \theta = \theta_0$ mot $H_0 : \theta > \theta_0$.

Vi har en testfunktion säg Z . Vi har obeserverat stickprovet och beräknat testfunktionen till ett värde z .

Vad är sannolikheten att observera detta värde eller värre under H_0 ? Dvs vad är

$$\Pr(Z > z \mid H_0 \text{ sann})$$

Detta är det sk *p*-värdet. Tänk så här:

Antag att det framräknade värdet z på testfunktionen Z sammanfaller med den kritiska gränsen C . Vad motsvarar detta för signifikansnivå α ?

Räkna "baklänges" för att ta fram ett α som ger just z som kritisk gräns. Detta är vårt *p*-värde.