

Datorövning 4

Statistikens Grunder 2

Syfte

1. Lära sig utföra goodness of fit-test
2. Lära sig utföra test av homogenitet
3. Lära sig utföra prövning av hypoteser om oberoende

Exempel

Utföra goodness of fit-test

Utgå från exempel 18.1 i kompendiet som handlar om tärningskast. Vi kan läsa in datat på följande sätt:

```
data work.dice;  
input number count;  
datalines;  
1 18  
2 23  
3 16  
4 21  
5 18  
6 24  
;  
run;
```

Vi vill nu testa om tärningen är symmetrisk eller ej. Om tärningen är symmetrisk, och vi kastar den 120 gånger förväntar vi oss att varje nummer ska komma upp 20 gånger. Vi får hypoteserna

$$H_0 : \text{Lika fördelning mellan nummer}$$
$$H_A : \text{Ej lika fördelning}$$

För att göra ett goodness of fit-test skriver vi koden

```
proc freq data=work.dice;  
weight count;  
tables number / testf=(20 20 20 20 20 20) chisq;  
run;
```

Eftersom datat är inläst med en variabel som innehåller frekvenser, måste vi ange "*weight count*". Efter "*tables*" anger vi den variabel som vi vill testa frekvenser för. "*testf*" står för testfrekvens, men det går även bra att skriva

"*testp*" för att istället ange testprocent. "*chisq*" skriver vi för att vi vill att SAS ska räkna ut chi-två värdet.

Vi får utskriften

The FREQ Procedure					
number	Frequency	Test Frequency	Percent	Cumulative Frequency	Cumulative Percent
1	18	20	15.00	18	15.00
2	23	20	19.17	41	34.17
3	16	20	13.33	57	47.50
4	21	20	17.50	78	65.00
5	18	20	15.00	96	80.00
6	24	20	20.00	120	100.00

Chi-Square Test
for Specified Frequencies

Chi-Square	2.5000
DF	5
Pr > ChiSq	0.7765

Sample Size = 120

Det observerade chi-två värdet blev 2.5, vi har 5 frihetsgrader och p-värdet är 0.7765. Bör vi förkasta H_0 ?

Utföra test av homogenitet

Utgå från exempel 18.2 i kompendiet som handlar om ätvanor framför TV:n inom olika länder. Vi kan läsa in datat på följande sätt:

```
data work.ex182 ;
input land$ intensitet$ antal;
datalines;
sverige ofta 512
sverige ibland 125
sverige sällan 212
italien ofta 395
italien ibland 128
italien sällan 404
england ofta 685
england ibland 130
england sällan 32
usa ofta 743
usa ibland 131
usa sällan 71
kina ofta 532
kina ibland 184
kina sällan 181
;
run;
```

Vi vill nu testa om fördelningen av intensiteten är densamma för länderna. Vi får hypoteserna

H_0 : Fördelning av intensitet är samma över länderna

H_A : Fördelning av intensitet är ej samma över länderna

För att testa detta i SAS skriver vi koden

```
proc freq data=work.ex182 order=data;
weight antal;
tables land*intensitet / chisq expected cellchi2 norow nocol;
run;
```

Kommandot "*order=data*" används för att vi vill att SAS ska skriva ut tabellen i samma ordning som datat är inläst. Vi skriver "*chisq*" för att få ut chi-två värdet, det vill säga

$$\sum_{i=1}^n \frac{(n_i - E[n_i])^2}{E[n_i]}.$$

Vi skriver "*expected*" för att vi vill få ut förväntat värde i varje cell, det vill säga $E[n_i]$. Vi skriver "*cellchi2*" för att få chi-två värdet i varje cell, det vill säga

$$\frac{(n_i - E[n_i])^2}{E[n_i]},$$

för varje i . Vi skriver "*norow*" och "*nocol*" för att vi INTE vill ha ut rad- och kolumnprocent, detta för att minimera antalet värden i varje cell.

Vi får dels en korstabell med olika värden i varje cell,

The FREQ Procedure

Table of land by intensitet

land	intensitet			
	ofta	ibland	sällan	Total
sverige	512 545.15 2.0155 11.47	125 132.72 0.4492 2.80	212 171.13 9.7602 4.75	849 19.01
italien	395 595.23 67.356 8.85	128 144.92 1.9744 2.87	404 186.85 252.35 9.05	927 20.76
england	685 543.86 36.626 15.34	130 132.41 0.0438 2.91	32 170.73 112.73 0.72	847 18.97
usa	743 606.79 30.576 16.64	131 147.73 1.8944 2.93	71 190.48 74.946 1.59	945 21.16
kina	532 575.97 3.3565 11.91	184 140.23 13.665 4.12	181 180.81 0.0002 4.05	897 20.09
Total	2867 64.21	698 15.63	900 20.16	4465 100.00

och dels en tabell med chi-två värden och p-värden.

The FREQ Procedure

Statistics for Table of land by intensitet

Statistic	DF	Value	Prob
Chi-Square	8	607.7415	<.0001
Likelihood Ratio Chi-Square	8	628.2416	<.0001
Mantel-Haenszel Chi-Square	1	93.3486	<.0001
Phi Coefficient		0.3689	
Contingency Coefficient		0.3461	
Cramer's V		0.2609	

Sample Size = 4465

Vi är bara intresserade av den första raden utskriften. Chi-två värdet blev

607.7415 och tillhörande p-värde är mindre än 0.0001. Vi bör alltså förkasta H_0 . Övertyga er om att ni kan få fram chi-två värdet genom att summera alla cellers chi-två värden.

Utföra prövning av hypoteser om oberoende

Utgå från exempel 18.3 i kompendiet som handlar om civilstånd och typ av bostad. Vi kan läsa in datat på följande sätt:

```
data work.ex183;  
input bostad$ civilstand$ antal;  
datalines;  
villa gift 25  
villa ogift 5  
villa frånskild 10  
hyreslägenhet gift 15  
hyreslägenhet ogift 25  
hyreslägenhet frånskild 20  
;  
run;
```

Vi vill nu testa om det råder oberoende mellan "*civilstånd*" och "*bostadstyp*", det vill säga

$$\begin{aligned} H_0 &: \text{oberoende mellan "civilstånd" och "bostadstyp"} \\ H_A &: \text{beroende mellan "civilstånd" och "bostadstyp"} \end{aligned}$$

För att testa skriver vi koden

```
proc freq data=work.ex183 order=data;  
weight antal;  
tables bostad*civilstand / chisq expected cellchi2 norow nocol;  
run;
```

Som vi ser använder vi samma kod som för exemplet "*utföra test av homogenitet*". Vi får en liknande utskrift

The FREQ Procedure

Table of bostad by civilstand

bostad	civilstand			
Frequency Expected Cell Chi-Square Percent	gift	ogift	frånskil	Total
villa	25 16 5.0625 25.00	5 12 4.0833 5.00	10 12 0.3333 10.00	40 40.00
hyresläg	15 24 3.375 15.00	25 18 2.7222 25.00	20 18 0.2222 20.00	60 60.00
Total	40 40.00	30 30.00	30 30.00	100 100.00

Statistics for Table of bostad by civilstand

Statistic	DF	Value	Prob
Chi-Square	2	15.7986	0.0004
Likelihood Ratio Chi-Square	2	16.4528	0.0003
Mantel-Haenszel Chi-Square	1	7.2337	0.0072
Phi Coefficient		0.3975	
Contingency Coefficient		0.3694	
Cramer's V		0.3975	

Sample Size = 100

p-värdet är så pass litet att vi kan förkasta H_0 på alla signifikansnivåer som är större än $\alpha = 0.0004$. Det råder alltså beroende mellan variablerna "*civilstånd*" och "*bostadstyp*".

Uppgifter

Basuppgifter

1. Utgå från övning 18.1 i kompendiet. Läs in datat och testa

H_0 : Lika fördelning mellan märke

H_A : Ej lika fördelning

Använd signifikansnivå $\alpha = 0.05$. Bör vi förkasta H_0 ?

2. Utgå från övning 18.8 i kompendiet. Läs in datat och testa

H_0 : Fördelning av "butik" är samma för båda könen

H_A : Fördelning av "butik" är ej samma för båda könen

Använd signifikansnivå $\alpha = 0.05$. Bör vi förkasta H_0 ?

3. Utgå från övning 18.10 i kompendiet. Läs in datat och testa

H_0 : Oberoende mellan "uppfattning om rökning" och "rökvana"

H_A : Beroende mellan "uppfattning om rökning" och "rökvana"

Använd signifikansnivå $\alpha = 0.05$. Bör vi förkasta H_0 ?

Extra uppgifter

1. Utgå från tentamen 2009-10-29, fråga 4. Läs in datat och testa

H_0 : proportionen över inkomstklasser är
 $\{0.18 \ 0.13 \ 0.19 \ 0.20 \ 0.3\}$

H_A : proportionen är ej $\{0.18 \ 0.13 \ 0.19 \ 0.20 \ 0.3\}$
över inkomstklasserna

Använd signifikansnivå $\alpha = 0.05$. Bör vi förkasta H_0 ?

2. Utgå från tentamen 2009-11-23, fråga 4. Läs in datat och testa

H_0 : oberoende mellan "inställning till småbilar"
och "personlighetstyp"

H_A : beroende mellan "inställning till småbilar"
och "personlighetstyp"

Använd signifikansnivå $\alpha = 0.01$. Bör vi förkasta H_0 ?

3. Utgå från övning 18.2a. Läs in datat och testa

H_0 : proportionen är $\{9:3:3:1\}$

H_A : proportionen är inte $\{9:3:3:1\}$

Använd signifikansnivå $\alpha = 0.05$. Bör vi förkasta H_0 ?