

Skriftlig prøve: 17.05.2025

Kursusnavn og -nr.: **Introduktion til Statistik (02323)**

Varighed: 4 timer

Tilladte hjælpemidler: Alle hjælpemidler - uden adgang til internettet

Dette sæt er besvaret af

(studienummer)

(underskrift)

(bord nr.)

Opgavesættet består af 30 spørgsmål af “multiple choice” typen, som er fordelt på 12 opgaver. For at besvare spørgsmålene skal du udfylde “multiple choice” siderne på eksamen.dtu.dk.

Der gives 5 point for et korrekt “multiple choice” svar og -1 point for et forkert svar. KUN følgende 5 svarmuligheder er gyldige: 1, 2, 3, 4 eller 5. Hvis et spørgsmål efterlades blankt eller et ugyldigt svar angives, gives der 0 point for spørgsmålet. Endvidere, hvis mere end et svar angives til det samme spørgsmål, hvilket faktisk er teknisk muligt i online-systemet, gives der 0 point for spørgsmålet. Det antal point, der kræves for at opnå en bestemt karakter eller for at bestå eksamen, afgøres endeligt ved censureringen.

Den endelige besvarelse af opgaverne laves ved at udfylde og aflevere online. Skemaet her er KUN et nød-alternativ til dette. Husk at angive dit studienummer, hvis du afleverer på papir.

Opgave	I.1	I.2	I.3	II.1	II.2	III.1	III.2	IV.1	IV.2	V.1
Spørgsmål	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
Svar										

Opgave	V.2	V.3	V.4	V.5	VI.1	VI.2	VII.1	VII.2	VII.3	VIII.1
Spørgsmål	(11)	(12)	(13)	(14)	(15)	(16)	(17)	(18)	(19)	(20)
Svar										

Opgave	VIII.2	VIII.3	IX.1	X.1	X.2	XI.1	XI.2	XII.1	XII.2	XII.3
Spørgsmål	(21)	(22)	(23)	(24)	(25)	(26)	(27)	(28)	(29)	(30)
Svar										

Eksamenssættet består af 22 sider.

Fortsæt på side 2

Multiple choice opgaver: Der gøres opmærksom på, at der i hvert spørgsmål er én og kun én svarmulighed, som er rigtig. Endvidere er det ikke givet, at alle de anførte alternative svarmuligheder er meningsfulde. Husk altid at afrunde dit eget resultat til antallet af decimaler givet i svarmulighederne før du vælger et svar. Husk også, at der kan forekomme små afvigelser mellem resultatet af bogens formler og tilsvarende indbyggede funktioner i Python.

Opgave I

Denne øvelse omhandler fundamentale statistiske koncepter og basale Python-funktionaliteter.

Spørgsmål I.1 (1)

Hvilken funktion i `scipy.stats`-pakken bruges til at finde teoretiske fraktiler?

- 1 ☐ `cdf`-funktionen
- 2 ☐ `pdf`-funktionen
- 3 ☐ `ppf`-funktionen
- 4 ☐ `rvs`-funktionen
- 5 ☐ `std`-funktionen

Spørgsmål I.2 (2)

Betragt stikprøven (1, 2, 3, 9, 10). Hvilket af følgende udsagn er falsk?

- 1 ☐ Den nedre stikprøvekvartil (Q_1) er 2.
- 2 ☐ Stikprøvemedianen er 3.
- 3 ☐ Stikprøvegennemsnittet er 5.
- 4 ☐ Stikprøvens variationsbredde (range) er 9.
- 5 ☐ Stikprøvevariansen er 16.

Spørgsmål I.3 (3)

Hvilken af følgende størrelser er ikke en stokastisk variabel?

- 1 ☐ Estimatoren $\hat{\beta}$
- 2 ☐ Populationsmiddelværdien μ
- 3 ☐ Stikprøve gennemsnittet $\bar{X}$
- 4 ☐ Stikprøve variansen S^2
- 5 ☐ Teststørrelsen T

Fortsæt på side 4

Opgave II

Lad Y følge en eksponentialfordeling med rateparameter 2, og lad U følge en kontinuert uniform fordeling på intervallet $[3, 6]$. De to stokastiske variable er uafhængige.

Spørgsmål II.1 (4)

Hvad er sandsynligheden for, at Y er større end 2?

1 ☐ $2^{-2} = 0.250$

2 ☐ $2^{-1} = 0.500$

3 ☐ $1 - e^{-4} = 0.982$

4 ☐ $e^{-4} = 0.018$

5 ☐ $2e^{-4} = 0.037$

Spørgsmål II.2 (5)

Hvad er standardafvigelsen for U ?

1 ☐ $\frac{3}{2}$

2 ☐ $\frac{3}{4}$

3 ☐ $\frac{\sqrt{3}}{2}$

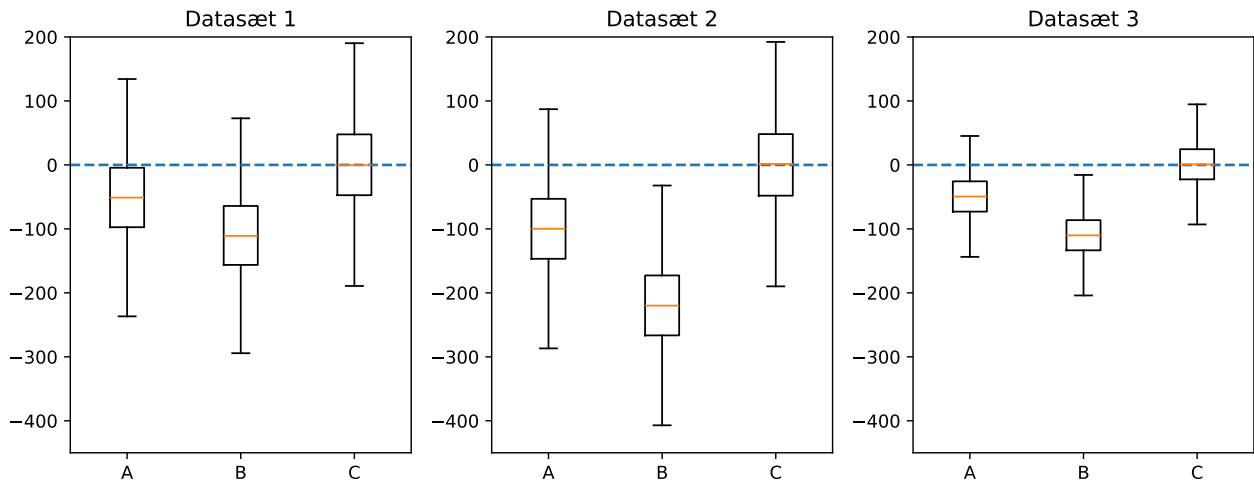
4 ☐ $\frac{\sqrt{3}}{4}$

5 ☐ $\frac{9}{2}$

Fortsæt på side 5

Opgave III

Tre eksperimenter sammenligner tre appetitnedsættende medikamenter (A, B og C). I alle eksperimenterne måles faldet i kalorieindtaget for hvert individ. Resultaterne er gemt i Datasæt 1, Datasæt 2 og Datasæt 3 og er visualiseret med boxplots nedenfor:



Spørgsmål III.1 (6)

Hvilket af følgende udsagn er falsk?

- 1 ☐ De gennemsnitlige effekter af medikamenterne i Datasættene 1 og 3 er sammenlignelige.
- 2 ☐ $SS(TR)$ - *Kvadratafgivelsessummen for behandlinger* - er størst i Datasæt 2.
- 3 ☐ SSE - *Residualkvadratafgivelsessummen* - er størst i Datasæt 3.
- 4 ☐ Inden for hvert datasæt er gruppernes (A, B og C) varianser tilnærmelsesvis ens.
- 5 ☐ I alle tre datasæt kunne medikament C være en placebo (dvs. være uden gennemsnitlig effekt).

Spørgsmål III.2 (7)

Stikprøvestørrelsen er den samme for hver gruppe i alle tre datasæt (dvs. alle ni grupper har den samme størrelse). Der er blevet foretaget en variansanalyse (ANOVA) for hvert datasæt, hvor den observerede F -størrelse og p -værdien for den sædvanlige nulhypotese er blevet udregnet. Hvilket af følgende udsagn er korrekt?

- 1 ☐ Den observerede F -størrelse er den samme for alle tre datasæt.
- 2 ☐ Den foretagne test er en Welch t -test.
- 3 ☐ Alle tre datasæt giver den samme p -værdi.
- 4 ☐ Antagelsen om varianshomogenitet på tværs af grupperne synes brudt i Datasæt 1.
- 5 ☐ Datasæt 1 giver anledning til en højere p -værdi end Datasæt 2 og 3.

Fortsæt på side 7

Opgave IV

En fitnesstræner påstår, at det gennemsnitlige vægttab for en ny kostplan er 4 kg. For at teste påstanden er der blevet udtaget en tilfældig stikprøve med 12 deltagere, der har fulgt kostplanen, hvor deres vægttab (i kg) er blevet målt: 3.5, 4.2, 3.8, 4.4, 4.1, 3.9, 4.0, 4.3, 3.6, 3.7, 4.5, 4.1.

Forskerne udfører derefter en t -test for at bestemme, om det gennemsnitlige vægttab er signifikant forskelligt fra 4 kg ved et signifikansniveau på 5%.

Spørgsmål IV.1 (8)

Lad t_{obs} være den observerede teststørrelse for t -testen. Hvad er afvisningskriteriet?

- 1 ☐ Afvis H_0 , hvis $t_{\text{obs}} \leq t_{1-\alpha}$, hvor $t_{1-\alpha}$ er $(1 - \alpha)$ fraktilen i t -fordelingen med 11 frihedsgrader.
- 2 ☐ Afvis H_0 , hvis $t_{\text{obs}} \leq t_{1-\alpha}$, hvor $t_{1-\alpha}$ er $(1 - \alpha)$ fraktilen i t -fordelingen med 12 frihedsgrader.
- 3 ☐ Afvis H_0 , hvis $t_{\text{obs}} \geq t_{1-\alpha}$, hvor $t_{1-\alpha}$ er $(1 - \alpha)$ fraktilen i t -fordelingen med 11 frihedsgrader.
- 4 ☐ Afvis H_0 , hvis $t_{\text{obs}} \leq -t_{\alpha/2}$ eller hvis $t_{\text{obs}} \geq t_{1-\alpha/2}$, hvor $t_{\alpha/2}$ og $t_{1-\alpha/2}$ er henholdsvis $\alpha/2$ fraktilen og $(1 - \alpha/2)$ fraktilen i t -fordelingen med 11 frihedsgrader.
- 5 ☐ Afvis H_0 , hvis $t_{\text{obs}} \leq -t_{\alpha/2}$ eller hvis $t_{\text{obs}} \geq t_{1-\alpha/2}$, hvor $t_{\alpha/2}$ og $t_{1-\alpha/2}$ er henholdsvis $\alpha/2$ fraktilen og $(1 - \alpha/2)$ fraktilen i t -fordelingen med 10 frihedsgrader.

Spørgsmål IV.2 (9)

Hvad er værdien af den observerede teststørrelse (t_{obs})?

- 1 ☐ $t_{\text{obs}} = 2.83$
- 2 ☐ $t_{\text{obs}} = 1.09$
- 3 ☐ $t_{\text{obs}} = 1.03$
- 4 ☐ $t_{\text{obs}} = 0.32$
- 5 ☐ $t_{\text{obs}} = 0.09$

Fortsæt på side 8

Opgave V

En softwareingeniør ønsker at teste en algoritmes effektivitet. Hun noterer udførelsestiden (**time**) for jobs med varierende kompleksitet (**complex**).

For at analysere sine data, kører hun koden nedenfor, hvor **dat** indeholder de registrerede data.

```
fit1 = smf.ols(formula = 'time ~ complex', data = dat).fit()
print(fit1.summary(slim=True))
```

OLS Regression Results						
=====						
Dep. Variable:	time		R-squared:	0.971		
Model:	OLS		Adj. R-squared:	0.971		
No. Observations:	50		F-statistic:	1631.		
Covariance Type:	nonrobust		Prob (F-statistic):	1.03e-38		
=====						
	coef	std err	t	P> t	[0.025	0.975]

Intercept	-0.1732	0.007	-23.620	0.000	-0.188	-0.158
complex	0.0006	1.47e-05	40.391	0.000	0.001	0.001
=====						

Spørgsmål V.1 (10)

Betragt teststørrelsen for nulhypotesen, at skæringspunktet med y-aksen (interceptet) er nul. Hvilken fordeling bruges til at beregne p -værdien for teststørrelsen?

- 1 ☐ En $t(50)$ -fordeling
- 2 ☐ En $t(48)$ -fordeling
- 3 ☐ En $F(1, 50)$ -fordeling
- 4 ☐ En $\mathcal{N}(0, 1^2)$ -fordeling
- 5 ☐ En $\mathcal{N}(0, 0.007^2)$ -fordeling

Spørgsmål V.2 (11)

Hvilket udsagn om modellens validitet (dvs. modellens antagelser) er korrekt baseret på ovenstående output?

- 1 ☐ Validiteten af modellens antagelser kan ikke vurderes baseret på ovenstående output.
- 2 ☐ Antagelserne må være opfyldt, da R^2 -værdien er tæt på 1.
- 3 ☐ Modellen bør udvides med et yderligere led, da p -værdierne er små.
- 4 ☐ Antagelserne må være opfyldt, da p -værdierne er tæt på nul.
- 5 ☐ Modellen bør udvides med et yderligere led, da R^2 -værdien er tæt på 1.

Tallene nedenfor kan være nyttige i de følgende spørgsmål.

```
print(np.mean(dat['complex']))  
484.58  
print(np.var(dat['complex'], ddof=1))  
12603.187346938777
```

Spørgsmål V.3 (12)

Lad $\hat{\sigma}$ betegne den estimerede standardafvigelse. Hvad er 95% konfidensintervallet for den gennemsnitlige udførelsestid af et job med kompleksitet 300 ifølge modellen?

- 1 ☐ $0.18 \pm 3.32\hat{\sigma}$
- 2 ☐ $0.0068 \pm 2.08\hat{\sigma}$
- 3 ☐ $0.18 \pm 0.82\hat{\sigma}$
- 4 ☐ $0.0068 \pm 1.27\hat{\sigma}$
- 5 ☐ $0.0068 \pm 0.55\hat{\sigma}$

Spørgsmål V.4 (13)

Hvad er estimatet af standardafvigelsen for fejlen ($\hat{\sigma}$)?

- 1 ☐ $2 \cdot 10^{-5}$
- 2 ☐ 0.01
- 3 ☐ 5.5
- 4 ☐ 0.97
- 5 ☐ 0.007

Ingeniøren opdager, at udførelsestiden er proportionel med kompleksiteten i tredje potens, og formulerer følgende model:

$$Y_i = \beta_0 x_i^{\beta_1} \varepsilon_i; \quad \varepsilon_i \sim LN(0, \sigma^2),$$

hvor Y_i og x_i er henholdsvis udførelsestiden og kompleksiteten for job i . Hun tester derfor hypotesen $\mathcal{H}_0 : \beta_1 = 3$ mod en tosidet modhypotese ved at køre koden nedenfor:

```
ltime = np.log(dat['time'])
lcomplex = np.log(dat['complex'])
dat2 = pd.DataFrame({'ltime' : ltime, 'lcomplex' : lcomplex})
fit2 = smf.ols(formula = 'ltime ~ lcomplex', data=dat2).fit()
print(fit2.summary(slim=True))
```

OLS Regression Results						
=====						
Dep. Variable:	ltime		R-squared:	0.999		
Model:	OLS		Adj. R-squared:	0.999		
No. Observations:	50		F-statistic:	8.032e+04		
Covariance Type:	nonrobust		Prob (F-statistic):	4.87e-79		
=====						
	coef	std err	t	P> t	[0.025	0.975]

Intercept	-18.0015	0.055	-325.241	0.000	-18.113	-17.890
lcomplex	2.5459	0.009	283.406	0.000	2.528	2.564
=====						

Spørgsmål V.5 (14)

Hvad er teststørrelsen (t_{obs}) for den sædvanlige test af $\mathcal{H}_0$, og hvad er testens konklusion ved et signifikansniveau på 5%?

- 1 ☐ Hypotesen afvises, da $t_{\text{obs}} = -5.645$.
- 2 ☐ Hypotesen accepteres, da $t_{\text{obs}} = 13.84$.
- 3 ☐ Hypotesen afvises, da $t_{\text{obs}} = -50.46$.
- 4 ☐ Hypotesen afvises, da $t_{\text{obs}} = 283.4$.
- 5 ☐ Hypotesen accepteres, da $t_{\text{obs}} = -5.645$.

Fortsæt på side 12

Opgave VI

De to spørgsmål i denne øvelse skal løses med simulation.

Antallet af (købende) kunder i et supermarked på en tilfældigt udvalgt dag følger en Poissonfordeling med gennemsnit 5000. Beløbet, en tilfældigt udvalgt kunde køber for, er eksponentialfordelt med et gennemsnit på 200 DKK. Det kan antages, at kunderne agerer uafhængigt af hinanden og uafhængigt af antallet af kunder.

Det følgende stykke kode kan bruges til at simulere k daglige omsætninger, men man skal udfylde de manglende dele ved at specificere parametrene i fordelingerne:

```
# Sæt seed
np.random.seed(2025)

# Antal stikprøver
k = 100000

# Antal kunder (array med kundetal for k dage)
N = stats.poisson.rvs(mu=____, size=k)

# Omsætninger (array med omsætninger for k dage)
Y = np.array([np.sum(stats.expon.rvs(scale=____, size=n)) for n in N])
```

Lad Y betegne supermarkedets omsætning på en tilfældigt udvalgt dag.

Spørgsmål VI.1 (15)

Hvad er forventningsværdien og standardafvigelsen for Y ?

- 1 ☐ $\mathbb{E}[Y] = 1,500,000$ DKK og $\text{SD}(Y) = 1,000,000$ DKK
- 2 ☐ $\mathbb{E}[Y] = 1,500,000$ DKK og $\text{SD}(Y) = 20,000$ DKK
- 3 ☐ $\mathbb{E}[Y] = 1,000,000$ DKK og $\text{SD}(Y) = 1,000,000$ DKK
- 4 ☐ $\mathbb{E}[Y] = 1,000,000$ DKK og $\text{SD}(Y) = 20,000$ DKK
- 5 ☐ $\mathbb{E}[Y] = 500,000$ DKK og $\text{SD}(Y) = 1,000,000$ DKK

Spørgsmål VI.2 (16)

Hvad er sandsynligheden for, at supermarkedet har en omsætning på mere end 1,500,000 DKK på en tilfældigt udvalgt dag?

- 1 ☐ Næsten 0%
- 2 ☐ Omkring 11%
- 3 ☐ Omkring 22%
- 4 ☐ Omkring 33%
- 5 ☐ Omkring 50%

Fortsæt på side 14

Opgave VII

Et shippingfirma har hyret en sælger til at rekruttere nye klienter. Sælgeren besøger n uafhængige klienter hver måned, og historiske data viser, at sandsynligheden for at rekruttere en klient efter et besøg er p . Lad X betegne antallet af rekrutterede klienter i en tilfældigt udvalgt måned.

Spørgsmål VII.1 (17)

Hvilken af følgende modeller er mest passende?

- 1 ☐ $X \sim U(0, n)$
- 2 ☐ $X \sim H(n, np, n)$
- 3 ☐ $X \sim \mathcal{N}(np, p^2)$
- 4 ☐ $X \sim \text{Pois}(np)$
- 5 ☐ $X \sim \text{Bin}(n, p)$

Sælgeren tjener 5000 USD per måned og får en bonus på 100 USD for hver klient, der er blevet rekrutteret i løbet af måneden. Sælgerens samlede løn for en tilfældigt udvalgt måned (Y) er derfor givet ved

$$Y = 100X + 5000,$$

hvor X angiver antallet af rekrutterede klienter i den gældende måned. Man kan antage, at standardafvigelsen for X er 2.

Spørgsmål VII.2 (18)

Hvad er variansen for Y ?

- 1 ☐ 5200
- 2 ☐ 20000
- 3 ☐ 25000
- 4 ☐ 40000
- 5 ☐ 45000

Spørgsmål VII.3 (19)

Hvilket af følgende udsagn om korrelationen mellem X og Y er korrekt?

- 1 ☐ Korrelationen mellem X og Y er en: $\rho(X, Y) = 1$.
- 2 ☐ Korrelationen mellem X og Y er mellem en og nul: $0 < \rho(X, Y) < 1$.
- 3 ☐ Korrelationen mellem X og Y er nul: $\rho(X, Y) = 0$.
- 4 ☐ Korrelationen mellem X og Y er negativ: $-1 \leq \rho(X, Y) < 0$.
- 5 ☐ Korrelationen kan ikke bestemmes uden yderligere information.

Fortsæt på side 16

Opgave VIII

En kommune har testet en service, der lader borgere rapportere skader på trafikale infrastruktur (f.eks. veje og trafiklys) via en app. De rapporterede skader blev kategoriseret i tre grupper: skal repareres øjeblikkeligt, kan udskydes og kan ignoreres.

Over en et-årig forsøgsperiode er der blevet rapporteret $n = 165$ skader. Antallet af de forskellige typer skader blev opgjort per årstid og skrevet i nedenstående tabel:

```
# Indlæs data i Python
data = np.array([[44, 10, 32, 18], [20, 8, 4, 2], [11, 2, 10, 4]])
# Konverter til en Pandas dataframe
data = pd.DataFrame(data, index=['Øjeblikkelig', 'Udskydes', 'Ignoreres'],
columns=['Forår', 'Sommer', 'Efterår', 'Vinter'])
print(data)
```

	Forår	Sommer	Efterår	Vinter
Øjeblikkelig	44	10	32	18
Udskydes	20	8	4	2
Ignoreres	11	2	10	4

Kommunen vil teste, om skadestype-fordelingen afhænger af årstiden. Specifikt vil den teste nulhypotesen

$$H_0 : p_{i1} = p_{i2} = p_{i3} = p_{i4} = p_i \text{ for } i = 1, 2, 3,$$

hvor p_{ij} repræsenterer andelen af observationer i kolonne j , der tilhører række i , dvs. andelen af skadestype i ud af alle skaderne i årstid j , og p_i repræsenterer andelen af alle observationer, der tilhører række i .

Spørgsmål VIII.1 (20)

Hvad er den estimerede andel af skader, der skal repareres øjeblikkeligt, under nulhypotesen?

- 1 ☐ 0.6303
- 2 ☐ 0.4231
- 3 ☐ 0.2061
- 4 ☐ 0.7939
- 5 ☐ Ingen af ovenstående svar er korrekte.

Spørgsmål VIII.2 (21)

Hvad er konklusionen på den sædvanlige test af nulhypotesen ved signifikansniveauet $\alpha = 0.05$? (Både argumentet og konklusionen skal være korrekte.)

- 1 ☐ Nulhypotesen afvises, da p -værdien er 0.04, hvilket indikerer en signifikant forskel på fordelingerne på tværs af årstiderne.
- 2 ☐ Nulhypotesen afvises, da p -værdien er 0.08, hvilket indikerer en signifikant forskel på fordelingerne på tværs af årstiderne.
- 3 ☐ Nulhypotesen accepteres, da p -værdien er 0.04, hvilket ikke indikerer nogen signifikant forskel på fordelingerne på tværs af årstiderne.
- 4 ☐ Nulhypotesen accepteres, da p -værdien er 0.08, hvilket ikke indikerer nogen signifikant forskel på fordelingerne på tværs af årstiderne.
- 5 ☐ Der kan ikke drages en konklusion, da nogle forventede celleantal under nulhypotesen er for lave, dvs. $e_{ij} < 5$ for nogle celler.

Et andet formål var at teste et AI-system til automatisk håndtering af de rapporterede skader ved at sammenligne dets resultater med manuelle vurderinger, hvor en fagperson har kategoriseret skaderne. Baseret på ekspertviden er det forventede antal af uoverensstemmelser mellem de to metoder blevet opgjort til 21 (ud af de 165 observerede skader).

Spørgsmål VIII.3 (22)

Lad $X \sim \text{Bin}(n, p)$ være antallet af uoverensstemmelser mellem de to metoder (man skal estimere p). Hvad er estimatet af standardafvigelsen for X/n ?

- 1 ☐ $\hat{\sigma} = 0.3343$
- 2 ☐ $\hat{\sigma} = 0.1273$
- 3 ☐ $\hat{\sigma} = 0.1118$
- 4 ☐ $\hat{\sigma} = 0.0259$
- 5 ☐ $\hat{\sigma} = 0.0007$

Fortsæt på side 18

Opgave IX

En forsker designer et forsøg til at estimere gennemsnitstiden for en kemisk reaktion. Forsøget skal kunne påvise en forskel i middelværdien på 0.5 sekunder under antagelsen, at standardafvigelsen for reaktionstiden er 1.2 sekunder (tidligere studier viser, at dette er det bedste estimat af σ). Forskeren vil bruge et konfidensniveau på 95% og ønsker, at forsøget skal have en styrke på 80%.

Spørgsmål IX.1 (23)

Hvad er den mindste stikprøvestørrelse, der kræves for at opnå de ønskede specifikationer? (Bemærk at facit er beregnet med Python. Derfor kan facit afvige en smule fra resultater, der findes med formler.)

1 ☐ 38

2 ☐ 47.17

3 ☐ 48

4 ☐ 63

5 ☐ 79

Fortsæt på side 19

Opgave X

Et forsøg tester $n = 85$ personer for et binært respons (positiv/negativ). Lad den binomialfordelte stokastiske variabel

$$X \sim \text{Bin}(n, p)$$

betegne antallet af positive svar (responser).

Følgende nulhypotese betragtes

$$H_0 : p = 0.2.$$

Spørgsmål X.1 (24)

Hvad er sandsynligheden for at observere 25 eller flere positive svar under nulhypotesen (dvs. antaget, at nulhypotesen er sand)?

1 ☐ 0.025

2 ☐ 0.030

3 ☐ 0.061

4 ☐ 0.132

5 ☐ 0.987

Spørgsmål X.2 (25)

Hvad er den observerede teststørrelse i den sædvanlige test af nulhypotesen, hvis der blev observeret præcis 25 positive svar?

1 ☐ 8.000

2 ☐ 2.169

3 ☐ 1.904

4 ☐ 0.057

5 ☐ 0.030

Fortsæt på side 20

Opgave XI

Forskere har foretaget en undersøgelse for at sammenligne gennemsnitsscoren i en matematiktest for to grupper studerende. Gruppe A består af 12 studerende med en gennemsnitsscore på 78 og en standardafvigelse på 10, medens Gruppe B består af 15 studerende med en gennemsnitsscore på 82 og en standardafvigelse på 8. Forskerne gennemfører en t -test med to stikprøver (two sample t -test) for at bestemme, om der er en signifikant forskel på middelscoren af de to grupper ved et signifikansniveau på 5%, dvs. de tester nulhypotesen

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

mod en tosidet modhypotese.

Spørgsmål XI.1 (26)

Hvad er den observerede teststørrelse i den passende test?

- 1 ☐ -3.422
- 2 ☐ -1.155
- 3 ☐ -1.127
- 4 ☐ -0.765
- 5 ☐ -0.317

Spørgsmål XI.2 (27)

Teststørrelsen følger en t -fordeling med hvor mange frihedsgrader?

- 1 ☐ 20.853
- 2 ☐ 22.382
- 3 ☐ 24.998
- 4 ☐ 25
- 5 ☐ 26

Fortsæt på side 21

Opgave XII

Betragt den stokastiske variabel Y og stikprøven $\mathbf{y} = (6, 5, 4, 4, 1, 10, 7, 6, 5, 3)$.

Spørgsmål XII.1 (28)

Hvilket 90% konfidensinterval for middelværdien af Y findes ved brug af ikke-parametrisk bootstrapping med 100,000 replikationer (bootstrap-stikprøver)?

(Bemærk facit bruger `np.random.seed(2025)`, og at dit svar kan variere afhængigt af seedet.)

1 ☐ [3.7,6.6]

2 ☐ [3.9,6.3]

3 ☐ [4.2,6.0]

4 ☐ [5.0,5.2]

5 ☐ [5.1,5.1]

Spørgsmål XII.2 (29)

Antag, at Y følger en Poissonfordeling. Hvilket 90% konfidensinterval for middelværdien af Y findes ved brug af parametrisk bootstrapping med 100,000 replikationer (bootstrap-stikprøver)?

(Bemærk facit bruger `np.random.seed(2025)`, og at dit svar kan variere afhængigt af seedet.)

1 ☐ [3.7,6.6]

2 ☐ [4.0,6.3]

3 ☐ [4.2,6.0]

4 ☐ [5.0,5.2]

5 ☐ [5.1,5.1]

Det følgende stykke kode er blevet kørt:

```
# Sæt seed
np.random.seed(2025)

# Data
y = np.array([6, 5, 4, 4, 1, 10, 7, 6, 5, 3])

# Estimeret gennemsnit
mu = y.mean()

# Simuler bootstrap-stikprøver
simsamples = stats.poisson.rvs(mu, size=(1000, len(y)))
simsamples = pd.DataFrame(simsamples)

# Udregn medianer
simmedians = simsamples.median(axis=1)

# Find konfidensinterval
print(np.quantile(simmedians, [0.05, 0.95],
method='averaged_inverted_cdf'))

[3.5 6.5]
```

Spørgsmål XII.3 (30)

Betragt nulhypotesen $\mathcal{H}_0 : \text{median}(Y) = 5$. Hvad er den korrekte konklusion om $\mathcal{H}_0$ baseret på koden?

- 1 ☐ Den parametriske bootstrapping indikerer, at $\mathcal{H}_0$ ikke kan afvises ved et signifikansniveau på 90%.
- 2 ☐ Den parametriske bootstrapping indikerer, at $\mathcal{H}_0$ afvises ved et signifikansniveau på 10%.
- 3 ☐ Den parametriske bootstrapping indikerer, at $\mathcal{H}_0$ ikke kan afvises ved et signifikansniveau på 10%.
- 4 ☐ Den ikke-parametriske bootstrapping indikerer, at $\mathcal{H}_0$ afvises ved et signifikansniveau på 5%.
- 5 ☐ Den ikke-parametriske bootstrapping indikerer, at $\mathcal{H}_0$ ikke kan afvises ved et signifikansniveau på 10%.

SÆTTET ER SLUT. Nyd ferien!