



Stockholms
universitet

Utökad Lee-Carter för gemensam dödlighets- modellering av likartade populationer

Alexander Käll

Masteruppsats 2025:9
Försäkringsmatematik
Juni 2025

www.math.su.se

Matematisk statistik
Matematiska institutionen
Stockholms universitet
106 91 Stockholm



Matematisk statistik
Stockholms universitet
Masteruppsats 2025:9
<http://www.math.su.se/matstat>

Utökad Lee-Carter för gemensam dödlighetsmodellering av likartade populationer

Alexander Käll*

Juni 2025

Sammanfattning

Den gemensamma faktormodellen är en naturlig utökning av Lee-Cartermodellen för att modellera dödlighet i flera populationer. Denna modell inkluderar både gemensamma och delpopulationsspecifika termer vilket tillåter olika delpopulationer att ha olika tidstrender i dödlighet på kort sikt, men en gemensam trend på lång sikt. I detta arbete utvärderas den gemensamma faktormodellen genom att undersöka hur skattningsfelet i parameterskattningar påverkas av varierande populationsstorlek, hur antalet delpopulationsspecifika termer påverkar dess prediktiva förmåga och hur stort prediktionsfelet blir vid prognoser framåt i tiden. Dödlighetsdata för svenska män och kvinnor och för den totala populationen i Sverige, Norge, Finland och Danmark har använts. Vi visar att skattningsfelet i parameterskattningarna för den gemensamma faktormodellen starkt beror på populationsstorlek. Vi visar även att antalet delpopulationsspecifika termer som ger lägst Poissondevians varierar beroende på vilket tidsintervall som används för att anpassa modellen men att en till tre delpopulationsspecifika termer kan ge en bra balans mellan prediktiv förmåga och att undvika överanpassning. Slutligen visar vi att prediktionsfel i prognoser för förväntad livslängd främst kommer från variation i de tidsserier som tillämpas för att göra prognoser, och att fler delpopulationsspecifika termer och Poissonvariation bidrar relativt lite till detta fel.

*Postadress: Matematisk statistik, Stockholms universitet, 106 91, Sverige.
E-post: alexanderkall42@gmail.com. Handledare: Mathias Millberg Lindholm.

Tackord

Jag vill tacka min handledare Mathias Millberg Lindholm för sin hjälp och vägledning under arbetets gång. Jag vill även tacka Stig H. Olsson som med sitt mod och sin ihärdighet i alla stunder är en ständig inspiration.

Redogörande för användning av AI

ChatGPT och liknande har använts som hjälpmedel i detta arbete, främst för att skriva R-kod, för hjälp med $\text{\LaTeX}$ och för språklig hjälp.

Innehåll

1	Inledning	4
2	Dödlighetsmodeller	5
2.1	Lee-Cartermodellen	5
2.2	Den gemensamma faktormodellen för flera populationer	6
3	Modellanpassning	7
3.1	Maximum likelihood	8
3.2	Algoritm	8
4	Prognoser	10
4.1	Skattnings- och prediktionsfel	11
4.2	Förväntad livslängd	11
5	Dödlighetsdata	13
6	Simuleringsstudie	15
6.1	Populationsstorlek	16
6.2	Antalet termer	17
6.3	Prognosernas fel	18
7	Resultat	19
7.1	Populationsstorlek	19
7.2	Antal specifika termer	21
7.3	Prognosernas fel	23
8	Diskussion	30
9	Slutsats	32
	Bilagor	35

1 Inledning

Mellan år 1920 och 2020 har den förväntade livslängden i Sverige ökat från 58 till 82 år (se Bilaga 1 för beräkningar) och 2120 förväntas den enligt Statistiska Centralbyrån [8] vara 93 år. Hur denna siffra utvecklas har stora konsekvenser inom livförsäkring och pension där försäkringstagarens förväntade livslängd är central för att göra korrekta beräkningar gällande framtida utbetalningar.

För att kunna säga något om hur medellivslängden i en population utvecklas över tid används dödlighetsmodeller. Dessa modeller inkluderar vanligtvis termer som beror på tid och ålder, men kan även inkludera termer som beror på andra faktorer. Dessa anpassas till historisk data och sedan kan de användas för att göra prognoser framåt i tiden.

En sådan modell är Lee-Cartermodellen (LC-modellen) som introducerades av Lee och Carter 1992 i [3]. Den inkluderar enbart termer som beror på tid och ålder för att skapa en enkel men kraftfull modell som ser stor användning än idag. Sedan dess introduktion har LC-modellen utvecklats vidare i försök att göra den ännu bättre i vissa avseenden. En sådan är vid modellering av olika men likartade populationer såsom män och kvinnor i samma land eller flera länder inom ett geografiskt område. Att anpassa separata LC-modeller till varje delpopulation kan leda till ökande skillnader i dödlighet, något som inte överensstämmer med data, se [5]. Flera förändringar till LC-modellen har föreslagits för att motverka detta problem.

En sådan förbättring introducerades av Li och Lee 2005 i [5], vars modell till stor del liknar LC-modellen men som innehåller termer som är gemensamma för de olika delpopulationerna. Denna modell vidareutvecklades av Li 2013 i [4] till det som kommer att kallas den gemensamma faktormodellen (GF-modellen, engelska: Poisson common factor model) hädanefter.

Syftet med detta arbete är att utvärdera GF-modellen. Specifikt kommer vi att utreda hur den påverkas av varierande populationsstorlek, hur antalet delpopulationsspecifika termer påverkar dess prediktiva förmåga och hur stort prediktionsfelet blir vid prognoser framåt i tiden.

Vi visar att skattningsfelet i parameterskattningarna för GF-modellen ökar när populationsstorleken minskar, men i olika grad för de olika parametrarna. Vi visar också att antalet delpopulationsspecifika termer som ger lägst Poissondevians ändras beroende på vilket tidsintervall som används för att anpassa modellen, men att en till tre termer kan ge en bra balans mellan prediktiv förmåga och att undvika överanpassning. Slutligen visar vi att

prediktionsfel av prognoser för förväntad livslängd vid födsel främst kommer från variation i de tidsserier som används för att göra prognoser. Flera delpopulationsspecifika termer och Poissonvariation bidrar relativt lite till detta fel.

Vi börjar med att beskriva LC-modellen och GF-modellen i del 2, i del 3 beskrivs hur modellen anpassas till data, del 4 förklarar hur modellen kan användas för att göra prognoser framåt i tiden och hur vi beräknar förväntad livslängd, i del 5 beskrivs dödlighetsdatan som används och i del 6 beskrivs hur simuleringarna gjordes. Resultaten presenteras i del 7 och diskuteras i del 8. Slutsats ges i del 9.

2 Dödlighetsmodeller

Dödlighetsmodeller är modeller som beskriver dödligheten i någon population. Vanligtvis inkluderar dessa termer som beror på ålder, x , och tid, t . Låt $\mu_{x,t}$ vara dödlighetsintensiteten för ålder x och tid t och låt $m_{x,t}$ vara kvoten mellan antalet dödsfall $D_{x,t}$ och riskexponeringen $E_{x,t}$ för x år gamla individer under år t , där riskexponeringen är den totala tid som har levts av x år gamla individer under år t , d.v.s.

$$m_{x,t} = \frac{D_{x,t}}{E_{x,t}}.$$

Fortsättningsvis kommer vi att använda $\mu_{x,t}$ för att beteckna dödlighetsintensitet som antas vara styckvis konstant över ettårsintervall för ålder och tid, även om sann dödlighetsintensitet inte är det.

2.1 Lee-Cartermodellen

I en artikel från 1992 introducerade Lee och Carter i [3] sin modell för dödlighet som numera kallas Lee-Cartermodellen (LC-modellen). Det som särskiljer deras modell från tidigare modeller är att den inte inkluderar vetenskap om till exempel medicinska eller beteendemässiga inflytanden på dödlighet utan gör minimala antaganden. Deras modell säger att

$$\log(\mu_{x,t}) = a(x) + b(x)k(t) + \epsilon_{x,t}$$

där $\epsilon_{x,t}$ är en felterm med väntevärde 0 och varians σ_ϵ^2 . Modellparametrarna kan tolkas så här: $a(x)$ är den grundläggande formen för dödlighetsintensiteten som en funktion av ålder, $k(t)$ beskriver hur dödlighetsintensiteten

utvecklas över tid och $b(x)$ beskriver hur mycket tidstrenden påverkar dödlighetsintensiteten för olika åldrar. Av dessa bidrar $a(x)$ som regel mest till dödlighetsintensiteten medan $b(x)$ och $k(t)$ tillåter den att förändras sakta över tid.

I originalartikeln från Lee och Carter används äldre metoder för modellanpassning men sedan 2002, då Brouhns et al. [2] introducerade ett Poissonantagande, där antalet dödsfall modelleras direkt enligt

$$D_{x,t} \mid E_{x,t} \sim \text{Poisson}(E_{x,t}\mu_{x,t}),$$

där $\mu_{x,t} = \exp(a(x) + b(x)k(t))$, möjliggjordes modellanpassning med hjälp av maximum-likelihood. LC-modellen har sett mycket användning tack vare sina stabila prognoser med tillhörande konfidensintervall [3] och ser fortsatt användning idag, men den har vissa tillkortakommanden. Ofta vill man modellera delpopulationer, såsom män och kvinnor i samma land, olika då deras dödlighet skiljer sig åt. Om man anpassar separata Lee-Cartermodeller till dessa delpopulationer har det visats vara problematiskt då dessa kan förutspå att skillnaden i dödlighet mellan delpopulationerna ökar över tid, något som inte stärks av data [5]. För att åtgärda detta kan vi använda en modell där de olika delpopulationerna har både gemensamma och specifika termer, och en sådan är den gemensamma faktormodellen.

2.2 Den gemensamma faktormodellen för flera populationer

För att åtgärda ett av Lee-Cartermodellens tillkortakommanden introducerade Li och Lee i [5] modellen

$$\log(\mu_{x,t,i}) = a(x, i) + B(x)K(t) + b(x, i)k(t, i) + \epsilon_{x,t,i}$$

där $B(x)$ och $K(t)$ är gemensamma för hela populationen och resten av termerna är specifika för delpopulation i . Tanken bakom denna modell är att olika delpopulationer kan ha olika trender på kort sikt som fångas av $b(x, i)$ och $k(t, i)$, men en gemensam trend på lång sikt. Modellparametrarna kan tolkas på ett liknande sätt som för Lee-Cartermodellen. Denna modell utvecklades vidare av Li i [4] som föreslog modellen

$$\log(\mu_{x,t,i}) = a(x, i) + B(x)K(t) + \sum_{j=1}^n b(x, i, j)k(t, i, j)$$

som kan innehålla ett godtyckligt antal specifika termer för varje delpopulation. Även här gör vi ett poissonantagande och modellerar antalet dödsfall direkt enligt

$$D_{x,t,i} \mid E_{x,t,i} \sim \text{Poisson}(E_{x,t,i}\mu_{x,t,i}), \quad (1)$$

där $\mu_{x,t,i} = \exp\left(a(x,i) + B(x)K(t) + \sum_{j=1}^n b(x,i,j)k(t,i,j)\right)$, vilket möjliggör modellanpassning genom maximum likelihood och simulering av ny data genom parametrisk bootstrap.

3 Modellanpassning

När Lee och Carter [3] introducerade sin modell skattade de parametrarna genom att använda metoder som sällan används idag. Numera anpassas den vanligtvis med maximum likelihood och det gör även den gemensamma faktormodellen. För att kunna göra det behöver vi log-likelihoodfunktionen. För en Poissonmodell ges den av

$$l(\hat{\mu}_{x,t,i}) = \sum_{x,t,i} (D_{x,t,i} \log(E_{x,t,i}\hat{\mu}_{x,t,i}) - E_{x,t,i}\hat{\mu}_{x,t,i} - \log(D_{x,t,i}!))$$

där

$$\hat{\mu}_{x,t,i} = \exp\left(\hat{a}(x,i) + \hat{B}(x)\hat{K}(t) + \sum_{j=1}^n \hat{b}(x,i,j)\hat{k}(t,i,j)\right). \quad (2)$$

Från detta kan vi beräkna Poissondeviansen som blir

$$\begin{aligned} d &= 2 \sum_{x,t,i} \left(l\left(\frac{D_{x,t,i}}{E_{x,t,i}}\right) - l(\hat{\mu}_{x,t,i}) \right) \\ &= 2 \sum_{x,t,i} \left(D_{x,t,i} \log\left(\frac{D_{x,t,i}}{E_{x,t,i}\hat{\mu}_{x,t,i}}\right) - D_{x,t,i} + E_{x,t,i}\hat{\mu}_{x,t,i} \right). \end{aligned} \quad (3)$$

Modellparametrarna skattas genom att minimera (3). Notera att den gemensamma termen är ekvivalent med motsvarande term i Lee-Cartermodellen när den anpassas till hela populationen.

3.1 Maximum likelihood

Poissonantagandet möjliggör modellanpassning genom maximum likelihood, men på grund av de bilinjära termerna $B(x)K(t)$ och $\sum_{j=1}^n b(x, i, j)k(t, i, j)$ går det inte att anpassa modellen direkt, utan istället skattas parametrarna iterativt genom Newton-Raphsonmetoden. De initialiseras till lämpliga värden, t.ex. konstanter och sedan uppdateras de enligt ekvationen

$$\theta^* = \theta - \frac{\frac{\partial l}{\partial \theta}}{\frac{\partial^2 l}{\partial \theta^2}} = \theta - \left(\frac{\partial^2 l}{\partial \theta^2}\right)^{-1} \frac{\partial l}{\partial \theta}$$

tills Poissondeviansen konvergerar. Här är $\frac{\partial l}{\partial \theta}$ en $(X(n+2) + T(n+1)) \times 1$ -vektor och $\frac{\partial^2 l}{\partial \theta^2}$ en $(X(n+2) + T(n+1)) \times (X(n+2) + T(n+1))$ -matris, där X är antalet av varje åldersberoende parameter och där T är antalet av varje tidsberoende parameter. Maximum likelihood-skattningarna som fås är endast unika upp till en multiplikativ konstant och behöver genomgå transformationer för att bli identifierbara. Li [4] föreslår en restriktion där $\sum_x B(x) = \sum_x b(x, i, j) = 1$ och $\sum_t K(t) = \sum_t k(t, i, j) = 0$, vilket uppfylls efter transformationerna

$$\begin{aligned} B(x) &\rightarrow B(x) / \sum_x B(x), & K(t) &\rightarrow \sum_x B(x)(K(t) - \bar{K}(t)), \\ b(x, i, j) &\rightarrow b(x, i, j) / \sum_x b(x, i, j), & k(t, i, j) &\rightarrow \sum_x b(x, i, j)(k(t, i, j) - \bar{k}(t, i, j)) \end{aligned}$$

där $\bar{K}(t)$ och $\bar{k}(t, i, j)$ är medelvärdena över tid av $K(t)$ och $k(t, i, j)$.

3.2 Algoritm

Li [4] beskriver en algoritm som använder (3.1) för alla termer och itererar tills Poissondeviansen konvergerar. Dock menar Li att användning av denna algoritmen kan resultera i problem med konvergens, och som ett alternativ föreslås att den gemensamma faktorn skattas först, följt av de delpopulations-specifika faktorerna en i taget. Denna alternativa algoritm har använts i detta arbete och ser ut så här:

Del 1: Initiering av parametrarna

1. Initiera parametrarna enligt $\hat{a}(x, i) = \frac{1}{T} \sum_t \log\left(\frac{D_{x,t,i}}{E_{x,t,i}}\right)$ där T är antalet år som summeras, $\hat{B}(x) = \hat{b}(x, i, j) = \frac{1}{90}$ och $\hat{K}(t) = \hat{k}(t, i, j) = 0$ för alla x , i och j , och beräkna $\hat{D}_{x,t,i}$.

Del 2: Skattning av $a(x, i)$ och gemensam term

2. Uppdatera $\hat{a}^*(x, i) = \hat{a}(x, i) + \sum_t (D_{x,t,i} - \hat{D}_{x,t,i}) / \sum_t \hat{D}_{x,t,i}$ för alla x och i , och beräkna $\hat{D}_{x,t,i}$.
3. Uppdatera $\hat{K}^*(t) = \hat{K}(t) + \sum_{x,i} (D_{x,t,i} - \hat{D}_{x,t,i}) \hat{B}(x) / \sum_{x,i} \hat{D}_{x,t,i} \hat{B}^2(x)$ för alla t , justerat så att $\sum_t \hat{K}^*(t) = 0$, och beräkna $\hat{D}_{x,t,i}$.
4. Uppdatera $\hat{B}^*(x) = \hat{B}(x) + \sum_{t,i} (D_{x,t,i} - \hat{D}_{x,t,i}) \hat{K}(t) / \sum_{t,i} \hat{D}_{x,t,i} \hat{K}^2(t)$ för alla x , och beräkna $\hat{D}_{x,t,i}$.
5. Beräkna Poissondeviansen.
6. Repetera steg 2-5 tills Poissondeviansen konvergerar.

Del 3: Skattning av delpopulationsspecifika termer

7. Uppdatera $\hat{k}^*(t, i, j) = \hat{k}(t, i, j) + \sum_x (D_{x,t,i} - \hat{D}_{x,t,i}) \hat{b}(x, i, j) / \sum_x \hat{D}_{x,t,i} \hat{b}^2(x, i, j)$ för alla t , i och $j = 1$, justerat så att $\sum_t \hat{k}^*(t, i, j) = 0$, och beräkna $\hat{D}_{x,t,i}$.
8. Uppdatera $\hat{b}^*(x, i, j) = \hat{b}(x, i, j) + \sum_t (D_{x,t,i} - \hat{D}_{x,t,i}) \hat{k}(t, i, j) / \sum_t \hat{D}_{x,t,i} \hat{k}^2(t, i, j)$ för alla x , i och $j = 1$, och beräkna $\hat{D}_{x,t,i}$.
9. Repetera steg 7 och 8 för alla värden av j .
10. Beräkna Poissondeviansen.
11. Repetera steg 7-9 tills Poissondeviansen konvergerar.
12. Justera parameterskattningarna så att $\sum_x B(x) = \sum_x b(x, i, j) = 1$ och $\sum_t K(t) = \sum_t k(t, i, j) = 0$.

I praktiken använde steg 6 och 11 tio repetitioner istället för tills konvergens. Detta gjordes för att förenkla implementering och visade sig fungera bra, särskilt för låga värden av n (se bilaga 2 för mer information). Li [4] initierar $\hat{B}(x)$ och $\hat{b}(x, i, j)$ till $\frac{1}{90}$ utan förklaring. Dock fungerade detta värde bra och har ej undersökts närmare.

De delpopulationsspecifika termerna anpassas till residualerna som fås av de redan skattade termerna. Man kan säga att varje term försöker förklara den variation som ännu inte har förklarats av tidigare termer. Detta resulterar i att varje delpopulationsspecifik term bidrar mindre och mindre till den totala dödlighetsintensiteten.

4 Prognoser

Utöver att bättre förstå en populations dödlighet syftar dödlighetsmodeller till att göra prognoser för framtida dödlighet. För att göra detta antas $a(x, i)$, $B(x)$ och $b(x, i, j)$ vara konstanta över tid medan $K(t)$ och $k(t, i, j)$ modelleras som tidsserier. Då $K(t)$ ska representera den långsiktiga trenden för hela populationen och $k(t, i, j)$ ska fänga kortsiktiga avvikelser används olika tidsserier av typen autoregressiva modeller. Till dessa tillhör ARIMA-modeller som bestäms av tre parametrar (p, d, q) där p är ordningen av den autoregressiva delen, d är graden av differentiering och q är ordningen av det glidande medelvärdet. Beroende på hur dessa parametrar väljs fås tidsserier med olika egenskaper. Lee och Carter [3] föreslår en slumpvandring med drift för $k(t)$ i sin modell, vilket är en ARIMA(0,1,0)-modell. Detta motsvarar en vanlig slumpvandring men som har en linjär, vanligtvis negativ, trend.

Li [4], liksom Lee och Carter, tillämpar en slumpvandring med drift (ARIMA(0,1,0)) för den gemensamma termen,

$$\tilde{K}(t+h) = C + \tilde{K}(t+h-1) + \epsilon_{t+h}, \quad (4)$$

där h är antalet år sedan prognosens start, t är det sista året som användes vid anpassning, C är driften och $\epsilon_{t+h} \sim N(0, \sigma^2)$. Detta kommer att ge alla delpopulationer en gemensam långsiktig trend. De delpopulationsspecifika termerna modelleras enligt,

$$\tilde{k}(t+h, i, j) = \alpha_{0,i} + \alpha_{1,i} \tilde{k}(t+h-1, i, j) + \epsilon_{t+h,i,j}$$

där $|\alpha_{1,i}| < 1$ gör att den här processen går mot en konstant över tid, vilket medför att det endast är den gemensamma termen som bidrar på lång sikt [4], och där $\epsilon_{t,i,j} \sim N(0, \sigma_{i,j}^2)$. Dessa kan nu användas för att göra prognoser för dödlighetsintensiteten enligt

$$\tilde{\mu}_{x,t+h,i} = \exp \left(\hat{a}(x, i) + \hat{B}(x) \tilde{K}(t+h) + \sum_{j=1}^n \hat{b}(x, i, j) \tilde{k}(t+h, i, j) \right) \quad (5)$$

som kan jämföras med observerade värden eller användas för att simulera ny data.

4.1 Skattnings- och prediktionsfel

Prediktionsintervall för $\tilde{K}(t+h)$ och $\tilde{k}(t+h, i, j)$ kan fås explicit via standardtekniker, men detta är betydligt svårare att göra för $\tilde{\mu}_{x,t+h,i}$ som innehåller flera tidsserier. Ett sätt att få en uppfattning om prediktionsfelet i prognoser för dödlighetsintensitet är att simulera trajektorier för de olika tidsserierna och sedan skatta den framtida dödlighetsintensiteten enligt (5). Prediktionsintervall för dödlighetsintensiteten kan sedan fås genom att titta på relevanta percentiler.

Om vi tittar på den gemensamma faktormodellen med $n = 0$ har vi bara en term som beror på tid, nämligen $K(t)$, som modelleras som en slumpvandring med drift enligt (4). Vi kan skriva om den i sluten form som

$$\tilde{K}(t+h) = (t+h)C + \sum_{i=1}^h \epsilon_i.$$

då vi påbörjar slumpvandringen vid $h = 1$. Variansen fås av

$$\begin{aligned} \text{Var}(\tilde{K}(t+h)) &= \text{Var}(hC + \sum_{i=1}^h \epsilon_i) \\ &= \text{Var}\left(\sum_{i=1}^h \epsilon_i\right) \\ &= h\sigma^2 \end{aligned}$$

då alla ϵ_i antas vara oberoende och likafördelade. Detta kan nu användas för att beräkna $(1 - \alpha)\%$ prediktionsintervall för $\tilde{K}(t+h)$ som fås av

$$[(t+h)C - \Phi^{-1}(1 - \frac{\alpha}{2})\sqrt{h}\sigma, (t+h)C + \Phi^{-1}(1 - \frac{\alpha}{2})\sqrt{h}\sigma] \quad (6)$$

där Φ är den kumulativa fördelningsfunktionen för en standard normalfördelning. Detta gäller som sagt för $n = 0$ men då varje ytterligare term bidrar mindre och mindre till dödlighetsintensiteten bör prediktionsintervallen för $n > 0$ följa ett liknande mönster, alltså att felet växer med roten ur tiden.

4.2 Förväntad livslängd

Våra prognoser resulterar i dödlighetsintensiteter men ofta är vi mer intresserade av förväntad livslängd. Dödlighetsintensiteterna kan användas för

att beräkna den förväntade livslängden, $e_{x,t}$, där $e_{0,t}$ är den förväntade livslängden vid födsel för någon som föds år t . Detta görs periodvis vilket betyder att endast dödlighetsintensiteter från det år som vi beräknar den förväntade livslängden för används. Om dödligheten minskar över tid i en population kan denna metod underskatta den förväntade livslängden. Ett alternativ är att göra beräkningar kohortvis vilket kan ge mer realistiska värden, men det är begränsande då vi måste göra prognoser alternativt vänta tills hela kohorten har dött för att kunna beräkna förväntad livslängd på detta vis. Sannolikheten att överleva från ålder x till $x + 1$ beskrivs i ekvation (1.5), kapitel 1.1.2 av Aalen et al. [1] som

$$P_{x,t} = \exp\left(-\int_x^{x+1} \mu_{y,t} dy\right)$$

och om vi antar att dödlighetsintensiteten är konstant över ettårsintervall kan vi approximera sannolikheten enligt

$$P_{x,t} \approx \exp(-\mu_{x,t}).$$

Sannolikheten att överleva från födsel till ålder x ges av

$$\prod_{x'=0}^x P_{x',t} \approx \prod_{x'=0}^x \exp(-\mu_{x',t}).$$

och nu kan vi beräkna väntevärdet av $e_{0,t}$ som är

$$e_{0,t} \approx \sum_{x=0}^{x_{\max}} \prod_{x'=0}^x \exp(-\mu_{x',t}), \quad (7)$$

där $x_{\max}$ är den högsta åldern som antas kunna uppnås. Här antas även att antalet dödsfall är likafördelat inom varje ettårsintervall för tid och ålder. Alla beräkningar av förväntad livslängd vid födsel i detta arbete har gjorts genom att tillämpa (7).

Låt oss återigen beakta den gemensamma faktormodellen när $n = 0$. Om vi antar att $a(x, i)$ och $B(x)$ är kända och att felet i prognoserna för den förväntade livslängden endast beror på $\tilde{K}(t)$ kan vi approximera variansen av e_{0t} . Enligt Lo et al. [6] ges variansen av $e_{0,t+h}$ av

$$\begin{aligned}
\text{Var}(e_{0,t+h}) &\approx \sum_x \left(\frac{\partial e_{0,t}}{\partial K(t)} \right)^2 \cdot \text{Var}(\tilde{K}(t+h)) \\
&= \sum_x \left(\frac{\partial e_{0,t}}{\partial K(t)} \right)^2 \cdot h\sigma^2 \\
&\propto h\sigma^2
\end{aligned} \tag{8}$$

vilket betyder att variansen av den förväntade livslängden växer linjärt med antalet år sedan prognosens start. Detta gäller endast för $n = 0$, men då varje ytterligare term bidrar mindre och mindre till den skattade dödlighetsintensiteten är det rimligt att variansen för $n > 0$ följer ett liknande mönster.

5 Dödlighetsdata

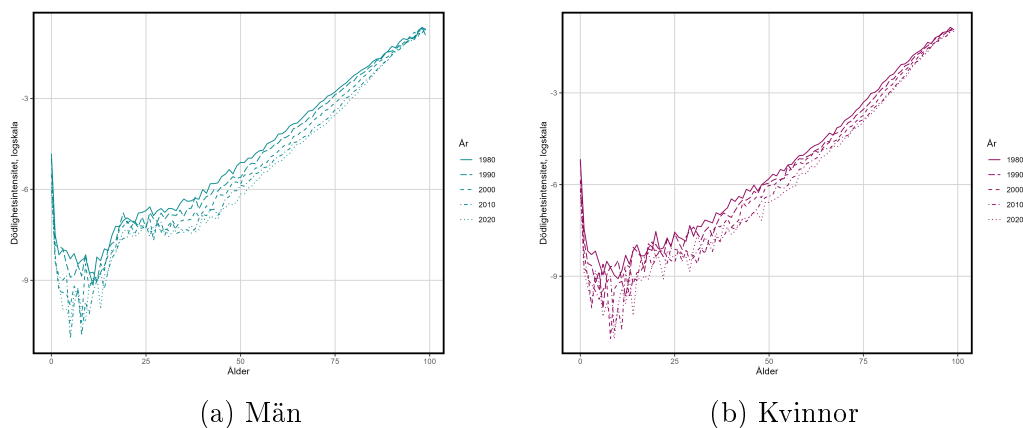
Dödlighetsdata hittas på Human Mortality Database (HMD) [7]. Här finns data från många olika länder och i många olika format, men i detta arbete används svensk dödlighetsdata för män och kvinnor och dödlighetsdata för den totala populationen i Sverige, Norge, Finland och Danmark. Specifikt är vi intresserade av antal dödsfall och riskexponering i ettårsintervall för tid och ålder. Även data för populationsstorlekar hittas på HMD. Dödlighetsdatan sträcker sig ända tillbaka till 1751, men i detta arbete kommer vi främst att använda data från 1950 och framåt. Vi kommer även att begränsa oss till att titta på data för åldrarna 0-99 då data för högre åldrar inte finns för alla årtal.

Kön	Antal
Män	5 312 519
Kvinnor	5 239 188
Totalt	10 551 707

Tabell 1: *Antal svenska invånare, 2023.*

Tabell 1 visar antalet svenska invånare 2023. Antalet svenska män och kvinnor är väldigt lika, så därför är det intressant att även titta på data från fyra olika populationer vars storlekar skiljer sig åt.

Figur 1 visar logaritmen av den observerade dödlighetsintensiteten som en funktion av ålder för åren 1980, 1990, 2000, 2010 och 2020 för män och kvinnor. För de flesta åldrar sjunker dödlighetsintensiteten över tid, vilket indikeras av att kurvorna skiftar nedåt. dödlighetsintensitetskurvorna har



Figur 1: *Logaritmen av dödlighetsintensiteten som en funktion av ålder för några olika år, svensk data för män och kvinnor.*

ungefär samma form för män och kvinnor med en relativt hög dödlighet för nyfödda som sedan sjunker för barn för att sedan börja öka igen. För kvinnor ökar dödligheten i en jämn takt medan den för män ökar snabbt runt 20 års ålder innan den planar ut för att sedan börja öka igen. dödlighetsintensiteten kan användas för att räkna ut förväntad livslängd, en storhet som är mer lättolkad.

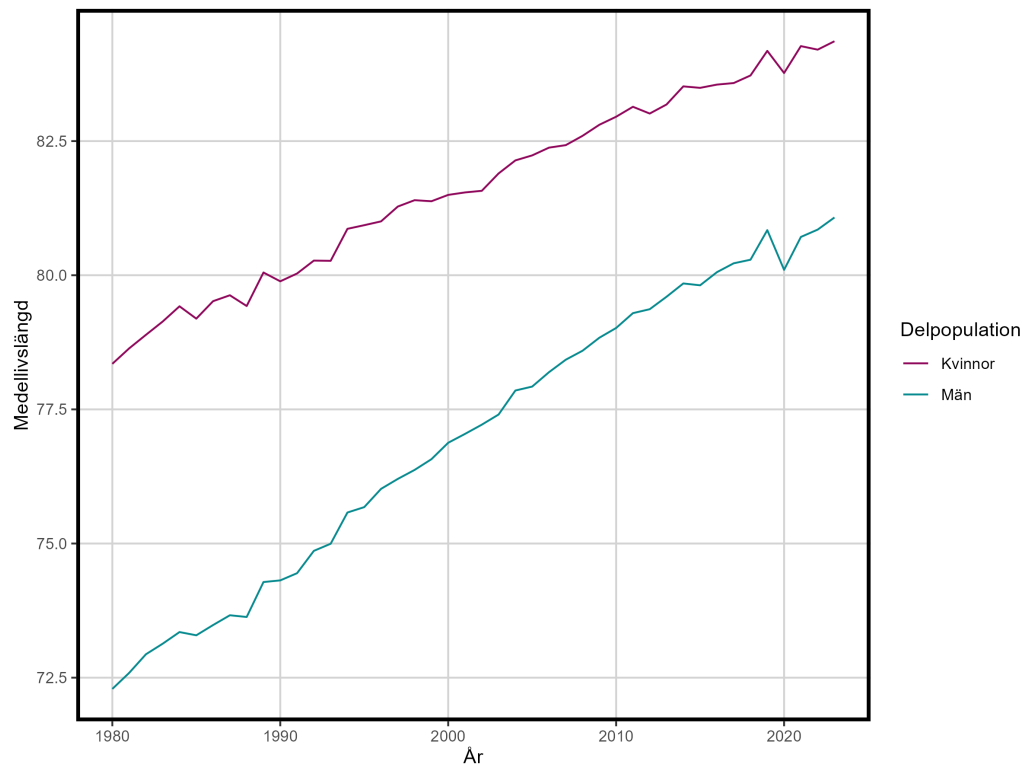
Figur 2 visar hur den förväntade medellivslängden har utvecklats för svenska kvinnor och män år 1980-2023. Skillnaden i förväntad livslängd vid födsel mellan män och kvinnor ser ut att sakta minska under denna tidsperiod.

Land	Antal
Sverige	10 551 707
Norge	5 550 203
Finland	5 603 851
Danmark	5 961 249
Totalt	27 667 010

Tabell 2: *Antal invånare, 2023.*

Tabell 2 visar antalet invånare i Sverige, Norge, Finland och Danmark 2023. Norge, Finland och Danmark har alla liknande antal invånare medan Sverige har betydligt fler.

Figur 3 visar logaritmen av dödlighetsintensiteten som en funktion av ålder för åren 1980, 1990, 2000, 2010 och 2020 för Sverige, Norge, Finland och Danmark. För alla fyra länder syns en skarp ökning i dödlighet runt 20 års



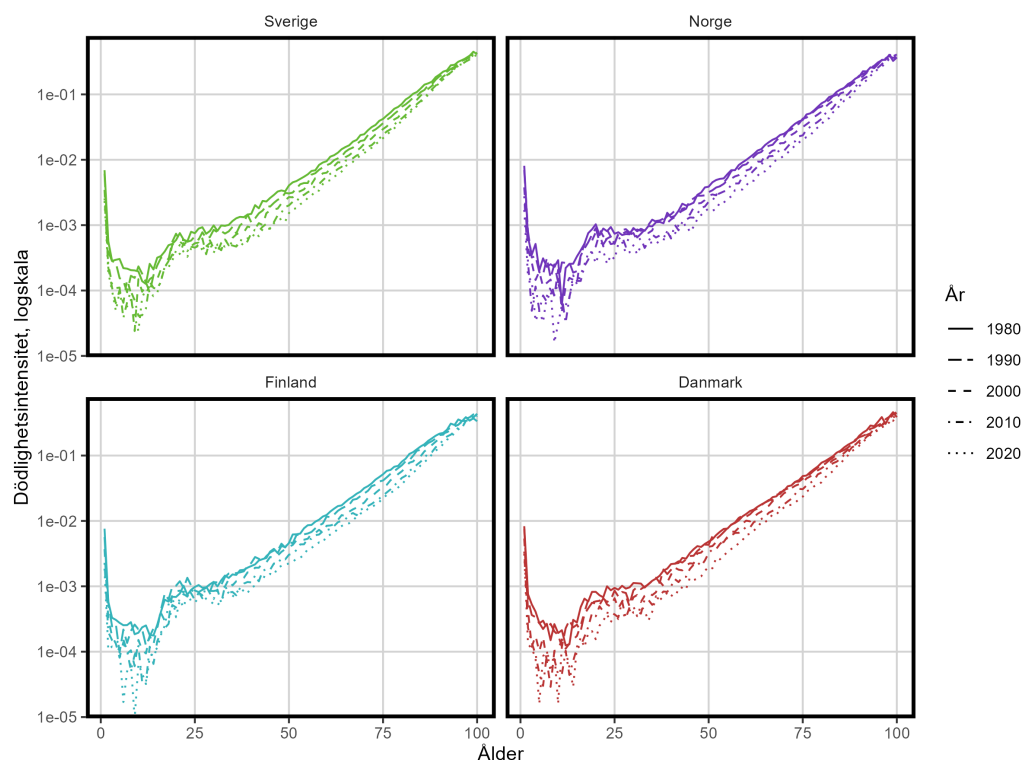
Figur 2: Förväntad livslängd vid födsel i Sverige, 1980-2023.

ålder som sedan planar ut, likt det som observerades i figur 1a.

Figur 4 visar hur den förväntade livslängden har utvecklats för Sverige, Norge, Finland och Danmark under åren 1950-2023. I början av tidsperioden var den förväntade medellivslängden snarlik i Sverige, Norge och Danmark medan den var betydligt lägre i Finland. Sedan har Finland sett en kraftig ökning och har gått om Danmark under delar av tidsperioden. Sedan 1980 har de fyra länderna haft liknande förväntad medellivslängd, ännu en indikation att gemensam faktormodellen är tillämplig.

6 Simuleringsstudie

För att testa den gemensamma faktormodellen (GF-modellen) i några olika avseenden kommer en simuleringsstudie att utföras. Syftet med detta är att få förståelse för hur populationsstorlek påverkar parameterskattningarna, att se hur modellens prediktiva förmåga påverkas av antalet delpopulationsspecifika termer och hur stort prediktionsfelet blir vid framåtblickande prognoser.

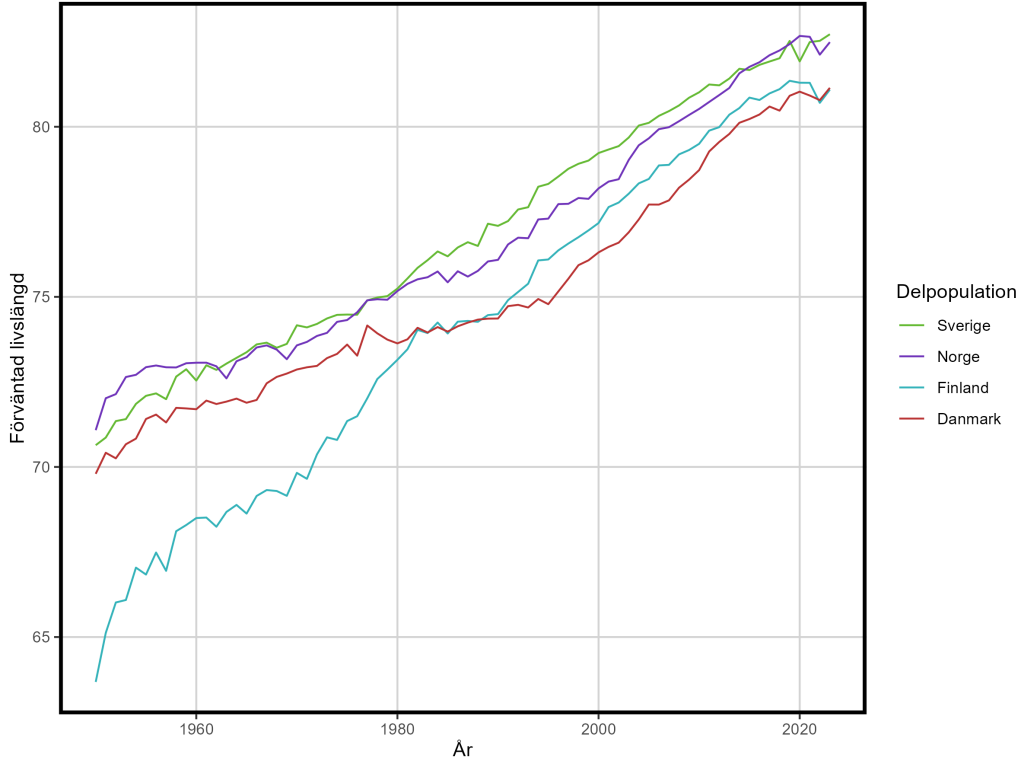


Figur 3: *Logaritmen av dödlighetsintensiteten per land som en funktion av ålder för några olika år, data från Sverige, Norge, Finland och Danmark.*

Simuleringarna gjordes i R och paketen StMoMo, Tidyverse, Demography, Forecast, Reshape2 och HMDHFDplus användes för datahantering och visualisering. Ett paket för anpassning av GF-modellen hittades ej, så algoritmen i del 3.2 behövde implementeras.

6.1 Populationsstorlek

För att undersöka hur populationsstorlek påverkar parameterskattningarna gjordes en simulering där GF-modellen anpassades till simulerad dödlighetsdata med olika populationsstorlekar. GF-modellen med $n = [0]$ anpassades till svensk dödlighetsdata för män och kvinnor från 1950-2000, och separat till dödlighetsdata från Sverige, Norge, Finland och Danmark, enligt algoritmen beskriven i del 3.2. Vi använder $n = 0$ här då det anses tillräckligt för att se effekten av en mindre population. Dessutom bidrar de delpopulations-specifika termerna mindre och mindre till dödlighetsintensiteten. Den anpassade modellen ger skattningar för de olika parametrarna som kan användas



Figur 4: *Förväntad livslängd vid födsel i några olika länder, 1950-2023.*

för att skatta $\hat{\mu}_{x,t,i}$ enligt (2). Med denna kan vi simulera ny dödlighetsdata genom att tillämpa parametrisk bootstrap. Om vi antar att riskexponeringen är given kan vi använda en variant av (1) vilket ger

$$D_{x,t,i}^f \mid E_{x,t,i}^f, \hat{\mu}_{x,t,i} \sim \text{Poisson}(E_{x,t,i}^f \hat{\mu}_{x,t,i})$$

där $E_{x,t,i}^f := \frac{E_{x,t,i}}{f}$ och $f = [1, 4, 16, 64]$ är en faktor som motsvarar förändringar av populationsstorleken. Dessa användes för att på nytt anpassa en GF-modell med $n = 0$ vilket upprepades 1000 gånger för varje f för att kunna skapa 95% konfidensintervall för varje parameter. Konfidensintervallen skapades genom att titta på de 2.5:e och 97.5:e percentilerna av de skattade parametrarna vilket gjordes punktvis för varje år/ålder.

6.2 Antalet termer

GF-modellen är en flexibel modell där antalet termer kan väljas fritt, men denna flexibilitet medför ett problem: hur ska detta antal bestämmas? Detta

undersöktes genom att GF-modellen med $n = [0, 1, \dots, 8]$ anpassades till svensk data och data från Sverige, Norge, Finland och Danmark från olika tidsintervall. Dessa användes sedan för att göra prognoser framåt i tiden, vilket beskrevs i del 4, som sedan utvärderades mot observerad data från 2001-2023, ibland kallat *out-of-sample*. Utvärdering skedde genom att beräkna Poissondeviansen för varje tidsintervall och värde av n . Prognoserna utvärderades mot observerad data från samma tidsperiod så att Poissondeviansen blir direkt jämförbar. Detta gjordes då ett av de huvudsakliga syftena med dödlighetsmodeller är att göra framåtblickande prognoser. Dessutom leder fler parametrar alltid till lägre Poissondevians när vi utvärderar modellen på samma data som användes för anpassningen, så därför är *out-of-sample* Poissondevians ett bättre mått på modellens prediktiva förmåga.

6.3 Prognosernas fel

För att undersöka prediktionsfel i framåtblickande prognoser anpassades GF-modellen med $n = [0, 2]$ till svensk dödlighetsdata och separat till dödlighetsdata från Sverige, Norge, Finland och Danmark från 1980-2000. Sedan simulerades 1000 trajektorier för $K(t)$ och $k(t, i, j)$ som användes för att beräkna den framtida dödlighetsintensiteten enligt (5) som i sin tur användes för att beräkna den förväntade livslängden vid födsel enligt (7). Från dessa fås ett medelvärde och 95% prediktionsintervall genom att titta på 2.5:e och 97.5:e percentilerna. Detta inkluderar endast prediktionsfel från tidsserierna och ej Poissonvariation.

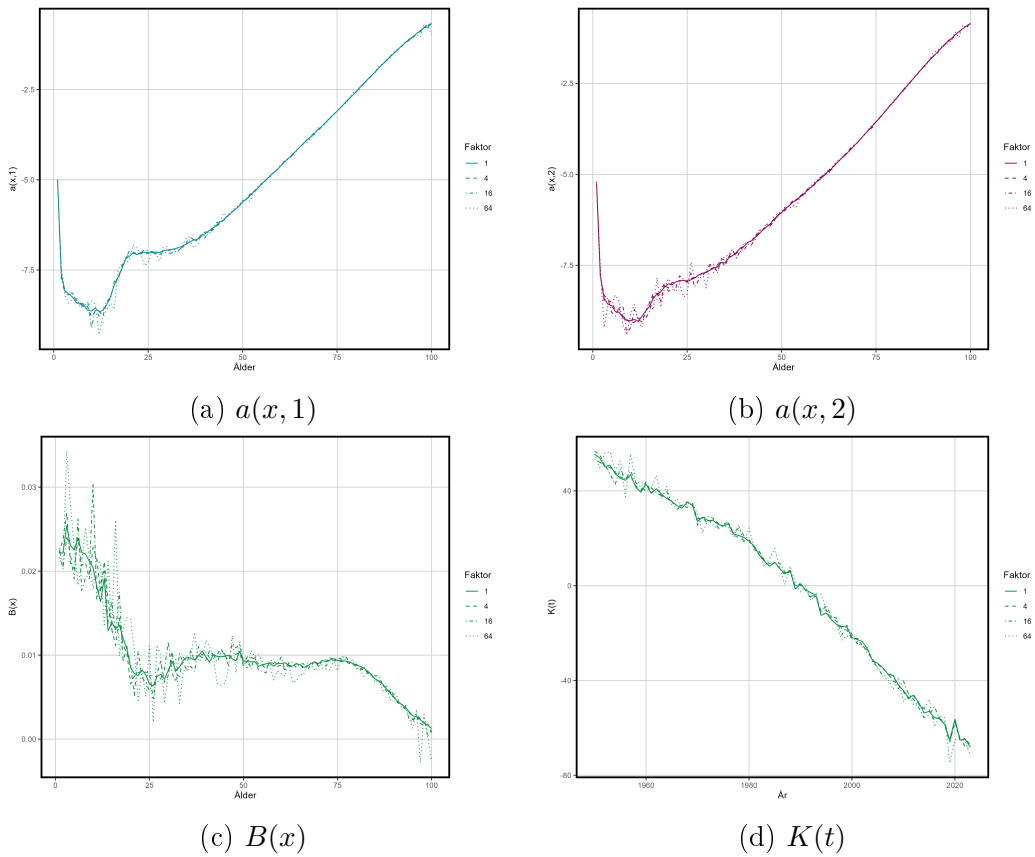
För att få med prediktionsfel från Poissonvariation gjordes ännu en simulering. Först anpassades GF-modellen med $n = [0, 2]$ till svensk dödlighetsdata och separat till dödlighetsdata från Sverige, Norge, Finland och Danmark från 1980-2000. Från den anpassade modellen skattades dödlighetsintensiteten enligt (2). Sedan gjordes en parametrisk bootstrap där nya dödsantal simulerades 1000 gånger enligt $D_{x,t,i}^{\text{sim}} \mid E_{x,t,i} \sim \text{Poisson}(E_{x,t,i} \hat{\mu}_{x,t,i})$. Tillsammans med den observerade riskexponeringen kan dessa användas för att beräkna nya simulerade dödlighetsintensiteter enligt $\mu_{x,t,i}^{\text{sim}} = \frac{D_{x,t,i}^{\text{sim}}}{E_{x,t,i}}$ som i sin tur användes för att på nytt anpassa en GF-modell med $n = [0, 2]$. Den anpassade modellen användes sedan för att göra prognoser för framtida dödlighet enligt (5) som i sin tur användes för att beräkna framtida förväntad livslängd vid födsel enligt (7). Från dessa simuleringar får vi ett medelvärde och 95% prediktionsintervall genom att titta på 2.5:e och 97.5:e percentilerna. Detta görs för att se om prediktionsfelet till följd av Poissonvariation bidrar på ett märkbart sätt till det totala prediktionsfelet i framåtblickande prognoser.

7 Resultat

I denna del kommer resultaten från föregående dels simuleringar att presenteras. Här motsvarar $i = [1, 2]$ svenska män respektive kvinnor för resultaten från svensk data, och $i = [1, 2, 3, 4]$ motsvarar Sverige, Norge, Finland och Danmark för resultaten från de fyra länderna.

7.1 Populationsstorlek

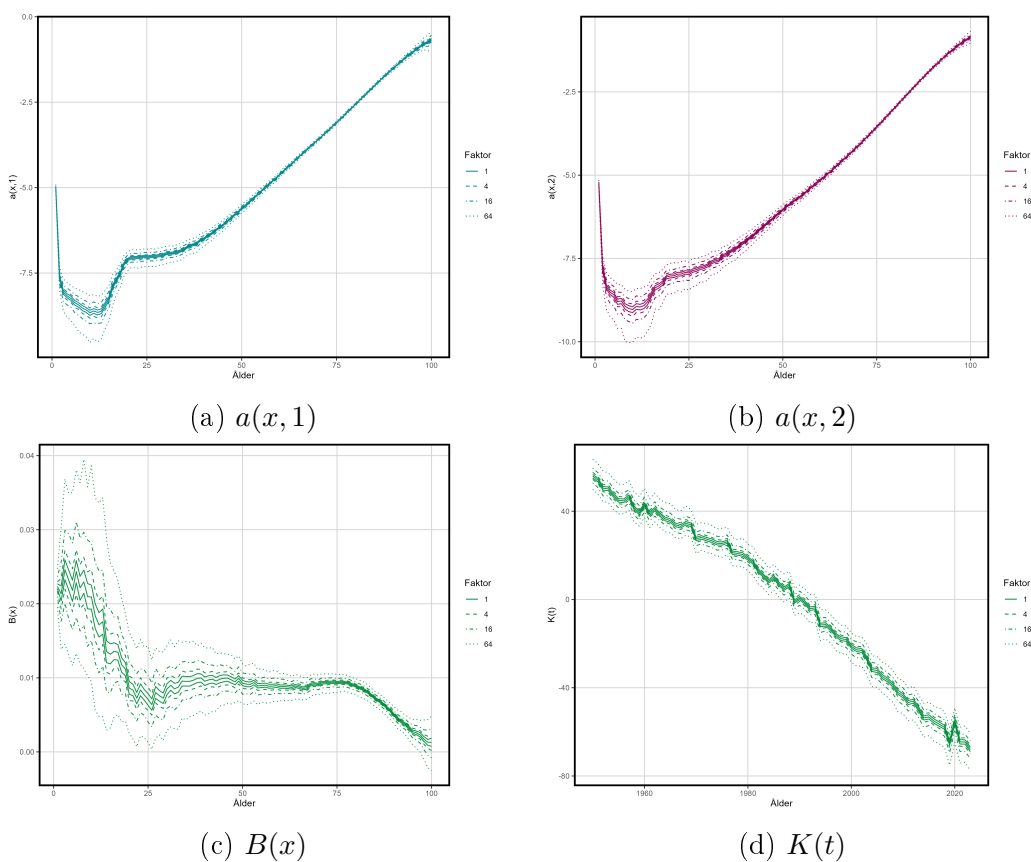
Inom statistik leder generellt sett mindre data till sämre skattningar. Sverige som land har en relativt liten population vilket kan leda till stora skattningsfel, särskilt för unga åldrar där antalet dödsfall är lågt.



Figur 5: Parameterskattningar för olika populationsstorlekar, svensk data för män och kvinnor.

Figur 5 visar skattningen av modellparametrarna för olika simulerade populationsstorlekar. Figur 5a, 5b och 5c visar skattningar för $a(x, 1)$, $a(x, 2)$ respektive

$B(x)$. Dessa parametrar beror på ålder. Här ser vi att en mindre populationsstorlek leder till mer skattningsfel då de hoppar upp och ner från en ålder till nästa. Detta är tydligast för yngre åldrar där antalet dödsfall redan är lågt. För högre åldrar verkar populationsstorleken ha mindre påverkan. Figur 5d visar skattningen för $K(t)$ som beror på tid. Även här medför mindre populationsstorlek större skattningsfel, men här verkar effekten vara konstant över tid. Detta kan troligtvis förklaras med att dödsfall är relativt jämnt utspridda över tid, medan det skiljer sig stort mellan olika åldrar.



Figur 6: Parameterskattningar med 95% konfidensintervall för olika populationsstorlekar, svensk data för män och kvinnor.

Figur 6 visar skattningarna av modellparametrarna på svensk dödlighetsdata för män och kvinnor tillsammans med 95% konfidensintervall från simulerad data med olika populationsstorlekar. I figur 6a och 6b ser vi att skattningsfelen för $a(x, 1)$ och $a(x, 2)$ är större för yngre åldrar, vilket är särskilt tydligt för $f = 64$. För högre åldrar verkar skattningsfelen vara mindre, även för mindre populationsstorlekar. Figur 6c visar skattningen av $B(x)$ där skattningsfelet

är större för yngre åldrar, men fortsatt relativt även för äldre åldrar och skillnaden mellan populationsstorlekarna ser ut att vara markant för alla åldrar. Figur 6d visar skattningen av $K(t)$ där skattningsfelet är konstant över tid och större för mindre populationsstorlekar.

Figur 7 visar skattningarna av modellparametrarna på dödlighetsdata från Sverige, Norge, Finland och Danmark tillsammans med 95% konfidensintervall från simulerad data med olika populationsstorlekar. I figur 7a, 7b, 7c och 7d ser vi att skattningsfelen för $a(x, 1)$, $a(x, 2)$, $a(x, 3)$ och $a(x, 4)$ är större för yngre åldrar, särskilt för mindre populationsstorlekar, vilket liknar det som observerades i figur 6a och 6b. För äldre åldrar verkar skattningsfelen vara små, även för mindre populationsstorlekar. Figur 7e visar skattningen av $B(x)$ där skattningsfelet är större för yngre åldrar, men relativt hög för alla åldrar, särskilt för $f = 64$. Dock är den mindre än den som observerades i figur 6c, vilket troligtvis beror på att den totala population är större. Figur 7f visar skattningen av $K(t)$ där skattningsfelet är konstant över tid och större för mindre populationsstorlekar, vilket liknar det som observerades i figur 6d.

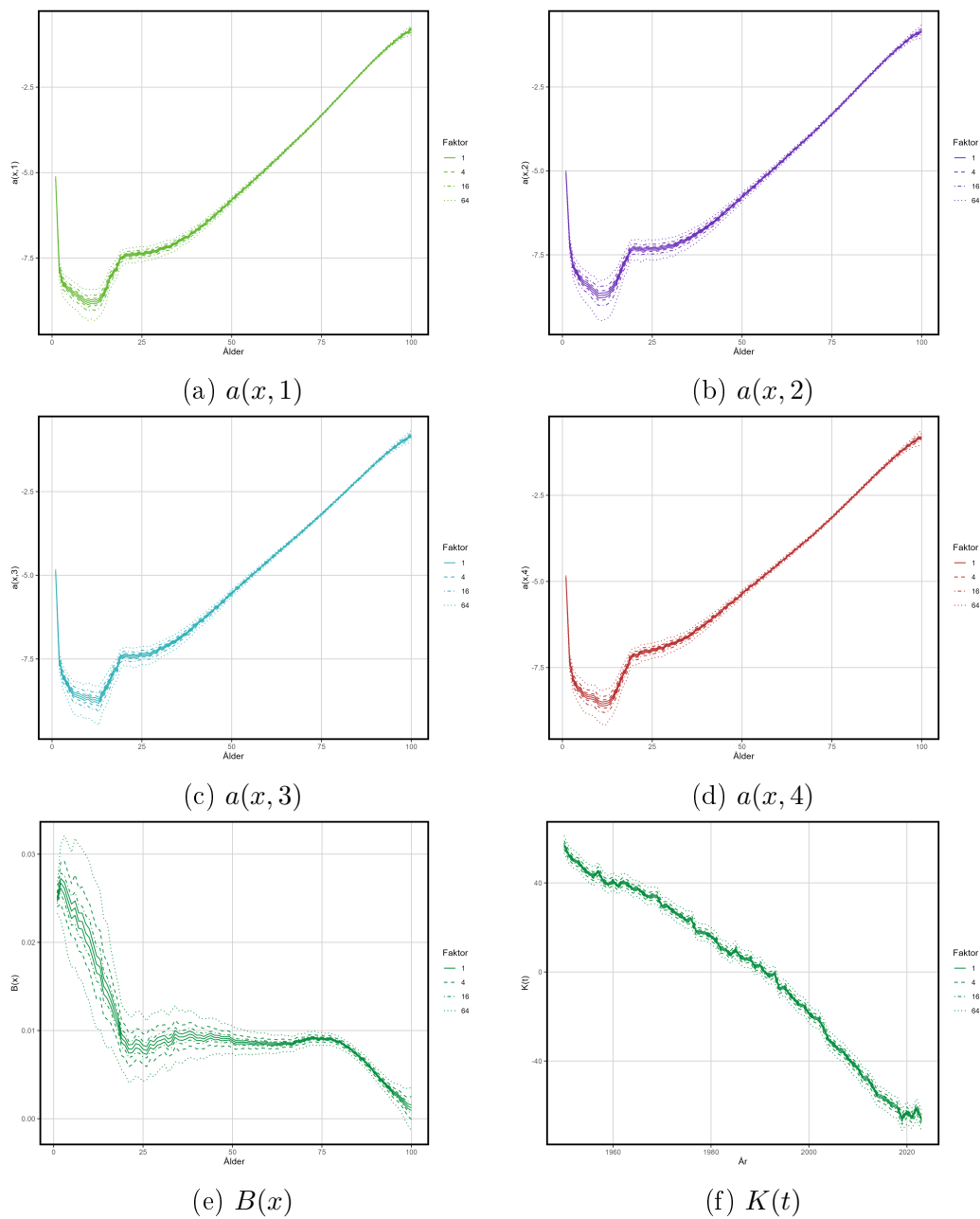
7.2 Antal specifika termer

Till skillnad från Lee-Carter-modellen kräver den gemensamma faktormodellen att vi gör ett aktivt val, nämligen att bestämma antalet delpopulations-specifika parametrar. Någon definitiv metod för att göra detta verkar ej finnas men ett tillvägagångssätt är att titta på modellens prediktiva förmåga.

n	0	1	2	3	4	5	6	7	8
1950-2000	25156	25817	22126	21934	21926	21928	21903	21902	21902
1960-2000	22626	21398	18585	18591	18595	18570	18569	18566	18564
1970-2000	22053	21167	19659	19670	19668	19653	19653	19653	19654
1980-2000	18375	18068	17967	17984	17969	17968	17970	17970	17993
1990-2000	24619	24386	24380	24379	24380	24372	24391	24414	24415

Tabell 3: *Out-of-sample Poissondevians för olika värden av n och olika tidsintervall för modellanpassning, svensk data för män och kvinnor. Fetstilt indikerar den lägsta Poissondeviansen för varje tidsintervall.*

Tabell 3 visar *out-of-sample* Poissondevians för olika antal delpopulations-specifika termer när den gemensamma faktormodellen har anpassats på dödlighetsdata för svenska män och kvinnor från olika tidsintervall. Värdet av n som ger lägst Poissondevians skiljer sig mycket beroende på vilket tidsintervall som användes vid modellanpassningen. När modellen anpassades med



Figur 7: Parameterskattningar för olika populationsstorlekar med 95% konfidensintervall, data från Sverige, Norge, Finland och Danmark.

data från 1950-2000 och 1960-2000 är Poissondeviansen som lägst för de högsta värdena av n som testades och skulle därmed kunna vara ännu lägre för högre värden av n . Dock är det stor skillnad i Poissondeviansens storlek mellan dessa. För anpassningsperioden 1970-2000 är det $n = [5, 6, 7]$ som ger lägst Poissondevians och för 1980-2000 och 1990-2000 är det $n = 2$ respektive $n = 3$ som ger lägst Poissondevians, där den förstnämnda även är den lägsta Poissondeviansen i hela simuleringen.

n	0	1	2	3	4	5	6	7	8
1950-2000	138994	91200	82711	82581	82300	82314	82289	82288	82286
1960-2000	116161	81080	77320	77299	77242	77196	77193	77177	77180
1970-2000	78829	63523	64850	64702	64596	64588	64585	64592	64595
1980-2000	66456	65770	65556	64935	64947	64958	64956	64960	64961
1990-2000	48936	47829	47872	47895	47868	47886	47892	47866	47843

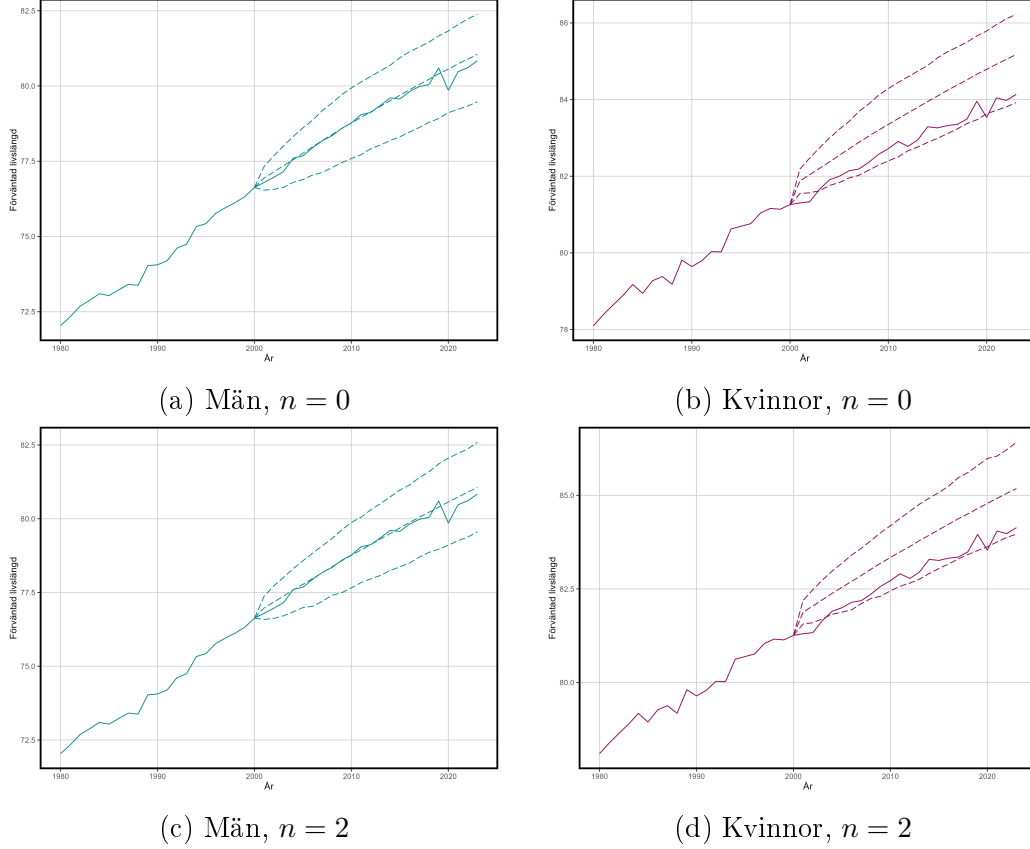
Tabell 4: *Out-of-sample Poissondevians för olika värden av n och olika tidsintervall för modellenanpassning, data från Sverige, Norge, Finland och Danmark. Fetstilt indikerar den lägsta Poissondeviansen för varje tidsintervall.*

Tabell 4 visar *out-of-sample* Poissondevians för olika antal delpopulations-specifika termer när den gemensamma faktormodellen har anpassats på dödlighetsdata från Sverige, Norge, Finland och Danmark från olika tidsintervall. Värdet av n som ger lägst Poissondevians skiljer sig även här mycket beroende på vilket tidsintervall som användes vid modellenanpassningen. När modellen anpassades med data från 1950-2000 och 1960-2000 är Poissondeviansen som lägst för $n = 8$ respektive $n = 7$ och dessa har relativt lika Poissondevians. För anpassningsperioderna 1970-2000 och 1990-2000 är det $n = 1$ som ger lägst Poissondevians och för 1980-2000 är det $n = 4$ som ger lägst Poissondevians. Poissondeviansen är klart lägst för tidsperioden 1990-2000, en skillnad mot tabell 3 där alla tidsperioder hade relativt lika lägsta Poissondevians. Detta kan bero på att befolkningarna i de olika länderna är rätt så olika fram till under lång tid, vilket vi såg i figur 4.

7.3 Prognosernas fel

I en del figurer i detta avsnitt ser prognoserna ut att vara skiftade upp eller ner från den observerade datan. Detta beror på hur prognoserna görs och är ett resultat av att hur $K(t)$ anpassas. När prognoserna görs modelleras $\tilde{K}(t+h)$ som en slumpvandring med drift med väntevärde $(t+h)C$ och då värdet av $\hat{K}(t)$ för det sista året i anpassningsperioden vanligtvis är lite lägre eller högre än tC resulterar det i att prognoserna ser ut att vara skiftade upp

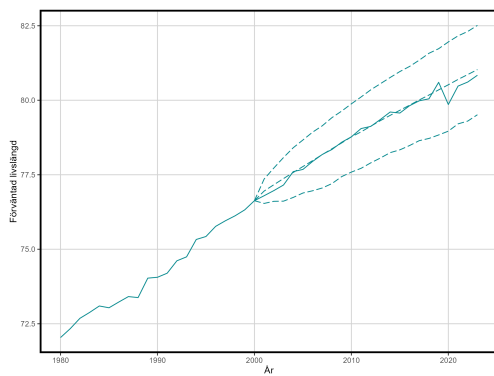
eller ner.



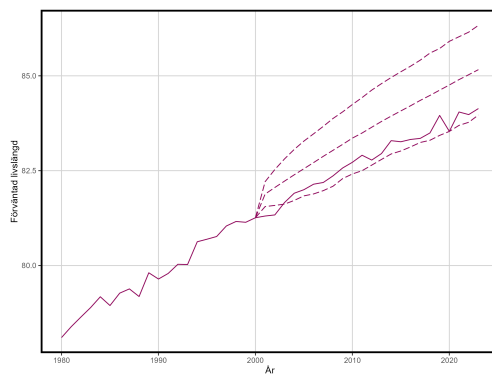
Figur 8: *Förväntad livslängd vid födsel, observerad och prognos med 95% prediktionsintervall, exklusive Poissonvariation.*

Figur 8 visar observerad förväntad livslängd vid födsel för 1980-2023 samt prognoser för åren 2001-2023 med 95% prediktionsintervall för män och kvinnor där gemensam faktormodellen med $n = [0, 2]$ har anpassats till data från 1980-2000. Prognoserna och tillhörande prediktionsintervall har fått hjälp av parametrisk bootstrap vilket beskrivs i första stycket av del 6.3. I figur 8a och 8b ser vi att prediktionsintervallen ser ut att vara proportionerliga mot roten ur tiden sedan prognosens start, vilket är väntat enligt (6). I figur 8c och 8d ser prediktionsintervallen ut att följa ett liknande mönster och det verkar inte vara någon större skillnad i prediktionsfel mellan $n = 0$ och $n = 2$ för varken män eller kvinnor. De observerade värdena är inom prediktionsintervallen hela tidsperioden för både män och kvinnor, både för $n = 0$ och $n = 2$.

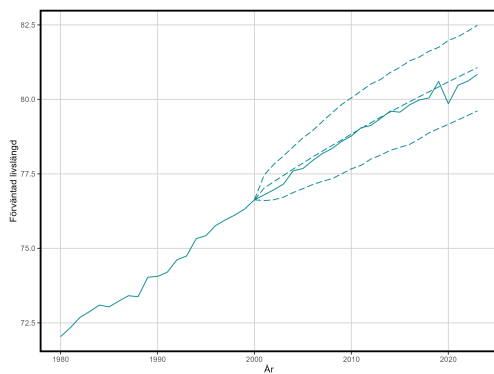
Figur 9 visar observerad förväntad livslängd vid födsel för 1980-2023 samt



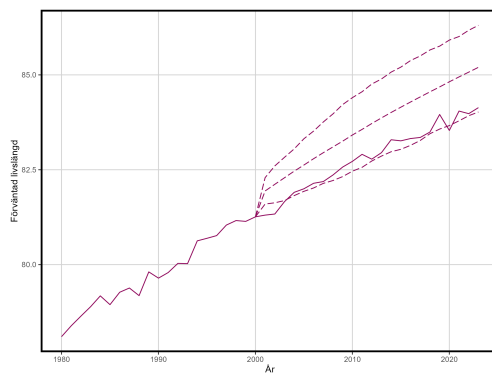
(a) Män, $n = 0$



(b) Kvinnor, $n = 0$



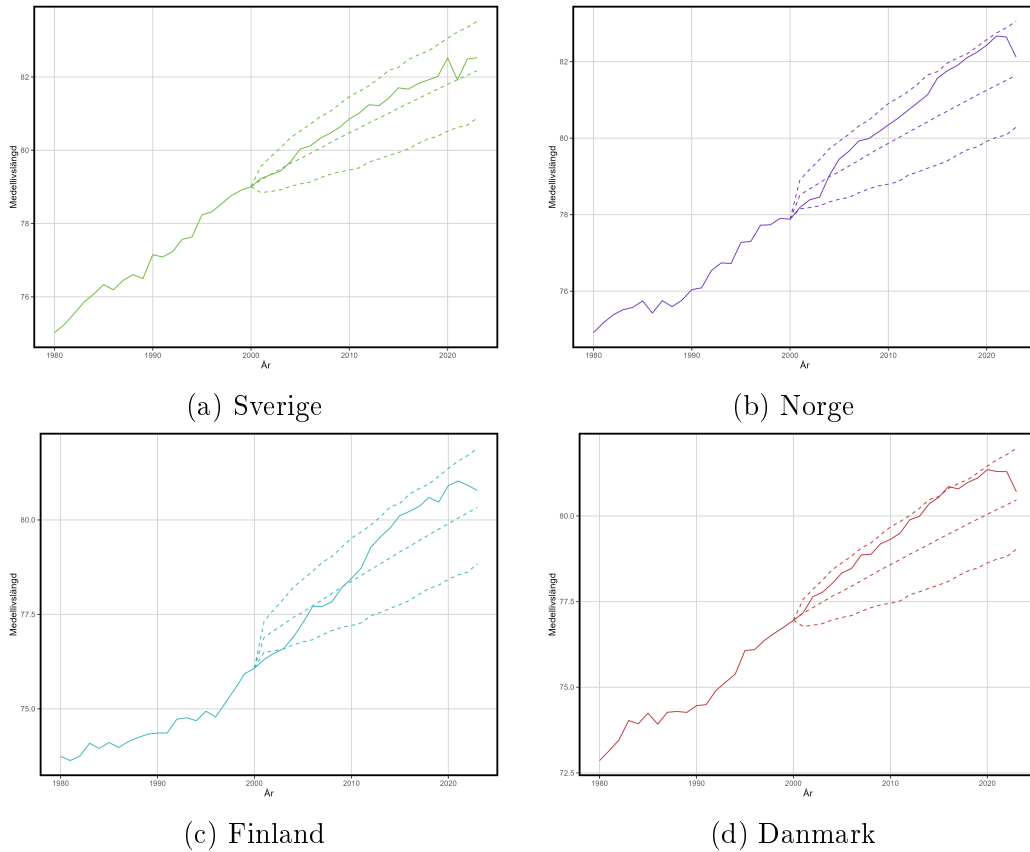
(c) Män, $n = 2$



(d) Kvinnor, $n = 2$

Figur 9: *Förväntad livslängd vid födsel med 95% prediktionsintervall, inklusive Poissonvariation.*

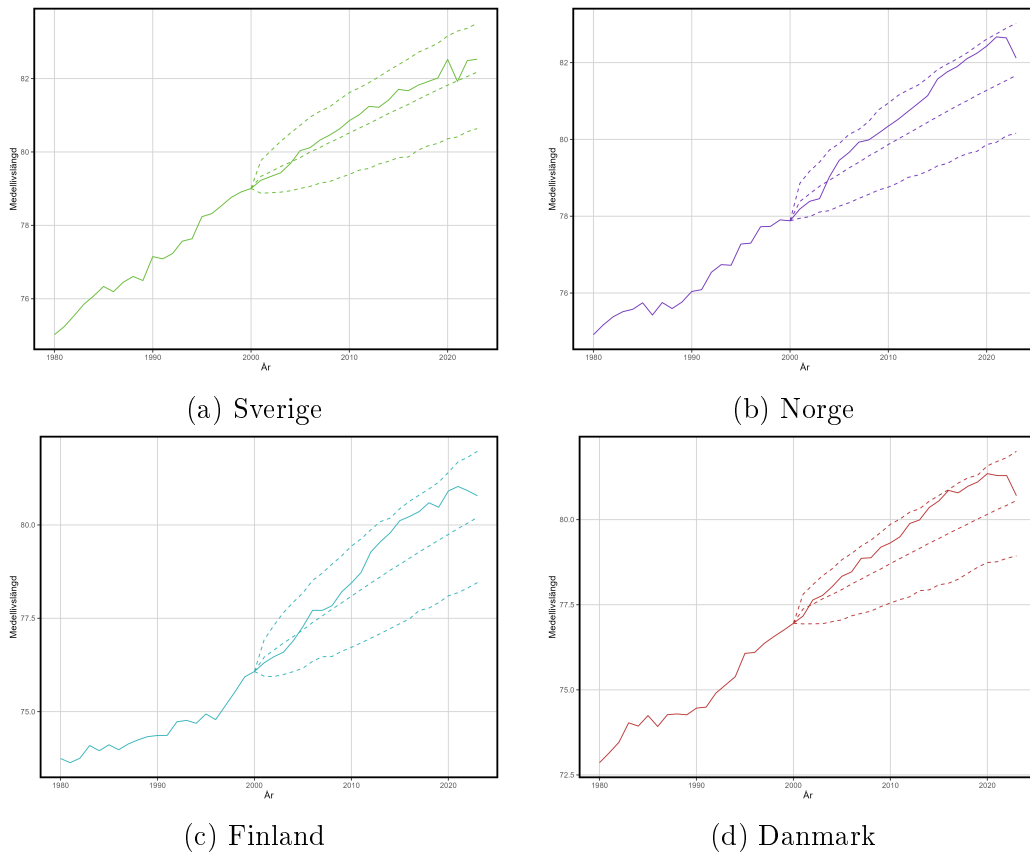
prognoser för åren 2001-2023 med 95% prediktionsintervall för män och kvinnor där gemensam faktormodellen med $n = [0, 2]$ har anpassats till data från 1980-2000. Prognoserna och tillhörande prediktionsintervall har fått hjälp av parametrisk bootstrap vilket beskrivs i andra stycket av del 6.3. I figur 8a och 8b ser vi att prediktionsintervallen ser åter igen ut att vara proportionerliga mot roten ur tiden sedan prognosens start, vilket är väntat enligt (6). I figur 8c och 8d ser prediktionsintervallen ut att följa ett liknande mönster. De observerade värdena är inom prediktionsintervallen hela tidsperioden för både män och kvinnor, både för $n = 0$ och $n = 2$. Till skillnad från figur 8 ser prediktionsintervallen ut att vara bredare för $n = 2$ än för $n = 0$ för både män och kvinnor.



Figur 10: *Förväntad livslängd vid födsel med 95% prediktionsintervall, data från Sverige, Norge, Finland och Danmark, exklusive Poissonvariation.*

Figur 10 visar observerad förväntad livslängd vid födsel för 1980-2023 samt prognoser för åren 2001-2023 med 95% prediktionsintervall för den totala populationen i Sverige, Norge, Finland och Danmark där gemensam faktor-

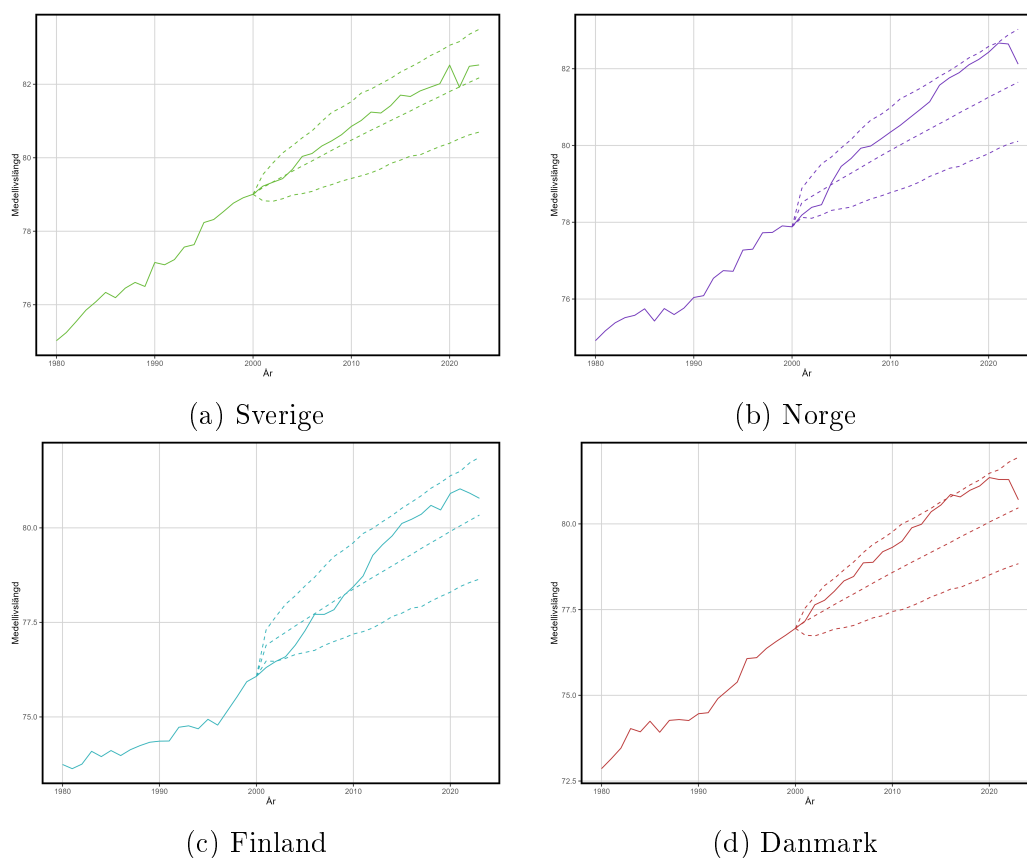
modellen med $n = 0$ har anpassats till data från 1980-2000. Prognoserna och tillhörande prediktionsintervall har fåttts med hjälp av parametrisk bootstrap vilket beskrivs i första stycket av del 6.3. I figur 10a, 10b, 10c och 10d ser vi att prediktionsintervallen ser ut att vara proportionerliga mot roten ur tiden sedan prognosens start, vilket är väntat enligt (6). Den observerade datan för Sverige är bekvämt inom prediktionsintervallet under hela tidsperioden, medan den är närmare gränserna för de andra länderna. Detta kan vara en konsekvens av att Sverige har en större population än de andra länderna. Formen på prediktionsintervallen är liknande för de fyra länderna.



Figur 11: *Förväntad livslängd vid födsel med 95% prediktionsintervall, data från Sverige, Norge, Finland och Danmark, exklusive Poissonvariation.*

Figur 11 visar observerad förväntad livslängd vid födsel för 1980-2023 samt prognoser för åren 2001-2023 med 95% prediktionsintervall för den totala populationen i Sverige, Norge, Finland och Danmark där gemensam faktor-modellen med $n = 2$ har anpassats till data från 1980-2000. Prognoserna och tillhörande prediktionsintervall har fåttts med hjälp av parametrisk bootstrap

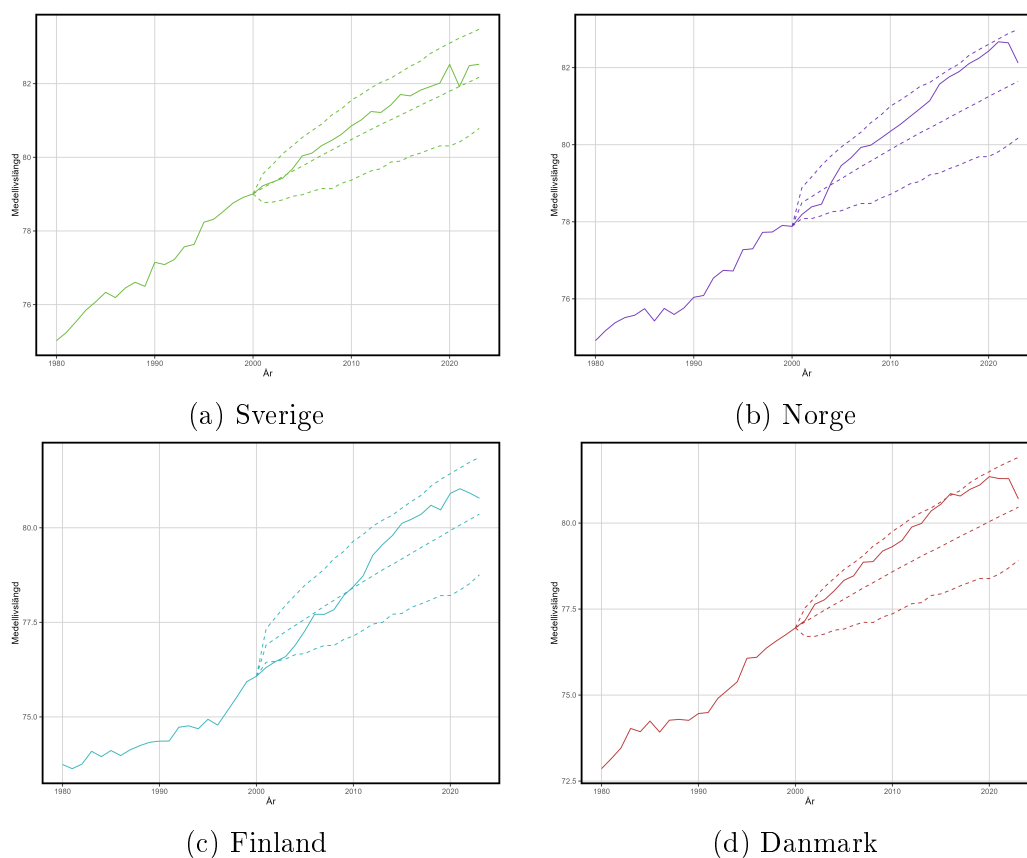
vilket beskrivs i första stycket av del 6.3. I figur 11a, 11b, 11c och 11d ser vi att prediktionsintervallen ser ut att vara proportionerliga mot roten ur tiden sedan prognosens start, vilket är väntat enligt (6). Den observerade datan för Sverige är bekvämt inom prediktionsintervallet under hela tidsperioden, medan den är närmare gränserna för de andra länderna. I stort är figur 11 väldigt lik figur 10 vilket tyder på att det inte är någon större skillnad i prediktionsfel mellan $n = 0$ och $n = 2$ för de olika länderna, likt det som observerades i figur 8.



Figur 12: Förväntad livslängd vid födsel med 95% prediktionsintervall, data från Sverige, Norge, Finland och Danmark, inklusive Poissonvariation.

Figur 12 visar observerad förväntad livslängd vid födsel för 1980-2023 samt prognoser för åren 2001-2023 med 95% prediktionsintervall för den totala populationen i Sverige, Norge, Finland och Danmark där gemensam faktor-modellen med $n = 0$ har anpassats till data från 1980-2000. Prognoserna och tillhörande prediktionsintervall har fåtts med hjälp av parametrisk bootstrap vilket beskrivs i andra stycket av del 6.3. I figur 12a, 12b, 12c och 12d ser

vi att prediktionsintervallen ser ut att vara proportionerliga mot roten ur tiden sedan prognosens start, vilket är väntat enligt (6). Den observerade datan för Sverige är bekvämt inom prediktionsintervallet under hela tidsperioden, medan den är närmare gränserna för de andra länderna. Formen på prediktionsintervallen är liknande för de fyra länderna. Om vi jämför med figur 10 ser vi att prediktionsfelen är ungefär lika stora vilket tyder på att det mesta av prediktionsfelen i prognoserna inte kommer från parameter-skattningarna utan från tidsserierna.



Figur 13: *Förväntad livslängd vid födsel med 95% prediktionsintervall, data från Sverige, Norge, Finland och Danmark, inklusive Poissonvariation.*

Figur 13 visar observerad förväntad livslängd vid födsel för 1980-2023 samt prognoser för åren 2001-2023 med 95% prediktionsintervall för den totala populationen i Sverige, Norge, Finland och Danmark där gemensam faktor-modellen med $n = 2$ har anpassats till data från 1980-2000. Prognoserna och tillhörande prediktionsintervall har fåttts med hjälp av parametrisk bootstrap vilket beskrivs i andra stycket av del 6.3. I figur 13a, 13b, 13c och 13d ser vi

att prediktionsintervallen ser ut att vara proportionerliga mot roten ur tiden sedan prognosens start, vilket är väntat enligt (6). Den observerade datan för Sverige är bekvämt inom prediktionsintervallet under hela tidsperioden, medan den är närmare gränserna för de andra länderna. I stort är figur 11 väldigt lik figur 10 vilket tyder på att det inte är någon större skillnad i prediktionsfel mellan $n = 0$ och $n = 2$ för de olika länderna, likt det som observerades i figur 8. Om vi jämför med figur 11 ser vi att prediktionsfelen är ungefär lika stora vilket tyder på att det mesta av prediktionsfelen i prognoserna inte kommer från parameterskattningarna utan från tidsserierna. Vi kan också jämföra med figur 12 där vi ser att prediktionsfelen är ungefär lika stora, vilket åter igen tyder på att det mesta av prediktionsfelen i prognoserna inte kommer från parameterskattningarna utan från tidsserierna.

8 Diskussion

Syftet med detta arbete var att utreda den gemensamma faktormodellen (GF-modellen). Som i alla statistiska modeller är datamängden avgörande för säkerheten i parameterskattningarna och så är även fallet för GF-modellen. Detta syns tydligt i figur 6 där konfidensintervallen är bredare för alla parametrar för mindre populationsstorlekar vid lägre åldrar, vilket beror på det låga antalet dödsfall i dessa åldrar. Den största påverkan syns i figur 6c där konfidensintervallet för $f = 64$ är brett för alla åldrar, men särskilt för yngre. Om vi jämför resultaten i figur 6a-6b med de i figur 7a-7d ser vi att skattningsfelen ser ut att vara jämförbara. Detta är väntat då dessa parametrar skattas individuellt för varje delpopulation och antalet invånare i de olika länderna och antalet svenska män och kvinnor är relativt lika, förutom Sveriges totala antal invånare som är ungefär dubbelt så stort som resten av populationerna men som ej har resulterat i en märkbar skillnad. Om vi istället jämför figur 6c och figur 6d med figur 7e och 7f ser vi att skattningsfelen är mindre för Sverige, Norge, Finland och Danmark än för svenska män och kvinnor. Detta beror på att dessa parametrar skattas gemensamt och den totala populationen är större för de olika länderna än för svenska män och kvinnor. GF-modellen kan därmed vara ett rimligt val för modellering av länder med små populationer då de kan modelleras gemensamt med andra länder, vilket minskar skattningsfelen av de gemensamma termerna. Dock är det viktigt att notera att olika delpopulationer som modelleras gemensamt bör ha liknande dödlighet, men det är svårt att avgöra. Män och kvinnor i samma land eller olika länder i ett geografiskt område känns som naturliga val för GF-modellen, men mer formella metoder för att avgöra om två eller flera populationer är tillräckligt lika för att modelleras gemensamt verkar saknas.

Att titta på om sådana metoder är rimliga och i så fall hur de kan se ut är en intressant framtida uppgift.

En nackdel med GF-modellen är att den kräver att vi gör ett aktivt val av hur många delpopulationsspecifika termer som ska användas. Li [4] säger att överparametrisering av modellen ska undvikas, men att tillräckligt många delpopulationsspecifika termer används för att fånga datans huvudsakliga attribut. I detta arbete har vi undersökt hur antalet av dessa termer kan väljas om vi bara tittar på GF-modellens prediktiva förmåga. I tabell 3 ser vi att antalet delpopulationsspecifika termer som ger lägst Poissondevians varierar mellan 2 och 8 beroende på vilket tidsintervall som används vid modellanpassning. En tumregel är att använda som antal år för modellanpassning som för prognosens längd och det tidsintervall som är närmast detta är 1980-2000 som faktiskt ger lägst Poissondevians för $n = 2$, vilket indikerar att lägre värden av n är att föredra. Däremot är Poissondeviansen inte särskilt mycket större när GF-modellen anpassas på data från 1960-2000 med $n = 8$ vilket tyder på motsatsen. Det som är sant för alla tidsintervall är att om Poissondeviansen minskar efter $n = 2$ eller $n = 3$ så är det endast marginellt, och då vi vill undvika överanpassning är troligtvis lägre värden av n att föredra. Något liknande kan ses i tabell 4 där värdet av n som ger lägst Poissondevians varierar mellan 1 och 8 beroende på vilket tidsintervall som används vid modellanpassning. Den lägsta Poissondeviansen fås för $n = 1$ och tidsintervallet 1990-2000, men till skillnad från tabell 3 är detta betydligt lägre än det lägsta värdet för de andra tidsintervallen. I tabell 4 är Poissondeviansen för de olika tidsintervallen relativt jämt för $n \geq 3$. Utifrån dessa resultat verkar det som att $n = [1, 2, 3]$ är rimliga värden som ger bra prediktiv förmåga utan att modellen överparametreras, men detta kan såklart se annorlunda ut för annan data. Det vore intressant att undersöka hur värdet av n kan bestämmas utifrån annat än prediktiv förmåga, exempelvis genom att titta på Poissonresidualer, eller genom att kombinera flera mått för att avgöra det optimala värdet av n .

En av de viktigaste tillämpningarna av dödlighetsmodeller är att göra prognoser framåt i tiden, och dessa prognoser innehåller prediktionsfel. I och med hur prognoserna görs är det flera faktorer som bidrar till prediktionsfelet och detta arbete har undersökt två av dessa, nämligen fel som är en följd av tidsserier och Poissonvariation. Antalet delpopulationsspecifika termers påverkan på prediktionsfelet har också undersökts då fler sådana termer resulterar i fler tidsserier som bidrar till prognoserna. Figur 8 visar 95% prediktionsintervall för $n = 0$ och $n = 2$ för svenska män och kvinnor när Poissonvariation är exkluderad och om vi jämför figur 8a och 8b med figur 8c och 8d vi ser att prediktionsintervallen möjligtvis är något bredare för $n = 2$

än för $n = 0$ för både män och kvinnor. Detta är dock inte helt tydligt och skulle kunna se annorlunda ut för en större simulering. Figur 9 visar 95% prediktionsintervall för $n = 0$ och $n = 2$ för svenska män och kvinnor när Poissonvariation är inkluderad och om vi jämför figur 9a och 9b med figur 9c och 9d ser vi återigen att prediktionsintervallen är något bredare för $n = 2$ än för $n = 0$. Detta indikerar att fler delpopulationsspecifika termer ger större prediktionsfel. Om vi jämför figur 8a-8d med figur 9a-9d är det svårt att se någon tydlig skillnad i prediktionsfelen mellan de två figurerna, vilket tyder på att Poissonvariation har en relativt liten påverkan på prediktionsfelen. Detta kan se annorlunda ut för mindre populationer då vi vet från resultaten av simuleringen beskriven i del 6.1 att populationsstorlek har en stor påverkan på prediktionsfel orsakad av Poissonvariation. Dessa skillnader, om de finns, är ännu mindre tydliga i simuleringen baserad på data från Sverige, Norge, Finland och Danmark. Om vi jämför figur 10a-10d med figur 11a-11d syns ingen tydlig skillnad i prediktionsintervallens bredd mellan $n = 0$ och $n = 2$ för något land, vilket tyder på att det i detta fall inte är någon skillnad i prediktionsfel mellan $n = 0$ och $n = 2$ för något land. Samma sak gäller för figur 12a-12d och figur 13a-13d där Poissonvariation är inkluderad. Om vi jämför figur 10a-10d med figur 12a-12d ser vi att det inte är någon märkbar skillnad i prediktionsfel mellan de två figurerna, vilket tyder på att Poissonvariation återigen har en relativt liten påverkan på prediktionsfelen. Samma sak gäller för figur 11a-11d och figur 13a-13d där Poissonvariation är inkluderad. Skillnaden mellan resultaten från simuleringen baserad på svensk data och simuleringen baserad på data från Sverige, Norge, Finland och Danmark kan bero på att den totala populationen är större för de olika länderna tillsammans än för Sverige, vilket skulle kunna leda till att Poissonvariation är mindre för de olika länderna tillsammans än för Sverige. Dock är det svårt att säga med säkerhet då prediktionsintervallen endast bygger på 1000 prognoser och det är möjligt att resultaten skulle se annorlunda ut med fler prognoser. Det vore intressant att undersöka hur prognosernas prediktionsintervall hade påverkats av större och mindre populationsstorlekar och om en större skillnad hade i prediktionsintervallens bredd hade synts för större värden av n . Då den störst bidragande faktorn till prediktionsfelen är tidsserierna vore det även intressant att se hur detta hade påverkats om vi valt att använda andra tidsserier än de som använts i detta arbete.

9 Slutsats

Den gemensamma faktormodellen är en naturlig utökning av Lee-Cartermodellen för gemensam modellering av likartade populationer. Den inkluderar en gemen-

sam faktor som skattas gemensamt för alla populationer och ett godtyckligt antal delpopulationsspecifika termer. I detta arbete har vi undersökt skattningsfel orsakad av Poissonvariation när populationsstorleken varierar när vi använder svensk data för män och kvinnor och data för den totala populationen i Sverige, Norge, Finland och Danmark. Skattningsfelen ökar i form av bredare konfidensintervall för de olika parametrarna när populationsstorleken minskar. För alla populationsstorlekar är det tydligt att konfidensintervallens bredd beror till stor del på hur mycket data som finns tillgängligt. De parametrar som beror på ålder har bredare konfidensintervall vid lägre åldrar, vilket beror på att det är färre dödsfall i dessa åldrar, medan skattningsfelen för de parametrar som beror på tid är jämna över tid. Skattningsfelen för de parametrar som skattas gemensamt för alla delpopulationer beror på populationsstorleken i den totala populationen, medan skattningsfelen för de parametrar som skattas individuellt för varje delpopulation beror på populationsstorleken i respektive delpopulation. Detta innebär att den gemensamma faktormodellen kan användas på länder med små populationer som kan modelleras gemensamt med andra länder, vilket minskar skattningsfelen för de gemensamma termerna.

Vi har också undersökt hur många delpopulationsspecifika termer som är rimliga att använda när vi gör prognoser framåt i tiden. Också här användes svensk data för män och kvinnor och data för den totala populationen i Sverige, Norge, Finland och Danmark. Det antal delpopulationsspecifika termer som ger lägst Poissondevians varierar starkt beroende på vilket tidsintervall som används vid modellanpassning och det skulle troligtvis se annorlunda ut för annan data. Utifrån resultaten i detta arbete verkar det som att $n = [1, 2, 3]$ är rimliga värden som ger bra prediktiv förmåga samtidigt som vi undviker att modellen överparametreras.

Till sist har vi undersökt hur prediktionsfelet i prognoserna påverkas av tidsserier och Poissonvariation när vi använder svensk data för män och kvinnor och data för den totala populationen i Sverige, Norge, Finland och Danmark.. Resultaten visar att det är svårt att se någon tydlig skillnad i prediktionsfel mellan $n = 0$ och $n = 2$ för någon delpopulation. Skillnaden kan möjligtvis bli större för $n > 2$ men då varje ytterligare delpopulationsspecifik term bidrar mindre och mindre till den skattade dödlighetsintensiteten är det inte säkert. Det går inte heller att se någon tydlig skillnad i prediktionsfel för de olika delpopulationerna med eller utan Poissonvariation vilket tyder på att Poissonvariation bidrar relativt lite till prediktionsfelet.

Referenser

- [1] Aalen, O., O. Borgan och H. Gjessing (2008). *Survival and Event History Analysis: A Process Point of View*. Springer.
- [2] Brouhns, N., M. Denuit och J. Vermunt (2002). "A Poisson log-bilinear regression approach to the construction of projected lifetables". I: *Insurance: Mathematics and economics* 31(3), 373–393.
- [3] Lee, R. och L. Carter (1992). "Modeling and forecasting US mortality". I: *Journal of the American statistical association* 87(419), 659–671.
- [4] Li, J. (2013). "A Poisson common factor model for projecting mortality and life expectancy jointly for females and males". I: *Population studies* 67(1), 111–126.
- [5] Li, N. och R. Lee (2005). "Coherent mortality forecasts for a group of populations: An extension of the Lee-Carter method". I: *Demography* 42(3), 575–594.
- [6] Lo, E., D. Vatnik, A. Benedetti och R. Bourbeau (2016). "Variance models of the last age interval and their impact on life expectancy at subnational scales". I: *Demographic Research* 35, 399–454.
- [7] Statistiska Centralbyrån (u.å.). *Dödlighetsdata från Statistiska Centralbyrån*. Data hämtad via Human Mortality Database, www.mortality.org, den 2025-03-17.
- [8] Statistiska centralbyrån (2025). *Befolkningens livslängd (medellivslängd) – Statistikdatabasen*. Hämtad 5 maj 2025. URL: https://www.statistikdatabasen.scb.se/pxweb/sv/ssd/START__BE__BE0401__BE0401A/BefProgLivslangdNb/table/tableViewLayout1/.

Bilagor

Bilaga 1 - Förväntad livslängd

Förväntad livslängd vid födsel har räknats ut med (7).

År	Förväntad livslängd
1920	58.37
2020	81.92

Tabell 5: *Förväntad livslängd vid födsel, Sveriges totala population.*

Tabell 5 visar den förväntade livslängden vid födsel för Sveriges totala population för åren 1920 och 2020. Den förväntade livslängden vid födsel har ökat med över 23 år under denna tidsperiod.

Bilaga 2 - Modellkonvergens

Alla tabeller i detta avsnitt visar värden som är avrundade nedåt till närmaste heltal.

Anpassning med svensk data för män och kvinnor

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	89159	22366	17583	17504	17499	17498	17498	17498	17498	17498	17498
$n = 1$	17498	7253	7197	7196	7196	7196	7196	7196	7196	7196	7196
$n = 2$	7196	6101	5767	5746	5744	5743	5743	5743	5743	5743	5743
$n = 3$	5743	5497	5421	5391	5365	5340	5324	5316	5312	5311	5310
$n = 4$	5310	5159	5083	5054	5036	5022	5013	5006	5002	5000	4999
$n = 5$	4999	4860	4793	4770	4756	4748	4743	4739	4737	4734	4732
$n = 6$	4732	4617	4557	4532	4512	4495	4482	4473	4467	4463	4460
$n = 7$	4460	4355	4300	4278	4260	4246	4236	4230	4226	4223	4221
$n = 8$	4221	4127	4075	4055	4041	4031	4023	4018	4015	4012	4011

Tabell 6: *Poissondevians efter varje iteration för olika värden av n , svensk data för män och kvinnor från 1950-2000.*

Tabell 6 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för svenska män och kvinnor från 1950-2000. För $n = 0$ konvergerar Poissondeviansen redan efter 5 iterationer. För $n = 1$ och $n = 2$ har ser Poissondeviansen ut att ha

konvergerat efter 10 iterationer, men för högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	54222	11575	9682	9663	9662	9662	9662	9662	9662	9662	9662
$n = 1$	9662	5376	5348	5346	5346	5345	5345	5345	5345	5345	5345
$n = 2$	5345	4429	4273	4257	4253	4252	4252	4252	4252	4252	4252
$n = 3$	4252	4085	4037	4018	4005	3994	3985	3976	3969	3963	3959
$n = 4$	3959	3802	3751	3731	3719	3711	3704	3700	3696	3694	3692
$n = 5$	3692	3565	3511	3486	3471	3460	3451	3444	3438	3435	3432
$n = 6$	3432	3320	3273	3251	3237	3229	3224	3220	3218	3217	3216
$n = 7$	3216	3122	3083	3067	3056	3050	3045	3042	3039	3036	3033
$n = 8$	3033	2949	2911	2893	2879	2869	2862	2857	2854	2852	2850

Tabell 7: *Poissondevians efter varje iteration för olika värden av n , svensk data för män och kvinnor från 1960-2000.*

Tabell 7 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för svenska män och kvinnor från 1960-2000. För $n = 0$ konvergerar Poissondeviansen redan efter 4 iterationer. För $n = 1$ och $n = 2$ har ser Poissondeviansen ut att ha konvergerat efter 10 iterationer, men för högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	28351	4987	4428	4426	4426	4426	4426	4426	4426	4426	4426
$n = 1$	4426	3604	3569	3565	3564	3563	3563	3563	3563	3563	3563
$n = 2$	3563	3144	3052	3036	3032	3031	3031	3031	3031	3031	3031
$n = 3$	3031	2850	2803	2786	2779	2776	2774	2773	2772	2772	2771
$n = 4$	2771	2656	2628	2614	2606	2599	2593	2585	2576	2567	2558
$n = 5$	2558	2448	2412	2392	2382	2375	2372	2369	2367	2365	2364
$n = 6$	2364	2268	2240	2224	2213	2205	2200	2196	2193	2191	2189
$n = 7$	2189	2101	2074	2054	2036	2021	2011	2003	1999	1995	1994
$n = 8$	1994	1909	1883	1867	1852	1841	1835	1832	1830	1828	1828

Tabell 8: *Poissondevians efter varje iteration för olika värden av n , svensk data för män och kvinnor från 1970-2000.*

Tabell 8 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för svenska män och kvinnor från 1970-2000. För $n = 0$ konvergerar Poissondeviansen redan

efter 3 iterationer. För $n = 1$ och $n = 2$ har ser Poissondeviansen ut att ha konvergerat efter 10 iterationer, men för högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	11997	2652	2500	2499	2499	2499	2499	2499	2499	2499	2499
$n = 1$	2499	2306	2241	2206	2179	2158	2145	2139	2136	2134	2133
$n = 2$	2133	1964	1934	1909	1893	1886	1883	1882	1881	1881	1881
$n = 3$	1881	1749	1728	1715	1705	1696	1688	1684	1681	1679	1679
$n = 4$	1679	1567	1548	1536	1529	1525	1522	1520	1519	1518	1518
$n = 5$	1518	1421	1407	1396	1387	1380	1373	1367	1362	1359	1356
$n = 6$	1356	1265	1252	1238	1223	1212	1204	1199	1196	1194	1192
$n = 7$	1192	1104	1094	1086	1078	1069	1060	1054	1050	1048	1047
$n = 8$	1047	964	957	951	946	941	936	932	928	925	923

Tabell 9: *Poissondevians efter varje iteration för olika värden av n , svensk data för män och kvinnor från 1980-2000.*

Tabell 9 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för svenska män och kvinnor från 1980-2000. För $n = 0$ konvergerar Poissondeviansen redan efter 3 iterationer. För högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	2935	1120	1094	1094	1094	1094	1094	1094	1094	1094	1094
$n = 1$	1094	948	913	898	891	889	888	887	887	887	887
$n = 2$	887	779	763	757	753	750	748	745	742	738	733
$n = 3$	733	636	621	608	600	595	593	592	591	591	590
$n = 4$	590	509	499	492	488	484	481	479	476	474	472
$n = 5$	472	397	386	378	373	369	367	365	364	363	362
$n = 6$	362	295	287	282	280	278	277	277	276	276	276
$n = 7$	276	220	215	211	209	207	205	204	203	203	202
$n = 8$	202	150	146	143	139	135	130	127	125	124	124

Tabell 10: *Poissondevians efter varje iteration för olika värden av n , svensk data för män och kvinnor från 1990-2000.*

Tabell 10 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för svenska män och kvinnor från 1990-2000. För $n = 0$ konvergerar Poissondeviansen

redan efter 2 iterationer. För $n = 1$ och ser Poissondeviansen ut att ha konvergerat efter 10 iterationer, men för högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

Anpassning med data från Sverige, Norge, Finland och Danmark

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	213174	51428	37440	36990	36939	36933	36932	36932	36932	36932	36932
$n = 1$	36932	18084	17298	17233	17227	17226	17226	17226	17226	17226	17226
$n = 2$	17226	13656	12860	12712	12689	12684	12683	12683	12683	12683	12683
$n = 3$	12683	11676	11271	11205	11182	11154	11118	11087	11067	11056	11051
$n = 4$	11051	10585	10282	10173	10121	10090	10073	10065	10062	10060	10060
$n = 5$	10060	9742	9591	9534	9509	9495	9486	9480	9477	9474	9473
$n = 6$	9473	9214	9103	9055	9029	9013	9000	8989	8980	8973	8965
$n = 7$	8965	8745	8655	8608	8573	8547	8529	8518	8510	8504	8499
$n = 8$	8499	8294	8210	8165	8130	8103	8083	8068	8057	8049	8043

Tabell 11: *Poissondevians efter varje iteration för olika värden av n , data från Sverige, Norge, Finland och Danmark från 1950-2000.*

Tabell 11 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för från Sverige, Norge, Finland och Danmark från 1950-2000. För $n = 0$ konvergerar Poissondeviansen redan efter 6 iterationer. För $n = 1$ och $n = 2$ har ser Poissondeviansen ut att ha konvergerat efter 10 iterationer, men för högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

Tabell 12 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för från Sverige, Norge, Finland och Danmark från 1960-2000. För $n = 0$ konvergerar Poissondeviansen redan efter 5 iterationer. För $n = 1$ och $n = 2$ har ser Poissondeviansen ut att ha konvergerat efter 10 iterationer, men för högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

Tabell 13 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för från Sverige, Norge, Finland och Danmark från 1970-2000. För $n = 0$ konvergerar Poissondeviansen redan efter 3 iterationer. För högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	132213	30873	25479	25406	25403	25402	25402	25402	25402	25402	25402
$n = 1$	25402	12664	12108	12064	12060	12060	12059	12059	12059	12059	12059
$n = 2$	12059	9159	9013	9010	9009	9010	9010	9010	9010	9010	9010
$n = 3$	9010	8498	8318	8254	8204	8152	8108	8082	8069	8063	8060
$n = 4$	8060	7698	7543	7485	7443	7412	7394	7383	7374	7368	7364
$n = 5$	7364	7094	6972	6937	6924	6916	6912	6909	6907	6905	6904
$n = 6$	6904	6712	6616	6566	6534	6513	6499	6490	6482	6477	6473
$n = 7$	6473	6297	6216	6173	6140	6116	6097	6083	6073	6065	6060
$n = 8$	6060	5890	5816	5775	5744	5721	5705	5694	5687	5682	5678

Tabell 12: *Poissondevians efter varje iteration för olika värden av n , data från Sverige, Norge, Finland och Danmark från 1960-2000.*

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	64965	16489	15076	15071	15071	15071	15071	15071	15071	15071	15071
$n = 1$	15071	8030	7904	7892	7884	7873	7854	7829	7806	7793	7788
$n = 2$	7788	6643	6396	6362	6357	6356	6356	6356	6356	6356	6356
$n = 3$	6356	5999	5855	5804	5776	5763	5758	5755	5753	5750	5746
$n = 4$	5746	5492	5359	5304	5274	5258	5248	5243	5240	5238	5237
$n = 5$	5237	5019	4923	4892	4878	4869	4864	4860	4857	4854	4852
$n = 6$	4852	4676	4585	4538	4513	4498	4488	4482	4477	4473	4470
$n = 7$	4470	4311	4234	4197	4179	4166	4154	4144	4136	4128	4122
$n = 8$	4122	3982	3910	3867	3843	3829	3820	3814	3809	3805	3803

Tabell 13: *Poissondevians efter varje iteration för olika värden av n , data från Sverige, Norge, Finland och Danmark från 1970-2000.*

Tabell 14 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för från Sverige, Norge, Finland och Danmark från 1980-2000. För $n = 0$ konvergerar Poissondeviansen redan efter 2 iterationer. För $n = 1$ har ser Poissondeviansen ut att ha konvergerat efter 10 iterationer, men för högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

Tabell 15 visar Poissondeviansen efter varje iteration för olika värden av n när gemensam faktormodellen har anpassats på dödlighetsdata för från Sverige, Norge, Finland och Danmark från 1990-2000. För $n = 0$ konvergerar Poissondeviansen redan efter 2 iterationer. För högre värden av n ser det ut som att Poissondeviansen inte har konvergerat helt efter 10 iterationer.

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	24454	7514	7245	7245	7245	7245	7245	7245	7245	7245	7245
$n = 1$	7245	4955	4828	4808	4803	4801	4800	4800	4799	4799	4799
$n = 2$	4799	4301	4178	4132	4117	4111	4107	4103	4100	4097	4095
$n = 3$	4095	3754	3677	3646	3633	3626	3623	3620	3619	3618	3618
$n = 4$	3618	3367	3330	3310	3297	3287	3280	3276	3273	3271	3269
$n = 5$	3269	3060	3024	3003	2989	2976	2964	2952	2940	2931	2924
$n = 6$	2924	2725	2689	2668	2654	2644	2636	2630	2626	2622	2620
$n = 7$	2620	2445	2413	2393	2379	2366	2354	2342	2331	2323	2317
$n = 8$	2317	2150	2114	2088	2069	2056	2048	2043	2039	2036	2033

Tabell 14: *Poissondevians efter varje iteration för olika värden av n , data från Sverige, Norge, Finland och Danmark från 1980-2000.*

<i>Iteration</i>	0	1	2	3	4	5	6	7	8	9	10
$n = 0$	7030	2855	2813	2813	2813	2813	2813	2813	2813	2813	2813
$n = 1$	2813	2228	2201	2191	2185	2179	2174	2170	2166	2164	2163
$n = 2$	2163	1845	1806	1786	1775	1766	1756	1745	1737	1732	1730
$n = 3$	1730	1484	1472	1463	1457	1452	1448	1445	1442	1440	1438
$n = 4$	1438	1212	1195	1185	1178	1173	1168	1164	1161	1158	1155
$n = 5$	1155	964	946	932	922	915	910	905	901	897	895
$n = 6$	895	729	715	706	697	691	686	683	681	679	678
$n = 7$	678	529	521	516	511	504	497	491	486	483	480
$n = 8$	480	341	325	322	319	316	313	310	307	305	303

Tabell 15: *Poissondevians efter varje iteration för olika värden av n , data från Sverige, Norge, Finland och Danmark från 1990-2000.*