

Statistiska metoder för härledning av indata till säkerhetsanalyser inom kärnkraftsområdet

Vidar Hedtjärn Swaling



UPPSALA
UNIVERSITET

Teknisk- naturvetenskaplig fakultet
UTH-enheten

Besöksadress:
Ångströmlaboratoriet
Lägerhyddsvägen 1
Hus 4, Plan 0

Postadress:
Box 536
751 21 Uppsala

Telefon:
018 – 471 30 03

Telefax:
018 – 471 30 00

Hemsida:
<http://www.teknat.uu.se/student>

Abstract

Statistiska metoder för härledning av indata till
säkerhetsanalyser inom kärnkraftsområdet

Statistical methods for deriving input to nuclear power safety assessments

Vidar Hedtjörn Swaling

Components in Swedish nuclear power plants are subject to elaborate testing and maintenance, all in accordance with the rigorous safety requirements of the Swedish Nuclear Power Inspectorate (SKI). Nevertheless, failures sometimes occur that hypothetically might lead to severe accidents. Information on such critical failures is collected and statistical measures of the rate of occurrence are computed.

Today, these computations are carried out using a two-stage Bayesian method developed by Kurt Pörn. Pörn's method is supposed to be particularly appropriate when data is extremely sparse. One objection is that it is conceptually non-transparent and therefore should be replaced by a simpler alternative, if there is one.

On the basis of an overview given in this thesis, a method proposed by Jussi Vaurio seems to be the most interesting alternative. Its main advantages are fast calculations, transparent mathematical expressions and straightforward replication. Moreover, the numerical results presented here are quite promising. Hence, though further comparisons are necessary, Vaurio's method is proposed as a simpler and in some respects even better alternative to Pörn's method.

Handledare: Anna Gabrielsson
Ämnesgranskare: Sven Erick Alm
Examinator: Elísabet Andrésdóttir
ISSN: 1650-8319, UPTec STS06 014

Populärvetenskaplig sammanfattning

Verksamheten på de svenska kärnkraftverken präglas av rigorösa säkerhetsrutiner. Dessa omfattar inte minst tester och underhåll av kärnkraftverkets många komponenter och delsystem, dvs. pumpar, ventiler, reservaggregat etc. Driftsäkerheten hos de enskilda komponenterna är generellt sett extremt hög. Ändå inträffar emellanåt fel som i förlängningen anses kunna leda till allvarliga olyckor. Information om sådana *kritiska* fel sammanställs i en databas och ligger till grund för framställning av olika mått på *hur ofta de kritiska felen inträffar*. Sådan ”felstatistik” används i sin tur som indata till analyser där motsvarande statistiska mått för hela kärnkraftverket beräknas.

Den statistiska bearbetningen av rådata (dvs. feldata på komponentnivå) har sedan början av 90-talet gjorts med en metod utvecklad av Kurt Pörn. Pörns metod anses vara bra för att beräkna sannolikheter för händelser som inträffar mycket sällan, eller som *aldrig* har inträffat. En nackdel är att metoden är matematiskt avancerad och därmed svårgenomskådlig. Därför finns också ett intresse av att söka efter enklare alternativ.

Syftet med detta examensarbete är att ge en översikt över tänkbara alternativa metoder, och att sedan välja ut den mest intressanta för tester och ytterligare jämförelser med Pörns metod. En ”intressant metod” är härvidlag en metod som uppfyller erforderliga krav på tillförlitlighet och relevans och som dessutom, i någon mening, är *enklare* än Pörns metod.

Den metod som slutligen föreslås har utvecklats av Jussi Vaurio och används sedan ett par decennier på kärnkraftverket Loviisa i Finland. Vaurios metod är relativt enkel både i teorin och praktiken. Dess styrkor är framförallt den matematiska genomskådligheten, användarvänligheten och de snabba beräkningarna. Vaurios metod bygger likväl på den metodik som föregick Pörns metod även på de svenska kärnkraftverken. Metoden är därmed väl förankrad i kärnkraftstraditionen.

Med reservation för att ytterligare tester måste genomföras föreslås Vaurios metod som ett fullgott och i vissa avseenden bättre alternativ till Pörns metod.

Förord

Föreliggande rapport är resultatet av dryga sex månaders intensivt arbete. Mina intryck från dessa sex månader går emellertid långt utöver vad som kan sammanfattas på 75 sidor text. Jag har fått ta del av en teknologisk kultur där perspektiven många gånger svindlar och där krass know-how lever jämsides med högtflygande matematik och filosofiska djupsinnigheter. Kort sagt; en värld där den nyfikne aldrig går lottlös.

De människor jag mött är många och alla förtjänar ett stort tack. Inte minst de som oförblommerat ställt upp på intervjuer och som kommit med uppmuntrande kommentarer under arbets gång.

Några personer förtjänar ett särskilt omnämnande. Till att börja med vill jag tacka min handledare Anna Gabrielsson, Lars Pettersson och Sven Göran Skagerman på Vattenfall Power Consultant för ett gott samarbete under hela projektiden. (Det är förövrigt tack vare Anna som detta har blivit en rapport och inte ett epos). Jag vill också tacka Kurt Pörn som med stort hjärta hjälpt mig i det matematiska tolkningsarbetet, samt Jussi Vaurio för värdefull input.

Ett särskilt stort tack riktar jag till min ämnesgranskare prof. Sven Erick Alm som har varit mitt viktigaste teoretiska bollplank och en stor inspiratör från början till slut.

Stockholm i maj 2006

Vidar Hedtjärn Swaling

Innehållsförteckning

1 INLEDNING	3
1.1 SYFTE OCH PROBLEMSTÄLLNING	4
1.2 GENOMFÖRANDE	4
1.2.1 Litteraturstudie	5
1.2.2 Intervjuer	5
1.2.3 Tester	5
1.3 AVGRÄNSNINGAR	5
1.4 RAPPORTENS UPPLÄGGNING	6
1.4.1 Läsråd	6
1.5 BEGREPP	6
1.6 FORMALIA	6
2 BAKGRUND OCH PROBLEMDISKUSSION	7
2.1 TUD-SYSTEMET	7
2.2 T-BOKEN	8
2.3 UPPDRAGET	9
3 INTRODUKTION TILL BAYESIANSK STATISTIK	9
3.1 BAYES SATS SOM LAG...	9
3.2 ... OCH SOM VERKTYG	11
4 BAYESIANSK VS. KLASSISK STATISTIK	13
5 BAYESIANSK STATISTIK – DEN ENKLA MODELLEN	15
5.1 LIKELIHOODFUNKTIONEN	16
5.2 A PRIORIFÖRDELNINGEN	17
5.2.1 Konjugerade fördelningar	17
5.2.2 Informativa a priorifördelningar	17
5.2.2.1 Problem vid härledning av informativa a priorifördelningar...	18
5.2.2.2 ... och några möjliga lösningar	19
5.2.3 Icke-informativa a priorifördelningar	19
6 BAYESIANSKA METODER INOM KÄRNKRAFTSOMRÅDET	20
6.1 INLEDNING	20
6.1.1 Vanliga antaganden	21
6.1.2 Gruppering av data	22
6.2 TVÅSTEGS BAYES	22
6.2.1 Tvåstegsmodellen	23
6.2.2 Prior och hyperprior	25
6.3 PARAMETRIC EMPIRICAL BAYES (PEB)	27
6.3.1 Skattningsmetoder	27
6.3.2 Problem med PEB	28
7 METODER I URVAL	29
7.1 TVÅSTEGS BAYES	30
7.1.1 Pörn	30
7.1.1.1 Bakgrund	31
7.1.1.2 Tvåstegsmodellen	33
7.1.1.3 Likelihoodfunktionen	33
7.1.1.4 Prior	34
7.1.1.5 Hyperprior	35
7.1.1.6 Tillämpning	37
7.1.2 ZEDB-metoden	38
7.1.3 Övriga	39
7.2 PARAMETRIC EMPIRICAL BAYES	39
7.2.1 Varios PREB	40

7.2.2 Övriga	42
7.3 ETT KLASSISKT ALTERNATIV	43
8 SYNUNKTER FRÅN ANVÄNDARNA	44
9 DISKUSSION OCH VAL AV "TESTMETOD"	45
10 TESTER	46
10.1 GENOMFÖRANDE OCH RESULTAT	47
10.1.1 Implementering	48
10.1.2 Redovisning av resultat	48
10.2 ANALYS	49
10.2.1 Pörn, Vaurio och ZEDB	49
10.2.2 Extremfall	51
10.2.3 Sammanfattning	52
11 AVSLUTANDE DISKUSSION	53
12 SLUTSATSER	55
13 FÖRSLAG TILL FORTSATT ARBETE	55
14 REFERENSER	56
14.1 TRYCKTA REFERENSER	56
14.2 INTERVJUER OCH KORRESPONDENS	58
15 APPENDIX	60
APPENDIX A: JÄMFÖRELSETABELLER	60
APPENDIX B: INDATA BENCHMARK	66
APPENDIX C: T-BOKSTABELL	67
APPENDIX D: MATLAB-KOD	68
APPENDIX E: VAURIOS "ALGORITM" (PREB)	70
APPENDIX F: MATEMATIK	71

1 Inledning

Detta examensarbete är ett led i en fortgående kvalitetssäkring av modeller och verktyg för hantering av driftsäkerhet inom kärnkraftsområdet. Arbetet har utförts på uppdrag av TUD-gruppen med Vattenfall Power Consultant AB som administratör.¹ Examensarbetet avser civilingenjörsexamen på STS-programmet (*System i Teknik och Samhälle*), Uppsala universitet.

TUD (Tillförlitlighet, Underhåll och Drift) är ett driftsäkerhetsdatasystem som ägs av de nordiska kärnkraftsbolagen.² Systemet samlar felhändelseinformation och komponentdata från tolv svenska och två finska kärnkraftverk. Utifrån dessa data beräknas tillförlitlighetsparametrar för särskilt kritiska fel hos komponenter i kärnkraftverk. Parametrarna sammanställs i *T-boken* och används sedan som indata till de för kärnkraftverken obligatoriska *probabilistiska säkerhetsanalyserna* (PSA).³

För härledning och beräkning av de aktuella parametrarna används en *bayesiansk* statistisk metod utvecklad av Kurt Pörn. Metoden presenterades 1990 i doktorsavhandlingen *On Empirical Bayesian Inference Applied to Poisson Probability Models* (Pörn 1990). Som titeln antyder rör det sig om ett sätt att med bayesiansk metodik skatta parametern i en Poissonprocess. Denna modell tillämpas på *kontinuerligt driftsatta* komponenter, medan en utvidgad variant av samma modell används för komponenter i *standby*. De parametrar som skattas är, enkelt uttryckt, *felintensitet* respektive *felsannolikhet per behov*. Beräkningarna görs i *T-Code*, en för ändamålet avsedd programvara som även den har utvecklats av Kurt Pörn.

I kärnkraftssammanhang är nämnda parametrar i regel mycket små och följaktligen svåra att beräkna. Dessutom måste beräkningsresultaten uppfylla högt ställda krav på tillförlitlighet och relevans. Detta kräver i sin tur minutiösa insatser vad gäller insamling och klassificering av data, samt att de statistiska slutsatserna vilar på en solid matematisk grund. En sådan grund har traditionellt tillskrivits bayesiansk statistik i olika tappningar. Bayesianska metoder har fördelen av att fungera även i fall med extremt torftigt dataunderlag. Dessutom vilar de på en av inferensteorins viktigaste satser; Bayes sats. Den bayesianska statistiken har varit så framgångsrik inom kärnkraftsområdet att den i dag kan sägas utgöra ett världsomspännande paradig.

Emellertid är inte alla problem lösta. Metoderna har stötts och blötts genom åren och olika alternativ, det ena mer sofistikerat än det andra, har föreslagits. I detta avseende är Pörns metod utan tvekan "top of the line". Men förädlingen har delvis skett på bekostnad av begreppslig transparens, vilket i sig kan vara en riskfaktor. I detta examensarbete undersöks möjligheten av att, trots allt, tillämpa en enklare metod än Pörns för härledning och beräkning av *T-boken*s parametrar.

¹ Vattenfall Power Consultant AB hette tidigare SwedPower AB. Namnbytet genomfördes den 24 april 2006.

² Forsmarks Kraftgrupp AB, OKG Aktiebolag, Ringhals AB och Teollisuuden Voima OY (TVO).

³ På engelska står PSA för *Probabilistic Safety Assessment*. En äldre term för samma sak är *Probabilistic Reliability Analysis* (PRA).

1.1 Syfte och problemställning

Examensarbetets syfte är att ge en överblick över tänkbara matematiska metoder för härledning av tillförlitlighetsparametrar från driftdata, samt, genom att välja den som tycks mest intressant och implementera den i ett verktyg, göra en jämförelse med den metod som för närvarande används.

Syftesformuleringen är hämtad från den ursprungliga projektbeskrivningen och gäller alltså jämt. Emellertid lämnar den ett par frågor obesvarade:

- Vad är en ”matematisk metod”? Frågan är väsentlig med tanke på att det är ”tänkbara matematiska metoder” som ska studeras, samt att ”metoden” ska implementeras i ett verktyg.
- Vad är ”mest intressant”? Dvs. utifrån vilka kriterier väljs en alternativ metod?

Den första frågan är fundamental men svår att ge ett definitivt svar på. Pörns metod skulle till att börja med kunna sägas omfatta allt från gruppering av data till de algoritmer i T-Code som så småningom spottar ut parametervärden. Pörn löser s.a.s. problemen ”från ax till limpa”. Men allt detta är inte vad som normalt kallas matematik. Den matematiska metoden borde rimligen komma in någonstans ”mellan data och numerik”. En tänkbar precisering är att med ”matematisk metod” i första hand avse *de ekvationer som beskriver hur statistiska slutsatser härleds utifrån teoretiska antaganden och ett givet dataunderlag*. I så fall utesluts t.ex. numeriska metoder och principer för gruppering av data. Emellertid tror jag att en sådan avgränsning kan vara svår att upprätthålla. Jag vill påpeka att det i sammanhanget först och främst handlar om *tillämpad matematik*, vilket betyder att frågor som rör data och numerik kan vara högst väsentliga för hur matematiken har formulerats. Jag tror därför att den föreslagna avgränsningen måste tolkas generöst, som en riktlinje snarare än ett direktiv.

Vidare går det knappast att tillskriva Pörn alla aspekter av den ifrågavarande ”matematiska metoden”. Dess byggstenar är i många fall allmängods. Även om Pörn står som ensam upphovsman till (och är berömd för) vissa led i metoden så är det snarare kompositionen som är unik. Pörns metod är såtillvida ett aggregat av utbytbara delar och detsamma antas gälla vilken tänkbar metod som helst. Det mest intressanta alternativet kan alltså mycket väl vara ett hopplock från det ”statistiska smörgåsbordet”.

Nu till den andra frågan: Vilka alternativ är intressanta? Svaret beror naturligtvis på *vem* som vill ha ett alternativ och *varför*. Sådana bakgrundsfakta presenteras utförligare i kapitel 2. Här nöjer jag mig med att säga att jag i samråd med min uppdragsgivare och handledare har bestämt att i första hand söka en metod som i någon mening är *enklare* än Pörns. Samtidigt får enkelheten inte utan vidare vara på bekostnad av noggrannhet och precision. En alternativ metod bör rimligen ha den prestanda som omständigheterna kräver. Därmed inte sagt att den nödvändigtvis måste ha samma prestanda som Pörns metod.

1.2 Genomförande

Arbetet med att hitta ett alternativ till Pörns metod har utgått från en litteraturstudie omfattande såväl primär- som sekundärlitteratur. Litteraturstudien har syftat till att ge en överblick över *tänkbara* alternativ till Pörns metod och har bildat underlag för ett urval av *särskilt in-*

tressanta alternativ. Det alternativ som sedan förefallit *mest intressant* har implementerats i ett verktyg och jämförts med Pörns metod med avseende på beräkningsresultat.

1.2.1 Litteraturstudie

I den ursprungliga projektbeskrivningen fastslogs att litteraturstudien skulle omfatta ”statistiska metoder för härledning av tillförlitlighetsparametrar från driftdata i en databas som TUD-databasen”. I första hand har artiklar och rapporter från olika vetenskapliga tidskrifter studerats. Jag har också använt mig av T-boken i olika versioner, tekniska rapporter och programvarumanualer (däribland dokumentationen till T-Code) liksom standardverk om både bayesiansk och klassisk statistik. En betydande del av litteraturstudien har ägnats åt Pörns metod.

I litteratursökningen har jag huvudsakligen använt mig av databaserna Compendex, Inspec och INIS. Litteratur har rekvirerats från Uppsala universitetsbibliotek, TUD-kansliet eller matematiska institutionen på Uppsala universitet.

1.2.2 Intervjuer

Som komplement till litteraturstudien har jag intervjuat personer med olika kopplingar till T-boken. Genom intervjuerna har jag dels kunnat kartlägga användningen av T-boken, dels fått en överblick över olika aktörers krav och förväntningar på en alternativ metod. Därtill har intervjuerna gett mig en djupare förståelse för branschens villkor och kopplingarna mellan dess olika aktörer. Intervjuerna har gjorts både skriftligt och muntligt. Muntliga intervjuer har refererats och referaten har sedan godkänts av den intervjuade. Spontan korrespondens har också förts med särskilt centrala aktörer, däribland Kurt Pörn. Information som erhållits från uppdragsgivaren (TUD-gruppen och TUD-kansliet) refereras inte i texten. Denna information gäller uteslutande bakgrunden till examensarbetet (kapitel 2) och data som använts i testerna (kapitel 10). Slutligen har ett sammanträffande med Nordiska PSA-gruppen givit värdefull input liksom ett seminarium på matematiska institutionen, Uppsala universitet.

1.2.3 Tester

Beräkningar har gjorts i MATLAB och toolboxen *Statistics*. För jämförelser med Pörns metod har jag använt resultat från en benchmark (T-book – ZEDB Benchmark, 2004) för Pörns metod och den metod som används inom ZEDB.⁴ Jämförelser har därmed kunnat göras även med ZEDB-metoden. Kompletterande tester har baserats på resultat i T-boken version 5 och 6. In- och utdatafiler för samtliga tester har erhållits från TUD-kansliet.

1.3 Avgränsningar

Litteraturstudien bedrevs till en början ganska förutsättningslöst. Dels eftersom jag initialt hade mycket begränsade kunskaper om Pörns metod, dess tekniska och teoretiska kontext liksom kärnkraftsbranschen i stort, dels eftersom det ingick i uppdraget att undersöka både kända och ”okända” alternativ. Spelrummet har därmed varit stort. Icke desto mindre har jag varit tvungen att börja och sluta någonstans. Jag har därför alltmör kommit att fokusera på metoder som används och har använts inom *kärnkraftsområdet och liknande områden*, dvs. områden med högt säkerhetstänkande och sparsamt dataunderlag. En av mina erfarenheter från litteraturstudien är att kärnkraftsbranschen varit starkt drivande i utvecklingen av sådana

⁴ ZEDB är, enkelt uttryckt, Tysklands motsvarighet till TUD.

metoder, och att de metoder som används inom t.ex. offshoreindustrin väsentligen är desamma.

Andra avgränsningar hänger samman med den precisering av syftet som gjordes inledningsvis: För det första kommer jag inte att undersöka alla aspekter av Pörns metod utan framförallt dem som rör den statistiska slutledningsprocessen, dvs. innehållet i avhandlingen (Pörn 1990) och ändringar som omedelbart rör detta innehåll. Med "tänkbara alternativ" förstås således metoder eller metodkomplex som motsvarar innehållet i Pörns avhandling, dvs. som i någon mening "gör samma sak men på ett annat sätt". Frågor om *common cause failures* (CCF), datas kvalitet (rapportering av felhändelser etc.) samt modellering av expertkunskap är viktiga men kan likväl endast behandlas i förbigående. För det andra kommer ett högst begränsat utrymme att ägnas åt alternativ som förefaller vara "minst lika komplicerade" som Pörns metod (t.ex. hierarkisk Bayes med fler än två steg eller modeller med tidsberoende felintensitet).

1.4 Rapportens uppläggning

I kapitel 2 ges till att börja med en generell bakgrundsbeskrivning. I kapitel 3 redogör jag sedan för den bayesianska statistikens grundvalar, detta för att bjuda in läsare som befinner sig på samma nivå som jag själv då detta arbete påbörjades. I kapitel 4 diskuteras skillnaden mellan klassisk och bayesiansk statistik. Där behandlas även kopplingen mellan bayesiansk statistik och PSA mer utförligt. Kapitel 5, 6 och 7 innehåller väsentligen resultat av litteraturstudien; kapitel 5 och 6 ägnas åt de principiella möjligheterna att framställa en alternativ metod medan kapitel 7 ägnas åt redan befintliga metoder. Särskilt stort utrymme ägnas åt Pörns metod. Resultat av intervjuerna redovisas huvudsakligen i kapitel 8. I kapitel 9 väljs (utifrån tidigare redovisade resultat) en metod för tester och ytterligare jämförelser med Pörns metod. Testerna redovisas och analyseras i kapitel 10. I kapitel 11 förs en avslutande diskussion. De viktigaste slutsatserna redovisas i kapitel 12 och i kapitel 13 ges slutligen några förslag till fortsatta efterforskningar.

1.4.1 Läsråd

Kapitel 6 och 7 är de teoretiskt mest avancerade och förmodligen de mest svårsmälta i denna rapport. Den läsare som vill ha en djupare förståelse för hur rapporten därefter utvecklar sig bör inte hoppa över dessa. Den som vill ha snabbare läsning bör åtminstone läsa kapitel 8, 9, 11 och 12.

1.5 Begrepp

Min ambition är att förklara specialtermer och ovanliga begrepp vid första förekomsten i texten. Den matematiska nomenklaturen utgår ifrån Råde & Westerberg (1998) men är delvis anpassad till konventioner i övrig studerad litteratur. Ambitionen har varit att hitta ett så homogent och enkelt skrivsätt som möjligt och att vara konsekvent i detta genom hela rapporten. Några särskilt viktiga *fördelningar* redovisas i appendix F.1.

1.6 Formalia

Referenser ges direkt i texten enligt Harvardsystemet. Vidare gäller en referens placerad *inne i en mening* bara för ifrågavarande mening. En referens placerad *efter en mening* gäller från tidigare referens i samma stycke. Intervjureferenser sätts inom klammer enligt principen [Namn X], där X används om upprepade intervjuer gjorts. De olika versionerna av T-boken betecknas T1, T2, ..., T6.

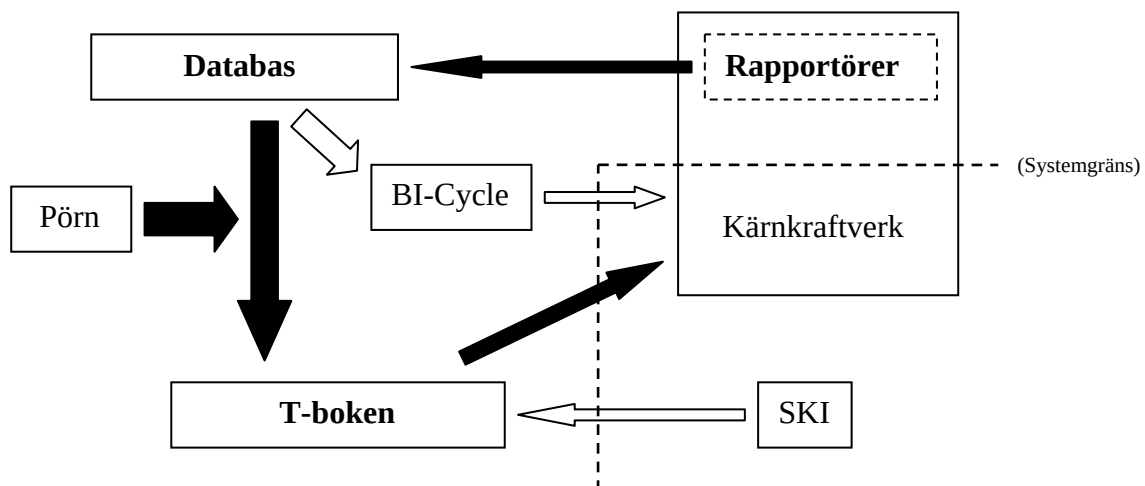
2 Bakgrund och problemdiskussion

Tidigt i detta arbete frågade jag mig vad det var för slags problem jag stod inför. Var problemet matematiskt, tekniskt eller organisatoriskt? Vem eller vilka hade problemet och varför? Sådana frågor har varit väsentliga för hur jag själv har förstått och därmed närmat mig problemet. För att ge läsaren en liknande utgångspunkt följer här en beskrivning av det sammanhang som idén till examensarbetet uppstått i.

2.1 TUD-systemet

TUD-databasen startades på frivilligt initiativ av de svenska kärnkraftsbolagen i mitten av 70-talet. 1981 tillkom det finska bolaget TVO som driver två reaktorblock av svensk design. TUD-databasen kan sägas vara kärnan i ett större system av teknisk såväl som social karaktär, med rapportörer på respektive kärnkraftsanläggning, en styrelse (TUD-gruppen) och ett kansli som administrerar databasen och ger ut T-boken.⁵ Till detta system kan även Pörns verksamhet räknas. I systemets närmaste omgivning finns kärnkraftverken och Statens kärnkraftsinpektion (SKI).

TUD-systemet behandlar information från alla process- och säkerhetsrelaterade system på de ingående kärnkraftsanläggningarna. Indata utgörs väsentligen av felrapporter, komponentbeskrivningar, underhållsdata och drifttidsavläsningar. Merparten av den sammanställda informationen återförs till anläggningarna via analysverktyget BI-Cycle som används för identifiering av olika nyckelproblemområden (komponenter med särskilt hög underhållskostnad o dyl.). En liten del av informationen, den som gäller särskilt kritiska fel, skickas till Pörn Consulting i Nyköping för statistisk behandling. ”Felstatistiken” återförs sedan till anläggningarna via T-boken.



Figur 1. TUD-systemet. De svarta pilarna illustrerar de kritiska felens väg ”från observationer till felstatistik”. I TUD-systemets närmsta omgivning finns Kärnkraftverken och SKI.

⁵ TUD-kansliet är förlagt till Vattenfall Power Consultant AB.

2.2 T-boken

Även om Pörns metod är tänkt att kunna användas inom många olika områden är den först och främst utvecklad för T-boken (Pörn 1990, s. 1).⁶ Omständigheterna runt T-boken är i sin tur ganska speciella.

Indata till T-boken är enkelt uttryckt *antal fel respektive drifttid* för komponenter, medan utdata är motsvarande felbenägenheter. Dessa utgör i sin tur indata till de *probabilistiska säkerhetsanalyser* (PSA) som är en väsentlig och obligatorisk del i det normala säkerhetsarbetet på kärnkraftverken. I dessa analyser konstrueras s.k. *felträd* där anläggningens olika system och delsystem ordnas hierarkiskt med avseende på felhändelser. Överst i trädet finns *topphändelsen* Q , tex. att reaktorhärden skadas. Realiseringen av Q antas i sin tur bero av realiseringen av andra händelser A och B som i sin tur beror av $C, D, E, \dots$ etc. På detta sätt bryts hela anläggningen upp i ett system av noder och länkar. På länkarna sitter s.k. *grindar* som definierar det logiska beroendet mellan händelserna. Grindarna är vanligen av typen ”och” eller ”eller”. Felträdet ger alltså information av typen ” Q inträffar om A och B inträffar, eller om C och D inträffar”. Felträdet analyseras i PSA-verktyget RiskSpectrum[®] (Relcon AB). Resultatet används för utvärdering av kärnkraftsanläggningens *riskprofil* dvs. hur olika delar av anläggningen bidrar till utsläpp, skador på reaktorhärden etc. Med hjälp av dessa analyser kan säkerhetshöjande åtgärder sättas in där de bäst behövs.

Längst ner i felträdet finns *bashändelserna*. Dessa uttrycker olika *felmoder* dvs. olika sätt att *fela* hos t.ex. en komponent eller en operatör. Exempel på felmoder hos en pump är ”obefogat stopp” och ”utebliven start”. Felmoder hos en ventil kan vara ”obefogad lägesändring”, ”utebliven öppning” etc. Dessa felmoder har i sin tur *felsannolikheter* eller *felintensiteter* beroende på om felen uppträder i diskret eller kontinuerlig tid. ”Utebliven start” är ett exempel på en diskret händelse och ”obefogad lägesändring” ett exempel på en kontinuerlig. Det är sådan *teknisk* felstatistik som finns upptagen i T-boken. Operatörsfel finns alltså inte med. I T-boken ges dessa felsannolikheter i form av *fördelningar*: Om felintensiteten betecknas λ så uttrycker *fördelningen* för λ sannolikheten att λ antar fördelningens olika värden. I T-boken presenteras fördelningarnas percentiler (5-, 50- och 95%) samt medelvärdet i *tabellform*. Fördelningarna är vidare sammanställda anläggningsvis och presenteras tillsammans med en s.k. *generisk* fördelning gällande för hela populationen (alla kärnkraftverk). Ett exempel på en T-bokstabel ges i appendix C. Vilka komponenter, felmoder etc. som ska finnas med i en ny utgåva av T-boken avgör TUD-gruppen tillsammans med PSA-avdelningarna på respektive anläggning.

Den första versionen av T-boken gavs ut 1982 och innehöll driftstatistik för 21 reaktorår.⁷ Version 6 gavs ut 2005 och täcker 315 reaktorår. Sverige är därmed ett föregångsland vad gäller dokumentation och uppföljning av driftförhållanden inom kärnkraftsområdet (Pörn 1990, s. 98). Förutom pappersversionen finns T-boken också i form av en CD och datafilen *Tbokrisk* som är kompatibel med RiskSpectrum[®]. T-bokens parametrar (t.ex. λ) skattas på basis av felrapporter i TUD-systemet samt de *rapporterbara omständigheter* (RO) som rapporteras direkt till Statens Kärnkraftsinspektion (SKI).

T-boken omfattar endast s.k. *funktionshindrande fel* dvs. ”[---] fel av sådan art, att komponentens funktion anses ha gått förlorad” (T6, s. 23). Typiskt för kärnkraftsområdet är att sådana fel är extremt ovanliga; fortfarande efter 315 reaktorår finns komponentgrupper som inte haft

⁶ Bokstaven ”T” står för ”tillförlitlighetsdata”.

⁷ Ett reaktorår motsvarar den normala drifttiden för en kärnkraftsreaktor under ett kalenderår.

ett enda funktionshinder fel. En sådan komponentgrupp skulle med klassiska statistiska metoder få felintensiteten noll, vilket är orealistiskt (T6, s. 31). *Bayesianska metoder* kringgår detta problem genom att de tillåter utnyttjande av s.k. *subjektiv* statistisk information; subjektiva sannolikheter är noll bara för *logiskt omöjliga händelser* och några sådana finns inte i T-boken (mer om subjektiva sannolikheter i kapitel 3 och 4). Typiskt för bayesianska metoder är också att de ger skattningar i form av fördelningar istället för punkter. Till de bayesianska metodernas nackdelar hör att de ofta leder till mycket komplicerade beräkningar.

2.3 Uppdraget

Sedan 1992 har Pörn Consulting anlitas för uppdatering av T-bokens parametrar. Pörns metod är implementerad i beräkningsprogrammet T-Code och beräkningarna har utförts av Pörn själv. Under 2005 köptes T-Code upp av TUD-gruppen. Därmed har både T-Code och den matematiska metoden på ett naturligt sätt blivit föremål för den kvalitetssäkringsprocess som TUD-systemet genomgår och som detta examensarbete är en del av. Emellertid handlar inte examensarbetet i första hand om kvalitetssäkring. TUD-gruppen arbetar generellt mot *ökad kontinuitet* i systemets olika delar. Vad gäller T-boken är det långsiktiga målet att den ska kunna uppdateras oftare, vilket gör det relevant att också söka efter *enklare alternativ* till de metoder som används idag. Ett problem med Pörns metod är vidare att den uppfattas som komplicerad och att beräkningarna kräver stor förtrogenhet med T-Code. Att T-Code dessutom har administrerats av ett enmansföretag har gett hela ”paketet” (metod och program) karaktären av ”svart låda”. Att minska beroendet av sådana svarta lådor ligger i linje med TUD-gruppens strävan mot kontinuitet. Det är snarast här examensarbetet kommer in i bilden. Vad som söks är ett alternativ till Pörns metod som ökar genomskådligheten och därmed underlättar uppdatering av T-bokens parametervärden.

Ytterligare ett motiv till examensarbetet är att Pörns metod utvecklades för en situation med extremt lite data. Metoden togs i bruk med ett empiriskt underlag motsvarande omkring 100 reaktorår. Idag är antalet reaktorår mer än det tredubbla. Därmed antas möjligheterna också vara större att använda enklare statistiska metoder.

3 Introduktion till bayesiansk statistik

3.1 Bayes sats som lag...

Bayesiansk statistik handlar enkelt uttryckt om att dra slutsatser med hjälp av Bayes sats. Bayes sats handlar i sin tur om *betingade sannolikheter*. Dessa definieras av sambandet

$$P(A|B) = \frac{P(A \cap B)}{P(B)} \quad (\text{Blom 1984, s. 33}), \quad (3.1)$$

där $P(A|B)$ förstås som ”sannolikheten att A är sann givet att B är sann” och $P(A \cap B)$ som ”sannolikheten att både A och B är sanna”. Bayes sats handlar om *logiska relationer* mellan händelser (eller utfallsrum) och säger därmed inget om orsakssamband.

Det var utifrån (3.1) som den engelska prästen Thomas Bayes (1702-1761) härledde sin berömda sats. I sin enklaste version är Bayes sats en direkt följd av (3.1):

$$P(A|B) = \frac{P(B|A)P(A)}{P(B)}. \quad (3.2)$$

Satsen kan tolkas som ett sätt att ”vända på betingningar” (Englund 2000, s. 178).

Exempel 3.1: En flock består av ett antal fåglar varav några är ankor. Om $P(A)$ är sannolikheten att en slumpmässigt vald fågel f är en anka och $P(B)$ är sannolikheten att f är sjuk så är $P(B|A)$ sannolikheten att ” f är sjuk givet att f är en anka”, dvs. *sjukdomsfrekvensen bland ankorna*. Bayes sats ger nu även $P(A|B)$ dvs. sannolikheten att ”om f är sjuk så är f en anka” dvs. *andelen ankor bland det totala antalet sjuka fåglar*.

□

Att A betingas på B betyder att B anger det relevanta utfallsrummet som i exempel 3.1 är mängden sjuka fåglar. $P(B)$ har därmed rollen av *normaliseringsfaktor* (se nedan).

Bayes visade att satsen även gäller vid diskretisering av utfallsrummet A (Blom 1984, s. 36). Bayes sats får då följande utseende:

$$P(A_i|B) = \frac{P(B|A_i)P(A_i)}{P(B)} \quad (3.3)$$

där

$$P(B) = \sum_{j=1}^n P(B|A_j)P(A_j) \quad (3.4)$$

enligt lagen om total sannolikhet (Blom 1984, s. 26). Här framgår innebörden av normaliseringsfaktorn tydligare. Att *betinga* är att definiera det relevanta utfallsrummet och detta görs genom att sannolikheten för samtliga relevanta händelser *summeras*. $P(B)$ uttrycker alltså sannolikheten för mängden *möjliga utfall* (dvs. B).⁸

Exempel 3.2: I en fabrik tillverkas komponenter vid maskinerna A_1 , A_2 och A_3 i proportionerna 20:30:50. Av produktionen är respektive 5%, 3% och 2% defekt. Komponenterna blandas innan de lämnar fabriken. Om en slumpmässigt vald komponent är felaktig, hur stor är sannolikheten att den har tillverkats vid maskin A_1 ?

Låt $P(A_i)$ vara sannolikheten att en komponent kommer från en viss maskin. Om vidare $P(B)$ är den *totala sannolikheten* att en komponent är defekt så är $P(B|A_i)$ sannolikheten att en komponent är defekt givet att den kommer från en viss maskin. Bayes sats ger nu:

$$P(A_1|B) = \frac{P(B|A_1)P(A_1)}{\sum_{j=1}^3 P(B|A_j)P(A_j)} = \frac{0,05 \cdot 0,20}{0,05 \cdot 0,20 + 0,03 \cdot 0,30 + 0,02 \cdot 0,50} = \frac{0,01}{0,029} \approx 0,345$$

På samma sätt är $P(A_2|B) = 0,009/0,029 \approx 0,31$ och $P(A_3|B) = 0,01/0,029 \approx 0,345$.

(Exemplet är baserat på Blom 1984, s. 37)

□

Av exempel 3.2 framgår att summan av alla $P(A_i|B) = 1$, dvs.

⁸ I en mening är alla sannolikheter betingade. $P(B|A)$ och $P(A)$ i exempel 3.1 är t.ex. betingade på händelsen att ” f är en fågel”. Detta är emellertid en *säker* händelse. Om händelsen att ” f är en fågel” betecknas F så är $P(F) = 1$. Normalisering med avseende på denna händelse ”märks” alltså inte.

$$\sum_{i=1}^n P(A_i|B) = 1. \quad (3.5)$$

Vidare är $P(B)$ konstant med avseende på A_i (i exempel 3.2 är $P(B) = 0,029$). Detta innebär att Bayes sats kan förenklas till

$$P(A_i|B) = K \cdot P(B|A_i)P(A_i) \quad (3.6)$$

eller ännu enklare

$$P(A_i|B) \propto P(B|A_i)P(A_i)$$

där " $\propto$ " betecknar *proportionalitet*. Normaliseringsfaktorn $P(B)$ är alltså en *proportionalitetskonstant* som garanterar att alla $P(A_i|B)$ summerar till 1. Genom att sätta $P(B)$ i nämnaren garanteras med andra ord att $P(A_i|B)$ blir en sannolikhet, dvs. $P(A_i|B)$ antar värden i intervallet $[0, 1]$ (vilket likväl är innebörden av Kolmogorovs första axiom, se kapitel 4). $P(B)$ kallas också för *marginalsannolikhet*.

3.2 ... och som verktyg

I bägge de föregående exemplen är all statistisk information given på förhand, slutsatsen finns s.a.s. implicit i premisserna. Därmed har Bayes sats karaktären av *lag* snarare än *metod*; satsen beskriver *den logiska relationen* mellan betingade sannolikheter och marginalsannolikheter. Sådillvida säger den också något om innebörden av själva begreppet *sannolikhet*. Det *empiriska problemet*, dvs. att beräkna sjukdomsfrekvensen hos ankor, felsannolikheter för maskiner etc., kvarstår dock.

Det karaktäristiska för bayesiansk statistik är att Bayes sats används för att lösa just sådana empiriska problem. Bayes sats uppfattas med andra ord som ett *verktyg* för att erhålla statistisk information utifrån erfarenheter. Samtidigt ses satsen som en garant för de statistiska slutsatsernas logiska riktighet (Atwood et al. 2003, ch. 6, s. 2. Kaplan 1986, s. 123). Bayes sats tolkas alltså dels som en *lag* beträffande logisk slutledning, dels som en *metod* för statistiska beräkningar. De två tolkningarna återspeglas i följande citat:⁹

1. "[---] Bayes' theorem is the fundamental law governing the process of logical inference." (Kaplan 1986, s. 123)
2. "The fundamental tool for the specialization (or 'updating') of probabilities, when new evidence becomes available, is Bayes' theorem" (Apolostakis et al. 1980, s. 321)

Huruvida dessa tolkningar är förenliga kan möjligen diskuteras. Det är hur som helst den senare tolkningen som ska behandlas här.

Exempel 3.3: En maskin tillverkar komponenter av vilka några är defekta. Det antas nu att maskinen har en viss *felintensitet*. De enda observerbara storheterna är emellertid antalet felaktiga komponenter och den tid som tillverkningen har pågått, dvs. drifttiden. Hur ska felintensiteten skattas?

⁹ Citaten är hämtade från två av den bayesianska statistikens förgrundsgestalter. Kaplan står som medförfattare även till det andra citatet.

Situationen kan modelleras enligt följande: Till att börja med antas felintensiteten vara konstant under den aktuella tiden. Intensiteten betecknas λ . Vidare betraktas antalet fel som utfall av en stokastisk variabel X . Tiden modelleras inte explicit eftersom den antas vara deterministisk och mätbar (jmf. Carlin & Louis 2000, s. 17). På något sätt måste nu X relateras till λ . Detta kan göras genom att observationerna antas vara realiseringar av en *Poissonprocess* där parametern λ ”styr” utfallet. Om detta uppfattas som en *betingning* kan sannolikhetsfunktionen för X betecknas $p(x|\lambda)$. Bayes sats ”vänder” nu på denna betingning så att istället λ erhålls som funktion av x . $\square$

Det är utmärkande för bayesiansk statistik att (som i exempel 3.3) observationerna uppfattas som *betingade* på modellparametern i en sannolikhetsfunktion (Englund 2000, s. 178). Denna har således rollen av *likelihoodfunktion* (se avsnitt 5.1). Genom tillämpning av Bayes sats erhålls sedan en sannolikhetsfunktion eller *fördelning* även för modellparametern. Om den sökta parametern är en intensitet λ och likelihoodfunktionen betecknas $p(x|\lambda)$ erhålls fördelningen för λ enligt

$$p(\lambda_i|x) = \frac{p(x|\lambda_i)p(\lambda_i)}{p(x)} \quad (3.7)$$

med *marginalfördelningen*

$$p(x) = \sum_{i=1}^n p(x|\lambda_i)p(\lambda_i). \quad (3.8)$$

Metoden ger alltså ingen direkt skattning av λ . Istället erhålls en fördelning som beskriver *sannolikheten för olika värden på λ* . Skillnaden mellan $p(\lambda)$ och $p(\lambda|x)$ är vidare att de uttrycker denna sannolikhet *före* respektive *efter* att data beaktats. Att sannolikheter kan ändras eller *uppdateras* på grundval av erfarenhet är en av de stora skillnaderna mellan bayesiansk och klassisk statistik. Bayesianska sannolikheter är i själva verket *subjektiva* vilket betyder att de uttrycker förväntan, kunskap eller osäkerhet om den storhet som ska skattas. (Det bayesianska sannolikhetsbegreppet behandlas utförligare i kapitel 4.) Fundamentalt är ansättandet av en s.k. *a priorifördelning* $p(\lambda)$. A priorifördelningen uttrycker i någon mening vad som på förhand (a priori) är känt om λ . Genuin osäkerhet kan t.ex. uttryckas genom att $p(\lambda)$ väljs likformig (Englund 2000, s. 177). Via likelihoodfunktionen $p(x|\lambda)$ uppdateras sedan a priorifördelningen till en *a posteriorifördelning* $p(\lambda|x)$. Detta sker genom att $p(\lambda)$ och $p(x|\lambda)$ vikts ihop för ett antal värden på λ .

I praktiken ansätts oftast en kontinuerlig a priorifördelning och likelihoodfunktionen utvärderas för så många parametervärden som möjligt, vilket i det ideala fallet betyder *för alla tänkbara parametervärden*. Härtill används den kontinuerliga versionen av Bayes sats:

$$p(\lambda|x) = \frac{p(x|\lambda)p(\lambda)}{\int_0^{\infty} p(x|\lambda)p(\lambda)d\lambda} \propto p(x|\lambda)p(\lambda). \quad (3.9)$$

Skillnaden gentemot den diskreta versionen är att sannolikhetsfunktionerna ersätts med *tät-hetsfunktioner* (Englund 2000, s. 179). Ett problem med den kontinuerliga versionen är att integralen i nämnaren (dvs. marginalfördelningen) i allmänhet måste beräknas numeriskt,

vilket i många fall kräver avancerad Monte Carlo-simulering (Carlin & Louis 2000, s. 10, 120f).

En central fråga inom bayesiansk statistik är hur a priorifördelningen ska framställas. Den ska som sagt representera kunskap som föregår observationerna. Frågan behandlas närmare i kapitel 5 och 6.

4 Bayesiansk vs. klassisk statistik

Litteraturen ger inget entydigt svar på frågan om skillnaderna mellan bayesiansk och klassisk (även kallad frekventistisk) statistik närmare bestämt består, eller hur fundamentala de är.¹⁰ Hos vissa författare framställs de två riktningarna som komplementära och valet mellan dem som betingat av den *praktiska situationen*. Hos andra blir valet snarast ett *filosofiskt problem* genom att det ena alternativet anses vara sannare än det andra. Ytterligare andra tycks se de båda riktningarna som två sidor av samma mynt; valet mellan dem blir så tillvida ett *val av perspektiv*. Generellt kan sägas att utvecklingen gått från filosofisk debatt mot ökad samsyn och pragmatism (Carlin & Louis 2000, s. 1). Vad Carlin & Louis (2000, s. 6) kallar "the Bayes-frequentist controversy" tillhör alltså framförallt det förgångna. Många grundläggande frågor debatteras emellertid fortfarande. I det följande försöker jag ge en någorlunda samlad bild av diskursen.

Carlin & Louis (2000, s. 5) ger följande beskrivning av skillnaden mellan bayesiansk och klassisk statistik: "The frequentist conditions on parameters and [integrates] over the data; the Bayesian conditions on the data and [integrates] over the parameters." Beskrivningen förklarar varför bayesiansk statistik ger parameterskattningar i form av fördelningar till skillnad från den klassiska statistikens punktskattningar. Ibland tolkas denna skillnad som ett resultat av *huruvida slumpmässighet antas i parametern eller i data*. Englund säger t.ex. att i bayesiansk statistik betraktas *data som fixa och parametern som slumpmässig*, medan det omvända gäller i klassisk statistik (Englund 2000, s. 178). Både Carlin & Louis och Englund framställer alltså de bägge riktningarna som i någon mening motsatta.

Att säga att data betraktas som "fixa" i bayesiansk statistik är emellertid något missvisande eftersom likelihoodfunktionen beskriver en *stokastisk process*. Likväl är det missvisande att säga att parametern är "slumpmässig"; parameterns fördelning tolkas sällan eller aldrig som en *stokastisk fördelning*. Snarare antas den uttrycka en *osäkerhet* om parameterns sanna värde (jmf. Atwood et al. 2003, ch. 6, s. 2). I någon mening överförs alltså datas fördelning (som beskrivs av likelihoodfunktionen) till parametern. Om data är ett stickprov ur en population antas vidare en stor del av osäkerheten härröra från parameterns *variation* i populationen (jmf. Pörn 1990, s. 64).

I klassisk statistik ansätts ingen a priorifördelning vilket betyder att all slump hänförs till data. I en vald population antas alltså individerna vara lika med avseende på den sökta parametern (Atwood et al. 2003, ch. 6, s. 2). Inte heller antas någon osäkerhet beträffande parametern *i sig*. Osäkerhet kan emellertid uttryckas med avseende på *skattningens precision*. Detta görs med hjälp av konfidensgränser. Klassiska och bayesianska konfidensgränser (där de senare utgörs av fördelningens percentiler) har således olika *meningsinnehåll*, vilket i sin tur kan hänföras till att de båda riktningarna håller sig med fundamentalt olika tolkningar av begreppet "sannolikhet".

¹⁰ Klassisk statistik kan betraktas som ett specialfall av frekventistisk statistik (Leonard & Hsu 1999, s. 4). De två begreppen används emellertid ofta synonymt, så även här.

I bayesiansk statistik är sannolikheter *subjektiva*, eller *epistemiska*, vilket betyder att de refererar till interna psykologiska tillstånd som uttrycker kunskap, tro eller förväntan om den sökta parametern (Siu & Kelly 1998, s. 90ff. Leonard & Hsu 1999, s. 5).¹¹ Strängt taget är det alltså fel att tala om "sannolikheten för händelsen H " eftersom sannolikhet inte är en egenskap hos H utan hos oss *själva* (Kaplan 1986, s. 124). Detta förklarar också att sannolikheter ändras då ny information tillkommer (jmf. Siu & Kelly 1998, s. 90).¹² I klassisk statistik är sannolikheter istället *frekventistiska*. En frekventistisk sannolikhet uttrycker "the long-run proportion of times the event occurs in a large number of replications of the experiment" (Leonard & Hsu 1999, s. 5), alternativt "the long-term fraction of times that the event would occur, in a large number of trials" (Atwood 2003, ch. 6, s. 2).

Skillnaden mellan de två begreppen blir särskilt tydlig då data saknas, dvs. vid skattning av sannolikheten för en händelse som efter ett antal försök (eller en viss observationstid) ännu inte inträffat. I sådana situationer använder bayesianen ofta en matematiskt framställd s.k. *icke-informativ* eller *objektiv* a priorifördelning (Vaurio 1990a, s. 55. Siu & Kelly 1998, s. 98). Tanken med en sådan fördelning är att den ska uttrycka att *alla parametervärden är lika troliga* vilket i sin tur antas återspegla ett tillstånd av "complete ignorance" (Siu & Kelly 1998, s. 105). Även om en sådan fördelning framställs på matematisk väg har den alltså en begreppslogisk koppling till kunskap; den är s.a.s. *logiskt subjektiv* (jmf. Vaurio 1990b, s. 127f).¹³

Med det frekventistiska sannolikhetsbegreppet är parametern en okänd *konstant*. Om den ifrågavarande händelsen ännu inte har inträffat ger traditionella skattningsmetoder parametervärdet noll, vilket enligt bayesianen är ett orimligt resultat eftersom det inte går att *tro på*. Även frekventisten skulle underkänna resultatet, men inte för att parametervärdet *i sig* är orimligt, utan för att den *statistiska situationen* är det; resultatet är orimligt eftersom det obefintliga dataunderlaget ger en skattning utan precision. Frekventistiska sannolikheter är såtillvida rent matematiska storheter utan koppling till något tänkbart subjekt. Därmed kan de också sägas vara *logiskt objektiva*.

Ett alternativ för frekventisten är att anta att *händelserna* inträffar enligt någon fördelning, t.ex. en Poissonfördelning, med en *given* parameter. Detta parametervärde är nu *sant eller falskt*. Bayesianen har å sin sida, via a priorifördelningen, tilldelat alla tänkbara parametervärden vissa sannolikheter. Därmed är likväl alla parametervärden *sannolika*; strängt taget kommer inget parametervärde någonsin att visa sig vara falskt eftersom detta skulle kräva oändligt många observationer (mer om detta i avsnitt 5.2.2).

De två sannolikhetsbegreppen medför som sagt skillnader i hur resultat med bayesiansk respektive klassisk statistik tolkas och används. Något som redan berörts är tolkningen av bayesianska respektive klassiska konfidensintervall: Ett bayesianskt konfidensintervall uttrycker att en viss del av parameterns sannolikhetsmassa är samlad i intervallet. Motsvarande klassiska intervall innebär å andra sidan att om ett statistiskt försök upprepas i oändlighet så kommer

¹¹ Siu & Kelly (1998) använder begrepp som "internal notions", "intellectual knowledge" och "beliefs".

¹² En subjektiv sannolikhet får emellertid inte vara uttryck för personligt godtycke e dyl., utan måste vara konsistent med de sannolikheteoretiska grundsatser som finns samlade i Kolmogorovs axiomsystem (Leonard & Hsu 1999, s. 5). Kolmogorovs axiom finns återgivna i Blom (1984, s. 23).

¹³ Att dylika fördelningar kallas "objektiva" kan förklaras med att de saknar *reell* koppling till kunskap, dvs. att de inte bygger på någon faktisk persons gissning, tro, förväntan e dyl. En sådan fördelning skulle alltså kunna kallas *reellt objektiv* men *logiskt subjektiv*.

parameterskattningen att hamna i intervallet *en viss andel av gångerna*. Sannolikheten att parametern ligger inom intervallet med avseende på ett enskilt försök är emellertid 0 eller 1, beroende på om den *ligger där eller inte*. Ett klassiskt konfidensintervall uttrycker hur bra skattningen är (med utgångspunkt från tillgången på data), inte vad parametern har för värde. Ofta anses klassiska konfidensintervall sakna den intuitiva direktet som bayesianska intervall har. Utifrån ett bayesianskt intervall är det t.ex. möjligt att *acceptera en nollhypotes*. I klassisk statistik kan en nollhypotes bara förkastas eller *inte förkastas*. (Carlin & Louis 2000, s. 6ff)

Den största fördelen med bayesiansk statistik anses emellertid vara möjligheten att väga in subjektiv information, dvs. även sådan (erfarenhetsbaserad) information som kan skilja sig från individ till individ. Denna möjlighet är särskilt attraktiv i situationer med sparsamt eller ”brusigt” dataunderlag, vilket i sin tur är en av anledningarna till den bayesianska statistikens starka ställning inom kärnkraftsområdet. Andra anledningar är kopplingen mellan Bayes sats och beslutsteori (PSA betraktas ofta som en metod för beslutsfattande) [Holmberg], samt att fördelningar behövs för att fortplanta sannolikheter genom felträd (se kapitel 8). (Siu & Kelly 1998, s. 89f. Atwood et al. 2003, ch. 6, s. 2)

Traditionen att väga in data som inte är direkt grundade på observationer är likväl vad som renderat den skarpaste kritiken från det klassiska lägret (Carlin 2000, s. 6ff. Kaplan 1983, s. 1). Som svar på dylika invändningar menar Kaplan att subjektiva antaganden görs även i klassisk statistik, om än implicit. Enligt Kaplan är en styrka hos bayesiansk statistik att subjektiva antaganden görs ”openly and explicitly”. (Kaplan 1983, s. 1) Även Holmberg framhåller detta som en stor fördel eftersom antaganden då blir ”spårbara” [Holmberg]. Emellertid är det fortfarande motiverat att ifrågasätta tillförlitligheten och relevansen av sådana subjektiva data, dvs. vilka de är, vad de grundar sig på och hur de implementeras i modellen (mer om detta i avsnitt 5.2.2.1).

Det ska påpekas att klassisk och bayesiansk statistik används parallellt även inom kärnkraftsområdet. Det är framförallt i samband med PSA som den bayesianska metodiken är förhärskande. Klassiska metoder används för hypotesprövning, preliminära utvärderingar av data, test av modellantaganden samt för att avgöra vilken matematisk modell som ska användas (Atwood 2003, ch. 6, s. 2). Slutligen används ofta medelvärdet i den erhållna bayesianska fördelningen som punktskattning av parametern. En sådan skattning kan emellertid vara problematisk om fördelningen är mycket sned (lång svans) vilket är typiskt i fall med svagt datastöd. Det är inte ovanligt att medelvärdet överstiger både medianen och 95%-percentilen (Siu & Kelly 1998, s. 104). I sådana fall ger medelvärdet ingen information om var parametrarnas sannolikhetsmassa är samlad vilket betyder att bayesianska punktskattningar (t.ex. de som ges i T-boken) måste tolkas med försiktighet.

5 Bayesiansk statistik – den enkla modellen

I detta kapitel görs en mer ingående studie av den modell som presenterades i kapitel 3 och som kan skrivas

$$p(\lambda|x) \propto p(x|\lambda)p(\lambda) \quad (5.1)$$

Cooke et al. (1995, pt. 2, s. 11) kallar denna modell för ”the Simple Bayesian model”. En annan vanlig benämning är ”One-stage Bayes” (Cooke et al. 2003, Siu & Kelly 1998). Terminologin förklaras väsentligen av att *Bayes sats tillämpas en gång*. Jag kommer fortsättningsvis att tala om ”den enkla modellen”.

(5.1) har tre komponenter: *A priorifördelningen* $p(\lambda)$ beskriver fördelningen för den sökta parametern innan några observationer har gjorts. Observationerna beskrivs av *likelihoodfunktionen* $p(x|\lambda)$ med vars hjälp *a priorifördelningen* uppdateras till en *a posteriorifördelning* $p(\lambda|x)$. I det följande ska jag redogöra för olika sätt att modellera de två komponenterna i högerledet.

5.1 Likelihoodfunktionen

Likelihoodfunktionen $L(\lambda)$ definieras i det diskreta fallet som

$$L(\lambda) = p(x_1; \lambda) \cdot p(x_2; \lambda) \cdot \dots \cdot p(x_n; \lambda),$$

där x_i är utfall av en stokastisk variabel. Likelihoodfunktionen anger så tillvida sannolikheten för att ett visst *stickprov* ska erhållas givet parametern λ . Vidare är maximum-likelihoodskattningen (ML-skattningen) det parametervärde som maximerar likelihoodfunktionen med avseende på det erhållna stickprovet. (Blom & Holmquist 1998, s. 62)

Modelleringen av likelihoodfunktionen kräver för det första att den stokastiska variabeln X definieras. Detta avgör i sin tur vilket slags process som kan tänkas generera X . Om den stokastiska variabeln är diskret och utfaller i *diskret tid* används vanligen en Bernoulliprocess (Siu & Kelly 1998, s. 95f). Bernoulliprocessen räknar observationerna av en *binomialfördelad* stokastisk variabel t.o.m. det n :te försöket, t.ex. *antal fel* (diskret variabel) vid n *behov* (diskret tid). "Behov" ersätter i tekniska sammanhang det statistiska begreppet "försök" och avser aktiveringar, startförsök, tester e.dyl. Likelihoodfunktionen betecknas i Bernoullifallet $p(x|q)$ där modellparametern q anger en *sannolikhet*, t.ex. felsannolikheten hos en komponent. Modeller med binomialfördelad stokastisk variabel kommer jag ibland att kalla " q -modeller".

Bernoulliprocessens "motsvarighet" i kontinuerlig tid är Poissonprocessen (Blom 1984, s. 197, Råde & Westergren 1998, s. 417). Poissonprocessen räknar observationerna av en *Poissonfördelad* stokastisk variabel fram till tiden T : Antag att en händelse uppträder slumpmässigt med exponentialfördelade tidsavstånd; då är *antalet händelser* fram till T Poissonfördelat (Blom 1984, s. 256ff). Likelihoodfunktionen i Poissonfallet betecknas $p(x|\lambda)$ där modellparametern λ anger en *intensitet*, t.ex. felintensiteten hos en komponent. Modeller med Poissonfördelad stokastisk variabel kommer jag ibland att kalla " λ -modeller". (En sådan beskrivs i exempel 3.3).

Om den stokastiska variabeln är kontinuerlig, t.ex. *tiden* till den första händelsen, väljs lämpligen någon process med exponentialfördelad variabel och *livslängd* som modellparameter, t.ex. en Weibullprocess (Siu & Kelly 1998, s. 96).

Att avgöra vilken process som är lämplig i ett visst sammanhang är inte trivialt. De stokastiska processerna är mer eller mindre goda approximationer av verkligheten. Modellering av en lämplig likelihoodfunktion kräver således god kunskap om den *verkliga* processen och vilka förenklingar som kan göras utan att viktig information går förlorad. (Siu & Kelly s. 95ff) Bernoulliprocessen och Poissonprocessen är flitigt använda i PSA-sammanhang och fullt relevanta i fråga om T-boken eftersom observationerna där avser *antal behovsrelaterade respektive tidsrelaterade fel*, dvs. antal fel i diskret respektive kontinuerlig tid. I detta arbete betraktas dessa modellval som tillfredsställande. Vissa kommentarer görs i avsnitt 6.1.1, liksom i kapitel 13.

5.2 A priorifördelningen

Valet av a priorifördelning är den kanske mest kontroversiella aspekten av bayesiansk statistik. Till att börja med är det just användningen av a priorifördelningar som ger bayesiansk statistik det förmenta övertaget över klassisk i situationer med ett sparsamt dataunderlag. Att dessa fördelningar dessutom uttrycker subjektiva sannolikheter ställer speciella krav på trovärdighet. Dessa två faktorer tillsammans, att ett av metodikens främsta attribut likväl är en potentiell källa till misstro, är kanske det som mer än något annat givit upphov till "The Bayes-frequentist controversy" (Carlin & Louis 2000, s. 6).

I litteraturen delas a priorifördelningarna ofta in i två grupper; *informativa* respektive *icke-informativa*. Distinktionen är inte självklar (se avsnitt 5.2.3), men preliminärt kan sägas att icke-informativa fördelningar framställs på matematisk väg för att (idealt sett) återspegla *total avsaknad av kunskap* alternativt *total osäkerhet* om den sökta parametern. Informativa fördelningar ska å andra sidan ge uttryck för att den sökta parametern i någon mening är känd. Vidare kan en a priorifördelning, oavsett om den är informativ eller icke-informativ, vara *konjugerad* med en viss likelihoodfunktion. Konjugerade fördelningar är av särskild betydelse i bayesianska sammanhang.

5.2.1 Konjugerade fördelningar

En a priorifördelning är *konjugerad* med en viss likelihoodfunktion om a posteriorifördelningen tillhör samma fördelningsfamilj som a priorifördelningen. T.ex. är gammafördelningen konjugerad med Poissonfördelningen och betafördelningen med binomialfördelningen. Detta betyder att om a priorifördelningen är en gammafördelning och likelihoodfunktionen en Poissonfördelning så kommer även a posteriorifördelningen att vara en gammafördelning. (Jmf. Pörn 1990, s. 13, samt Carlin & Louis 2000, s. 25f)

Konjugerade fördelningar har en särställning inom bayesiansk statistik eftersom de gör uppdateringsprocessen analytiskt lätthanterlig. I gamma-Poissonfallet kan Bayes sats skrivas:

$$\begin{aligned} p(\lambda|x) &\propto (\lambda T)^x e^{-\lambda T} \cdot \lambda^{\alpha-1} e^{-\lambda\beta} \\ &\propto \lambda^{\alpha+x-1} e^{-\lambda(\beta+T)} \end{aligned} \quad (5.2)$$

eller enklare

$$\begin{aligned} p(\lambda|x) &\propto Po(\lambda) \cdot \Gamma(\alpha, \beta) \\ &\propto \Gamma(\alpha + x, \beta + T) \end{aligned} \quad (5.3)$$

Efter uppdatering erhålls alltså en ny gammafördelning med parametrarna $(\alpha + x)$ och $(\beta + T)$. Om a priorifördelningen är $\Gamma(0, 0)$ så kommer a posteriorifördelningen att bli $\Gamma(x, T)$. Av denna anledning kan parametrarna i a priorifördelningen antas beskriva "fiktiva observationer" (Cooke et al. 1995, pt. 2, s. 11).

5.2.2 Informativa a priorifördelningar

Tanken med en informativ a priorifördelning är att den ska återspegla befintlig kunskap om den sökta parametern. Den ska med andra ord ge uttryck för all *på förhand given information*. Vanligen ansätts någon standardfördelning som generellt antas kunna beskriva kunskap av det relevanta slaget. I tekniska sammanhang används ofta lognormal-, gamma- och betafördel-

ningarna (se appendix F.1). Den valda fördelningen kalibreras sedan på något sätt utifrån den befintliga kunskapen. Ett enkelt sätt att framställa en ”hyfsad” a priorifördelning är t.ex. att skatta de övre och undre konfidensgränserna, och anpassa en lognormalfördelning till dessa. (Siu & Kelly 1998, s. 98, 103f) En annan möjlighet är att skattningar av fördelningens procentiler får ligga till grund för något slags punktdiagram. En sådan *diskret fördelning* kan sedan interpoleras beroende på om uppdateringen ska göras i diskret eller kontinuerlig tid. (Atwood, 1986)

Uppdatering med empiriska data garanterar att a priorifördelningen modifieras så att ”degree of belief’ is rational, not merely personal opinion” (Atwood et al. 2003, ch. 6, s. 2). Ju större dataunderlag, desto snabbare konvergerar fördelningen mot den klassiska ML-skattningen x_i/T_i (vilket också kan uttryckas med att a posteriorifördelningen går mot en s.k. Dirac-funktion, dvs. en ”spik”). Den empiriska informationen kommer alltså successivt att ”dränka” den subjektiva informationen. Den avvikelse från ML-skattningen som i praktiken alltid återstår är rester av a prioriantagandet. (Jmf. Siu & Kelly 1998, s. 94f)

5.2.2.1 Problem vid härledning av informativa a priorifördelningar...

Oavsett vilken ansats som väljs i framställandet av en informativ a priorifördelning, så måste expertens uppfattningar transformeras till mätbara storheter. I någon mening krävs alltså en övergång från *kvalitativa till kvantitativa data*. Det finns en mängd etablerade metoder för att hantera denna övergång, även om meningarna går isär om vilken eller vilka som är att föredra (Siu & Kelly 1998, s. 103). Atwood kritiserar ansatsen i stort och menar att dylika metoder bara kan jämföras med avseende på nackdelar eftersom det saknas kriterier för vad som är korrekt. Istället slår han ett slag för klassisk metodik. (Atwood 1986, s. 148ff) Även Siu & Kelly väljer att fokusera på problemen snarare än de eventuella möjligheterna. Här följer några vanliga invändningar:¹⁴

1. Experter tenderar att lita för mycket på sin uppfattning vilket leder till för snäva fördelningar. (Siu & Kelly 1998, s. 104)
2. Risken finns att insamlade data beskriver en annan situation än den som ska analyseras. I kärnkraftssammanhang samlas data typiskt nog in under normal drift. Resultaten generaliseras sedan till att gälla även extraordinära situationer. Apostolakis et al. visar att även om experter ombeds skatta felintensiteter gällande för *extraordinära* situationer, så är dessa i regel underskattningar även i relation till det normala fallet, dvs. jämfört med den a posteriorifördelning som härleds med data från normala driftsituationer. (Apostolakis et al. 1980, s. 328)
3. Den eller de parametrar som ska skattas är ofta rent matematiska storheter som inte kan observeras. En möjlighet är att parametrarna relateras till egenskaper hos fördelningen som är mer intuitivt ”meningsfulla”, t.ex. medelvärdet eller felfaktorn (95%-percentilen dividerad med 5%-percentilen). Den subjektiva bedömningen av dessa kan sedan föras över på de aktuella parametrarna. Emellertid är även medelvärdet en abstrakt storhet utan någon uppenbar fysikaliskt mening, särskilt i fråga om starkt sneda fördelningar (som inte är helt ovanliga i dessa sammanhang). (Siu & Kelly 1998, s. 104)
4. A posteriorifördelningen kommer att vara nollvärd överallt där a priorifördelningen är det, oavsett tillgången på data. Observationerna kommer ju in i Bayes sats via likelihoodfunktionen som sedan multipliceras med a priorifördelningen. En a priorifördelning med ”kort svans” riskerar t.ex. att vara helt och hållet okänslig för avvikande

¹⁴ För mer ingående studier av de metoder som står till buds hänvisar jag till Winkler, R.L. and Hays, W.L., *Statistics: Probability, Inference and Decision*, 2nd edn. Holt, Rinehart and Winston, 1975.

data. Problemet är emellertid särskilt överhängande i fråga om diskreta a priorifördelningar. (Siu & Kelly 1998, s. 104)

Nämnda svårigheter kan tänkas uppstå när förväntningar, kunskap, gissningar etc. ska modelleras. Problemet är därvidlag hur a priorifördelningen ska kunna uttrycka *befintlig information* om parametern. En vanlig hållning är emellertid att det är *bristen på information*, dvs. *osäkerheten* som ska modelleras. Frågan blir då snarare hur mycket tilltro som ska sättas till den information som finns: Hur korrekt är klassificeringen av data? Är alla rapporterade ”kritiska fel” verkligen kritiska? Hur relevanta är expertutlåtanden och observationer i förhållande till den aktuella situationen? (Atwood et al. 2003 kap 6, s. 3f) Dessa frågor avspeglar en mer konservativ hållning som möjligen kan leda till något generösare a priorifördelningar. De principiella problemen beträffande övergången från kvalitativa till kvantitativa data är emellertid desamma.

5.2.2.2 ... och några möjliga lösningar

Siu och Kelly (1998, s. 98) anför ett par argument för att framställandet av en informativ a priorifördelning i praktiken sällan medför några större bekymmer. För det första är uppdateringen en iterativ process som bara i undantagsfall börjar från scratch; a priorifördelningen representerar kunskapen före de senaste observationerna, inte före *alla observationer*. I många fall är därför den informativa a priorifördelningen inget annat än en tidigare härledd a posteriorifördelning. För det andra: Ju fler observationer som kan förväntas, desto mindre precis behöver a priorifördelningen vara. En mycket ungefärlig representation av befintlig kunskap kan vara tillräcklig. (Siu & Kelly 1998, s. 98) I vissa fall kan det vara motiverat att ansätta en fördelning som inte ger någon information alls och s.a.s. ”låta data tala för sig själva”. Detta är tanken med s.k. icke-informativa a priorifördelningar.

5.2.3 Icke-informativa a priorifördelningar

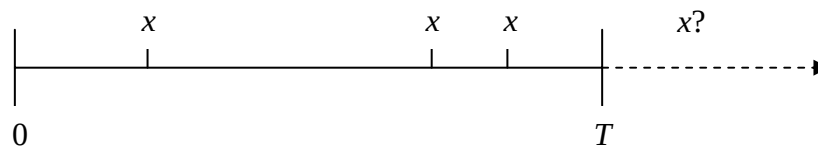
Om den förväntade tillgången på data är stor samtidigt som ingen uppenbar informativ a priorifördelning finns tillhands kan det vara lämpligt att ansätta en icke-informativ a priorifördelning (Atwood et al. 2003 kap 6, s. 14). En icke-informativ a priorifördelning ska idealt sett uttrycka *total osäkerhet* om den aktuella parametern, dvs. att alla parametervärden är lika sannolika. Emellertid finns inget entydigt sätt att framställa en fördelning som i alla avseenden uppfyller detta krav (Pörn 1990, s. 61). En likformig fördelning, det intuitivt sett mest tilltalande alternativet, är t.ex. *informativ* om den logaritmeras. Detta strider mot ett annat intuitivt krav nämligen att *om kunskap saknas om parametern θ så saknas också kunskap om varje transformation av θ* . Alltså måste ”icke-informativitet” definieras på något annat sätt. (Siu & Kelly 1998, s. 105)

Box och Tiao (1973) föreslår att en icke-informativ fördelning inte måste uttrycka okunskap i absolut mening utan istället ”[---] an amount of prior information which is small relative to what the particular projected experiment can be expected to provide” (Box & Tiao 1973, citat i Pörn 1990, s. 17). Box och Tiao ställer sig därmed frågan: *I vilken transformation av θ ska fördelningen vara likformig?* Resultatet kallas principen om ”data-translated likelihood”. I Poissonfallet leder den till följande a priorifördelning (Siu & Kelly 1998, s. 106):

$$p(\lambda) \propto \lambda^{-\frac{1}{2}} \quad (5.4)$$

dvs. en gammafördelning med $\alpha = \frac{1}{2}$ och $\beta = 0$ (Atwood et al. 2003, ch. 6, s. 14). Principen om data-translated likelihood är flitigt använd inom bayesiansk statistik. En annan vanlig princip är *Jeffreys regel* för framställning av approximativt "data-translated likelihoods". Även Jeffreys regel ger i Poissonfallet $p(\lambda) \propto \lambda^{-\frac{1}{2}}$. (Siu & Kelly 1998, s. 107)

En heuristisk tolkning av dessa resultat är följande (gamma-Poissonfallet förutsatt): Om processen observeras fram till tiden T (s.k. tidstrunkering) och nästa fel inträffar vid tiden $T+$ kan detta kompenseras med att $\alpha = 1$, dvs. ett fiktivt fel som ersättning för det "missade". Om felet istället inträffar vid tiden $T-$ behövs ingen kompensation, dvs. $\alpha = 0$. Detta motsvarar situationen vid "feltrunkering", dvs. när processen observeras t.o.m. att ett fel inträffar. Vid tidstrunkering är bägge dessa situationer orealistiska; det verkar extremt pessimistiskt att inleda observationen med ett fiktivt fel, men lika optimistiskt att inleda utan något fel alls. Som rimligt alternativ framstår istället medelvärdet av dessa extremer, nämligen $\alpha = \frac{1}{2}$. (Vaurio & Jänkälä 2006, s. 210f) Att inget fel observerats betyder alltså inte att felintensiteten är noll utan att felet "missades". En rimlig gissning är då att $x = \frac{1}{2}$, vilket likväl är konsistent med Jeffreys regel.



Figur 2. Jeffreys regel kan anses lösa problemet med "missade" fel vid tidstrunkering.

Ett problem med icke-informativa fördelningar är att de i allmänhet är *oäkta* dvs. integralen är oändlig (i gammafallet då $\alpha < 1$). Emellertid räcker det med en enda observation för att fördelningen ska bli äkta: $\Gamma(0, 0)$ är oäkta men $\Gamma(1, T)$ är äkta. Enligt Cooke et al. berättigar detta faktum användandet av icke-informativa fördelningar i samband med *den enkla modellen*: "With the simple Bayesian model, improper priors are justified by the fact that the improper priors become proper after updating on one failure." (Cooke et al. 1995, s. 11)

I kärnkraftssammanhang finns emellertid inga skäl att vänta sig annat än ytterst sporadiska observationer, åtminstone inte då det gäller händelser som är relevanta för T-boken. Att ansätta en icke-informativ a priorifördelning enligt ovan är därför knappast någon lösning. Målet måste istället vara att framställa *en så informativ a priorifördelning som möjligt*. Detta görs i praktiken genom utnyttjande av s.k. *generisk information*, dvs. information från en superpopulation (en överordnad population). Detta innebär likväl att den enkla modellen måste överges.

6 Bayesianska metoder inom kärnkraftsområdet

6.1 Inledning

Typiskt i kärnkraftssammanhang är bristen på empiriskt underlag. Det anses därför nödvändigt att komplettera den specifika informationen med data från liknande, om än inte identiska,

källor (Atwood et al. 2003, kap 8, s. 1). Metoder för insamlandet av sådan information har följt två olika utvecklingslinjer: *Hierarkisk Bayes* och *Parametric Empirical Bayes* (PEB). Bägge metodgenrerna bygger på den enkla modellen och antas lösa problemet med hur en informativ a priorifördelning ska framställas. Detta görs emellertid på radikalt olika sätt: I hierarkisk Bayes utvecklas a priorifördelningen genom att Bayes sats ansätts på flera hierarkiskt ordnade nivåer. På varje nivå härleds en a priorifördelning till nivån under. I praktiken förekommer hierarkisk Bayes bara i sin ”tvåstegsvariant”, dvs. Bayes sats ansätts två gånger; en för generella, s.k. *generiska data*, och en för specifika data (Siu & Kelly 1998, s. 100). Den gängse benämningen på denna variant är *Two-stage Bayes*. Metoden refereras ofta till som en *superpopulationsmetod* eftersom de två nivåerna antas motsvara ”grupp” respektive ”individ” (Carlin & Louis 2000, Vaurio & Jänkälä 2006). *Two-stage Bayes* är idag den helt dominerande approachen för sammanställning av data från flera källor (se t.ex. Vaurio & Jänkälä 2006, s. 209).

Även PEB är ett slags superpopulationsmetod med två nivåer. Skillnaden är att a priorifördelningen framställs med hjälp av *klassisk metodik*. Bayes sats uppträder alltså bara en gång.

I det följande kommer bägge dessa metodgenrer att ges en närmare beskrivning. Framställningen av hierarkisk Bayes kommer att begränsas till ”tvåstegsfallet”. Först ska jag emellertid redogöra för några vanliga antaganden som görs vid parameterskattning inom kärnkraftsområdet. Jag kommer också att ta upp den viktiga frågan om hur data ska grupperas.

6.1.1 Vanliga antaganden

De mest grundläggande antagandena rör dels den stokastiska variabeln, dels intensitetsfunktionen för den skattade parametern. Till att börja med antas observationerna så gott som alltid komma från en Poissonprocess vid tidsrelaterade fel, och en Bernoulliprocess vid behovsrelaterade fel. Det vanliga är också att processen *tidstrunkeras*, dvs. att observationen pågår fram till en viss tidpunkt T . Drifttiden betraktas såtillvida som *deterministisk och mätbar* (jmf. Cooke et al. 2003, s. 4, samt Carlin & Louis 2000, s. 17).

Olika antaganden beträffande den skattade parametrarnas intensitetsfunktion kan komma ifråga i samband med tidsrelaterade fel. Det vanligaste är att intensiteten antas vara konstant, dvs. $\lambda(t) = \lambda$. Den ”badkarskurva” som ofta används för att modellera utvecklingen av en komponents felbenägenhet används alltså inte i samband med superpopulationsmetoder. Det finns väsentligen två motiv till detta: För det första gör redan mycket enkla tidsberoenden att beräkningarna avsevärt kompliceras. För det andra kan komponenterna antas vara driftsatta under badkarskurvans flacka (konstanta) fas. Antagandet stöds av att komponenter vanligen testkörs innan de tas i drift, vilket eliminerar s.k. ”barnsjukdomar”, samt att de tas ur drift i god tid innan de börjar åldras.

En konsekvens av den konstanta felintensiteten är att det inte spelar någon roll *när* en komponent byts ut eller tas i drift. En komponent kan såtillvida betraktas som identisk med sin ”ersättare”. Strängt taget är det alltså inte själva komponenten som modelleras utan snarare dess ”plats” (sockel, uttag, hylsa e dyl.). (Cooke et al. 1995, pt. 1) Vanligen antas den enda relevanta informationen om en komponents felintensitet vara *antalet fel och den totala drifttiden* (alt. antalet aktiveringar). Bortfall pga. förebyggande underhåll, reparation av smärre fel och liknande antas såtillvida inte ge någon information om felintensiteten (Cooke et al. 1995, s. 4ff).

6.1.2 Gruppering av data

Frågan om hur data ska grupperas bestäms utifrån antaganden om var den väsentliga *variationen* finns. För det första antas komponenter av "samma sort" vara besläktade med avseende på felintensitet; en backventil antas ha mer gemensamt med andra backventiler än med t.ex. skalventiler. För det andra antas i regel den väsentliga variationen finnas *mellan anläggningar*. Superpopulationen utgörs således av en *grupp av anläggningar* medan delpopulationen utgörs av *komponenter på en specifik anläggning*. Detta betyder att komponenter av samma slag (t.ex. en grupp av backventiler) antas ha samma felintensitet och således kan grupperas (felintensiteten tas då som summan av alla fel dividerat med den totala drifttiden). Pörn hör till dem som avviker från denna konvention genom att han antar en variation *mellan individer*: "[O]ur basic model assumes a variability between the individual components [...] unlike many other applications which are based on a variability between the plants." (Pörn 1990, s. 100) Något förenklat kan Pörns superpopulation sägas bestå av alla komponenter på alla anläggningar medan delpopulationen består av en enskild komponent. Variationen kan i sin tur antas bero på:

- Fysikaliska skillnader.
- Skillnader i omgivning.
- Skillnader i underhållsstrategier och arbetslag.
- Orealistiska statistiska modeller.

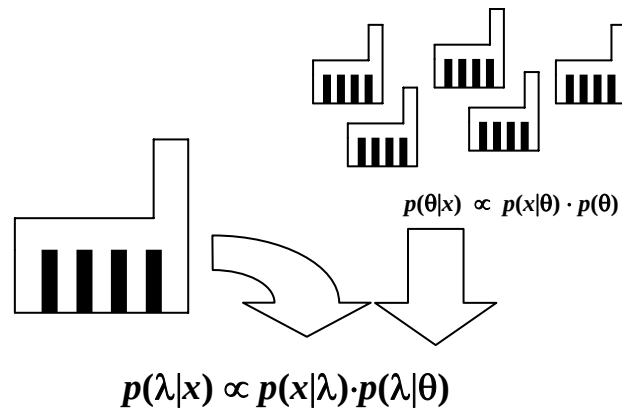
(Nyman 1995)

Andra vanliga antaganden rör det interna beroendet mellan den stokastiska variabeln och de parametrar som ingår i tvåstegsmodellen. Dessa tas upp i avsnitt 6.2.1.

6.2 Tvåstegs Bayes

1983 publicerades Stan Kaplans *On a "Two-stage" Bayesian Procedure for Determining Failure Rates from Experimental Data*. Han var därmed först med att föreslå en "hierarkisk" metod för parameterskattning i PSA-sammanhang (Cooke et al. 2003, s. 2. Heising & Shwayri 1987, s. 24). Tidigare hade han, tillsammans med bl.a. Apostolakis, propagerat för tillämpning av bayesianska metoder i liknande sammanhang, men då var det snarast tal om den enkla modellen (Apostolakis et al. 1980. Siu & Kelly 1998, s. 1). Kaplans metod har sedan slutet av 80-talet vidareutvecklats av bl.a. Pörn, Iman och Hora (Cooke et al. 2003, s. 2). Tvåstegsmetodiken är idag helt dominerande för härledning av parametrar av det slag som presenteras i T-boken. Förutom i Sverige och Finland används metodiken bl.a. i Tyskland, USA och Frankrike (T6, s. 31).

Kaplan presenterar sin metod som en lösning på problemet med att bestämma en a priorifördelning vid konventionell tillämpning av Bayes sats (Kaplan 1983, s. 1). Idén är att a priorifördelningen i den enkla modellen skrivs som a posteriorifördelning i en "överordnad" bayesiansk relation. På denna övre nivå utnyttjas data från en superpopulation för härledning av a priorifördelningens *parametrar*. I Kaplans exempel används en lognormalfördelning som a priorifördelning vilket betyder att osäkerheten överförs till parametervektorn $\theta = (\mu, \sigma)$. Detta betyder likväl att problemet med hur a priorifördelningen ska bestämmas förskjuts uppåt i hierarkin eftersom även θ antas ha en fördelning $p(\theta)$. Kaplans poäng är nu att denna fördelning kan bestämmas på grundval av vagare information eftersom tillgången på data är större i superpopulationen än i det specifika fallet (Kaplan 1983, s. 1ff). Ett försök till illustration av superpopulationsprincipen i tvåstegs bayesiansk analys görs i figur 3.



Figur 3. Tvåstegsmodellen: Hyperparametern θ betingas på data från "andra källor" och betingar i sin tur λ som uppdateras med data från en specifik källa.

Kaplans tvåstegsmodell betraktas ofta som ett specialfall av hierarkisk Bayes (Siu & Kelly 1998, s. 100. Carlin & Louis 2000, s. 57). Atwood et al. (2003, ch. 8, s. 20f) hävdar emellertid att det finns en väsentlig begreppslig skillnad mellan Kaplans *Two-stage Bayes* och *tvåstegsvarianten av hierarkisk Bayes*. Den förmenta skillnaden är relaterad till frågan om huruvida individspecifika data ska ingå i superpopulationen eller inte. Problemet är närmare bestämt följande: Om data från det aktuella systemet räknas in i populationsdata så kommer samma data att uppträda i likelihoodfunktionen på bägge nivåerna, dvs. både i $p(x|\lambda)$ och $p(x|\theta)$ (jmf. figur 3). Enligt den bayesianska traditionen är sådan *dubbelräkning* felaktig (detta kommer att diskuteras i flera av de återstående kapitlen). Kaplan löser problemet genom att explicit utesluta systemspecifika data ur populationsdata (Kaplan 1983, s. 2). Atwood et al. (2003) hävdar emellertid att Kaplan tvingas till detta eftersom hans modell inte är korrekt baserad på Bayes sats. Av satsen följer nämligen *logiskt* att data *inte dubbelräknas*. Atwood et al. förespråkar istället *tvåstegsvarianten av hierarkisk Bayes*: "[T]he hierarchical Bayes method is based directly on Bayes' Theorem, and therefore does not involve double counting. Therefore, the two-stage Bayesian method should no longer be used, but should be replaced by the conceptually cleaner hierarchical Bayes method." (Atwood et al. 2003, ch. 8, s. 20) Skillnaden är med andra ord att Kaplan framställer dubbelräkningen som i någon mening valfri (detta gör även Siu och Kelly 1998, s. 100) medan Atwood et al. menar att den är *logiskt oförenlig* med bayesiansk metodik. Denna konsekvens av Bayes sats explicitgörs bland annat hos Pörn (se avsnitt 7.1.1.2).

Atwood et al. (2003) förkastar alltså Kaplans *Two-stage Bayes* till förmån för hierarkisk Bayes. För att undvika terminologiskt och begreppsligt trassel kommer jag i fortsättningen att använda begreppet *tvåstegsmetod* (-modell etc.) och därmed avse tvåstegsvarianten av hierarkisk Bayes såsom den framställs av Cooke et al. (1995), Cooke et al. (2003), Carlin & Louis (2000) och Pörn (1990).

6.2.1 Tvåstegsmodellen

Tvåstegsmodellen ger svar på frågan om hur den enkla modellens a priorifördelning ska härledas. Den enkla modellen kan skrivas på formen

$$p(\lambda|x) \propto p(x|\lambda)p(\lambda). \quad (6.1)$$

Vad som söks är alltså $p(\lambda)$. Den enkla modellen med avseende på θ kan skrivas

$$p(\theta|\lambda) \propto p(\lambda|\theta)p(\theta) \quad (6.2)$$

helt analogt med (6.1). Marginalfördelningen i (6.2) är

$$p(\lambda) = \int p(\lambda|\theta)p(\theta)d\theta \quad (6.3)$$

(jmf. Pörn 1990, s. 47). Insättning av (6.3) i (6.1) ger nu

$$p(\lambda|x) \propto p(x|\lambda) \int p(\lambda|\theta)p(\theta)d\theta, \quad (6.4)$$

vilket är en *principiell* representation av tvåstegsmodellen. I (6.4) saknas emellertid en viktig detalj. För att data från superpopulationen ska kunna utnyttjas måste θ vara betingad på x . Det är emellertid inte "tillåtet" att sätta $p(\theta) = p(\theta|x)$ i (6.4). Då är det nämligen en öppen fråga huruvida data dubbelräknas. I själva verket måste redan (6.3) involvera data. Den "korrekta" versionen av (6.3) är såtillevda

$$p(\lambda|x) = \int p(\lambda|x, \theta)p(\theta|x)d\theta \quad (6.5)$$

där $x = \{x_1, \dots, x_{n+1}\}$ dvs. data för *samtliga anläggningar* (jmf. Pörn 1990, s. 11). Utifrån (6.5) kan tvåstegsmodellen härledas till

$$p(\lambda_{n+1}|x_1, \dots, x_{n+1}) \propto p(x_{n+1}|\lambda_{n+1}) \int p(\lambda_{n+1}|\theta)p(\theta|x_1, \dots, x_n)d\theta \quad (6.6)$$

där

$$p(\theta|x_1, \dots, x_n) \propto p(x_1, \dots, x_n|\theta)p(\theta) \quad (6.7)$$

(Pörn 1990, s. 11f. Cooke et al. 2003, s. 5f). Observera att specifika data x_{n+1} nu separerats från generiska data $\{x_1, \dots, x_n\}$. A priorifördelningen, dvs. integralen i (6.6), baseras bara på generiska data och uppdateras med specifika data som kommer in via likelihoodfunktionen. Att data separeras på detta sätt är en följd av att

- $\{\lambda_i\}$ är betingat oberoende och likafördelade givet θ .
- $\{x_i\}$ är betingat oberoende givet $\{\lambda_i\}$ och θ .
- x_i är betingat oberoende av θ och alla λ utom λ_i .

(Pörn 1990, s. 10. Cooke et al. 1995, pt 2, s. 5). Dessa villkor är vedertagna i så gott som alla tvåstegsmetoder (Cooke et al. 2003, s. 5).

I hierarkisk Bayes tolkas nu θ som en multiparametrisk vektor av godtycklig ordning. I tvåstegsmodellen tolkas θ som en vektor innehållande parametrarna i $p(\lambda)$. Vilka dessa är beror

på vilken fördelning som väljs för att representera $p(\lambda)$. Gammafördelningen har $\theta = (\alpha, \beta)$ och lognormalfördelningen $\theta = (\mu, \sigma)$.

Anmärkning: Begreppet ”steg” kan vara något missvisande eftersom det lätt associeras till något slags serie, iteration e dyl. Emellertid innebär tvåstegsmetodiken att Bayes sats *nästlas* snarare än itereras. Det kunde därför vara lämpligare att tala om *nivåer* istället för steg. Att jag ändå valt begreppet ”tvåstegs-” beror för det första på att jag inte hittat något smidigt alternativ på svenska; begrepp som ”tvåvånings-”, ”tvånivå-”, ”tvåplans-” o dyl. ter sig ganska otympliga och saknar dessutom förankring i matematiskt språkbruk. För det andra är begreppet ”tvåstegs-” den enda någotsånär etablerade svenska översättningen av ”two-stage” genom att Pörn använder den i T-boken. När jag anser det passande kommer jag emellertid att bryta denna konvention och använda ”nivåer” tillsammans med ”tvåstegs-”.

6.2.2 Prior och hyperprior

Tvästegsmodellen har a priorifördelningar på två nivåer. I engelskspråkiga texter används begreppen ”prior” respektive ”hyperprior” för att skilja dessa åt. På svenska finns emellertid ingen etablerad vokabulär. Jag kommer därför att använda de engelska uttrycken om inte risken för sammanblandning på annat sätt är undanröjd (t.ex. i samband med den enkla modellen). Begreppet ”a priorifördelning” kommer även fortsättningsvis att användas som samlingsnamn. På motsvarande sätt kommer jag ibland att använda begreppet ”hyperparameter” om θ och ”hyperposteriorifördelning” för den fördelning som beräknas i steg ett (motsvarande (6.7)).¹⁵ Jag kommer också att referera till de två ”stegen” som *steg ett* och *steg två* där steg två motsvarar den enkla modellen. Hyperprioriorn $p(\theta)$ ansätts alltså i steg ett och priorn $p(\lambda)$ i steg två.

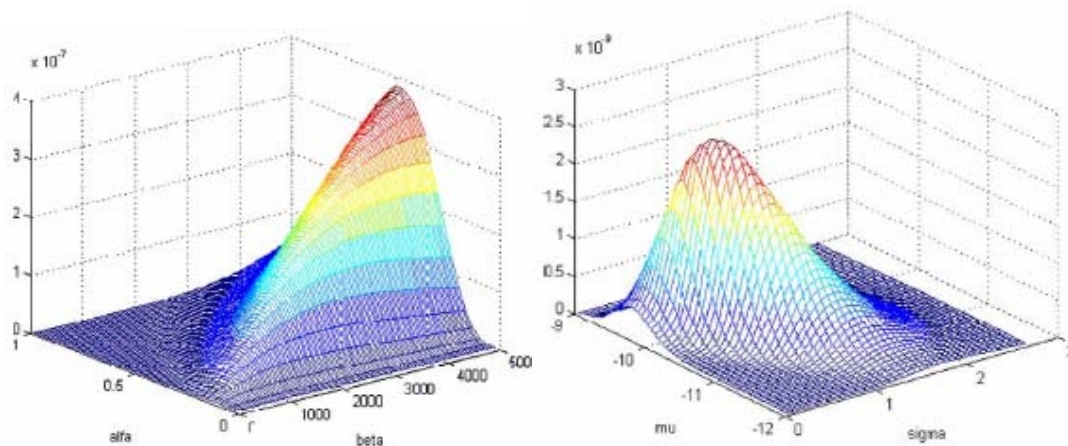
Som prior för λ väljs i allmänhet en familj av gamma- eller lognormalfördelningar. Gammafördelningen har fördelen av att vara konjugerad med Poissonfördelningen vilken så gott som alltid används som likelihoodfunktion för tidsrelaterade fel. För behovsrelaterade fel, med felsannolikheten q som modellparameter, väljs ibland en betafördelning tillsammans med en binomialfördelning (Bernoulliprocess). Även en *trunkerad lognormalfördelning* kan användas tillsammans med binomialfördelningen för ”mycket små q ” (Siu & Kelly 1998, s. 99).

I tvåstegsfallet tolkas ibland priorn som en s.k. *population variability curve* (PVC). Priorn antas såtillvida beskriva hur parametern varierar i superpopulationen. Denna tolkning görs av bl.a. Kaplan (1983) och Siu & Kelly (1998). Alternativet är att i första hand tolka priorn som en *osäkerhetsfördelning* där osäkerheten i sin tur kan bero av många olika faktorer (däribland populationsvariabiliteten). Denna tolkning företräds av Pörn (1990).

Hyperprioriorn $p(\theta)$ är så gott som alltid en icke-informativ fördelning (Siu & Kelly 1998). För det första är θ en mycket abstrakt storhet som inte utan vidare kan tilldelas erfarenhetsbaserade sannolikheter (jmf. problempunkt 3 i avsnitt 5.2.2.1). För det andra är själva tanken med tvåstegsmodellen att superpopulationen ska ge ett tillräckligt empiriskt underlag för att data ska kunna tillåtas ”tala för sig själva”. Framställningen av hyperprioriorn görs således på matematisk väg, företrädesvis genom tillämpning av Jeffreys regel. Ett problem är att hyperparametern är tvådimensionell i både gamma- och lognormalfallet och att den ”multiparametriska” versionen av Jeffreys regel är svår att tillämpa (Pörn 1990, s. 21). Box & Tiao föreslår istället att den enparametriska versionen tillämpas på parametrarna var för sig (Cooke et al. 2003, 19ff). Emellertid förutsätter detta att hyperparametrarna är oberoende (Pörn 1990, s. 21) vilket torde ge en viss fördel i fallet med en lognormalfördelad prior.

¹⁵ Terminologin är i överensstämmelse med Cooke et al. (1995, s. 12).

Liksom i fallet med icke-informativa a priorifördelningar (i den enkla modellen) är icke-informativa hyperpriors i allmänhet oäkta. Skillnaden gentemot den enkla modellen är att de inte blir äkta efter uppdatering med en observation. Cooke et al. (2003) visar att det i själva verket inte finns någon garanti för att de blir äkta ens efter oändligt många observationer (Cooke et al. 2003, s. 6f, 37. Jmf. Carlin & Louis 2000, s. 29). Icke-informativa hyperpriors måste därför trunkeras. Ett resultat hos Cooke et al. är vidare att denna trunkering är enklare med lognormalfördelad än med gammafördelad prior: "The lognormal model [...] enjoys a significant advantage over the gamma model in that, as observation time increases, a natural truncation of the hyperpriors μ, σ is possible." (Cooke et al. 2003, s. 37) Trunkeringen av hyperpriorerna görs med utgångspunkt från hyperposteriorifördelningen eftersom denna visar var någonstans "datastödet" är stort. Hyperposteriorifördelningen (6.7) beror i sin tur av likelihoodfunktionen $p(x|\theta)$. Cooke et al. visar nu att likelihoodfunktionen för gammafördelningen "does not peak but 'ridges'", dvs. likelihoodfunktionen i α och β har inget tydligt maximum (Fig. 4, vänstra bilden) (Cooke et al. 2003, s. 16f). Det finns alltså inte något naturligt sätt att avgränsa sannolikhetsmassan i hyperposteriorifördelningen och därmed inte heller i a posteriorifördelningen: "The inability to localize the hyperposterior mass for (α, β) means that we cannot localize the posterior mass [...]" (Cooke et al. 2003, s. 15). Likelihoodfunktionen för μ och σ , dvs. hyperparametrarna i lognormalfördelningen, bildar däremot ett tydligt maximum då T växer (Fig. 4, högra bilden) (Cooke et al. 2003, s. 22). Skillnaden torde bero på att beroendet är starkare mellan α och β än mellan μ och σ (Cooke et al. 1995).



Figur 4. Likelihoodfunktionen $p(x|\theta)$ i gammafallet (t.v.) och i lognormalfallet (t.h.) (Cooke et al. 2003, s. 16, 22).

Att en "naturligare" trunkering (en naturlig avgränsning av sannolikhetsmassan) kan göras i lognormalfallet betyder enligt Cooke et al. att skillnaden mellan olika godtagbara trunkeringar är mycket liten (Cooke et al. 2003, s. 25). I gammafallet ger å andra sidan olika "lika plausibla" trunkeringar skillnader på en faktor ≥ 5 (Cooke et al. 2003, s. 18). Detta ger enligt Cooke et al. en betydande fördel för lognormalfördelningen: "The possibility of truncating the domains of integration so as to include the bulk of the mass around the maximum of the prior is a significant argument in favour of the lognormal prior over the gamma prior." (Cooke et al. 2003, s. 25) Cooke et al. ställer sig emellertid tveksamma till att över huvud taget använda icke-informativa hyperpriors om de leder till oäkta fördelningar (Cooke et al. 2003, s. 37).

6.3 Parametric Empirical Bayes (PEB)

Tänkbara alternativ till tvåstegsmetodiken är olika former av *Parametric Empirical Bayes* (PEB). Dessa användes bl.a. i de första två upplagorna av T-boken (T1, T2). Även PEB bygger på ett slags superpopulationsprincip. Skillnaden är enkelt uttryckt att den bayesianska skattningen av hyperparametrarna ersätts med klassiska punktskattningar. Steg två är sedan identiskt med den enkla modellen. PEB betraktas därför ofta som ett slags blandmetodik: "[I]t is a kind of hybrid, involving a non-Bayesian step followed by a Bayesian step" (Atwood et al. 2003, ch. 8, s. 3).

Metodiken har fått sitt namn utifrån att a priorifördelningen inte bygger på subjektiv information. Istället skattas a priorifördelningen utifrån *existerande* data och bestäms alltså *empiriskt* (Atwood et al. 2003, ch. 8, s. 3). Pörn uttrycker detta med orden: "[The] terminology derives from the fact that the prior distribution of the primary parameter $[\lambda]$ is specified to its form, the (secondary) parameters $[\theta]$ of which are estimated on the basis of empirical evidence and sampling theory models." (Pörn 1990, s. 2)

Liksom i tvåstegsmodellen kan a priorifördelningen i princip betraktas som antingen en PVC eller en osäkerhetsfördelning. Skillnaden är att eventuell osäkerhet beträffande hyperparametrarna ignoreras. Därmed finns ingen osäkerhet beträffande PVC i sig, men PVC kan tolkas som en osäkerhetsfördelning för de specifika fördelningarna inom populationen. Typiskt för PEB är också att hyperparametrarna skattas ur den stokastiska variabelns *marginalfördelning* $p(x) = \int p(x|\theta)p(\theta)d\theta$ (Carlin & Louis 2000, s. 31). Fördelningen för respektive individ kan såtillvida betraktas som ett stickprov ur en gemensam fördelning. Att a priorifördelningen framställs ur marginalfördelningen $p(x)$ betyder vidare att individspecifika data *dubbelräknas*. Denna dubbelräkning är omdiskuterad eftersom den av många anses strida mot bayesianska principer (Carlin & Louis 200, s. 31f).

Traditionellt särskiljs *parametriska* empiriska modeller (PEB) från *icke-parametriska* empiriska modeller (NPEB). Bägge är s.k. EB-metoder (*Empirical Bayes*). En fördel med de senare är att a priorifördelningen inte måste antas tillhöra en specifik familj av fördelningar. I PSA-sammanhang är det framförallt PEB som tillämpas även om det finns sentida exempel på tillämpning av icke-parametriska modeller (Siu & Kelly s. 101). I standardverk om EB-metoder behandlas NPEB emellertid mycket översiktligt.

6.3.1 Skattningsmetoder

För skattning av hyperparametrarna används vanligen *momentmetoden* eller *ML-metoden* (maximum-likelihood). Med ML-metoden väljs den parameter θ som maximerar likelihood-funktionen $p(x|\theta)$.

Exempel 6.1: Antag att x är en observation från en Poissonprocess. Då är likelihoodfunktionen $L(\lambda) = p(x|\lambda)$. Med en gammafördelning $\Gamma(\lambda|\alpha, \beta)$ som a priorifördelning blir marginalfördelningen $p(x|\alpha, \beta) = \int p(x|\lambda)\Gamma(\lambda|\alpha, \beta)d\lambda$, dvs. en gamma-Poissonfördelning eller *negativ binomialfördelning* (i själva verket medelvärde av x givet alla tänkbara λ). $p(x|\alpha, \beta)$ kan därmed skrivas

$$p(x|\alpha, \beta) = \frac{\Gamma(\alpha + x)}{\Gamma(x + 1)\Gamma(\alpha)} \left(\frac{\beta}{\beta + T} \right)^\alpha \left(\frac{T}{\beta + T} \right)^x.$$

ML-skattningarna är nu de α och β som maximerar $p(x|\alpha, \beta)$. De löses ut ur

$$\frac{\partial}{\partial \theta} \ln[p(x|\theta)] = 0,$$

där $\theta = (\alpha, \beta)$.

(Atwood et al. 2003, ch. 8, s. 3ff)

□

Med momentmetoden löses parametrarna istället ut ur ett antal ekvationer (lika många som antalet parametrar) som erhålls genom att de *teoretiska momenten* (medelvärde, varians etc.) sätts lika med motsvarande stickprovsmoment. Den engelska benämningen är *moment matching method*.

Exempel 6.2: Detta exempel illustrerar principen för momentmetoden. Här används a priorifördelningen för skattning av hyperparametrarna (s.k. Prior Moment Matching Method). I praktiken används i allmänhet marginalfördelningen.

Antag en betafördelning som prior och att x är observationer från en Bernoulliprocess (n = antal behov). I följande ekvationer utnyttjas betafördelningens medelvärde (1:a moment) och varians (2:a moment) för bestämning av hyperparametrarna α och β . Koefficienten $1/(m-1)$ ger en s.k. "unbiased" variansskattning:

$$\frac{p}{p+q} = \bar{\phi}$$

$$\frac{pq}{(p+q)^2(p+q+1)} = \frac{1}{m-1} \sum_{i=1}^m (\phi_i - \bar{\phi})^2$$

där

$$\bar{\phi} = \frac{1}{m} \sum_{i=1}^m \phi_i.$$

ϕ_i skattas utifrån tillgängliga data dvs.

$$\phi_i \approx \frac{x_i}{n_i}.$$

(Exemplet är baserat på Siu & Kelly 1998, s. 101).

□

Skattningen av hyperparametrarna är numeriskt sett enklare med momentmetoden än med ML-metoden eftersom det i praktiken kan vara svårt att hitta maximum av den logaritmerade likelihoodfunktionen. ML-skattningarna är emellertid oftast bättre.¹⁶ Ett alternativ är därför att utnyttja momentmetodens enkelhet för att härleda *utgångspunkter* för de numeriska metoder som används för att beräkna ML-skattningarna. (Siu & Kelly 1998, s. 109)

6.3.2 Problem med PEB

PEB-metoder medför vissa svårigheter varav jag här tar upp några av de viktigaste.

- Ingen av de ovan föreslagna metoderna *garanterar* rimliga parameterskattningar. ML-skattningen existerar t.ex. inte då $x_i = 0$. Konvergensproblem kan också uppstå om li-

¹⁶ Vid stora stickprov är ML-skattningarna *alltid* bättre.

likelihoodfunktionen är komplicerad. Dessutom kan momentskattningarna bli negativa även om modellerna kräver att de är strikt positiva. (Siu & Kelly 1998, s. 102)

- PEB tenderar att ge snävare a priorifördelningar än hierarkiska modeller vilket kan hänföras till att inga osäkerheter vägs in i hyperparametrarna. Detta kan innebära att variabiliteten i superpopulationen underskattas: "The method [...] underestimates the uncertainty in the answers, because it treats [$p^*(\lambda)$] as if it were equal to the true distribution [...]." (Atwood et al. 2003, ch. 8, s. 3) Den generiska fördelningen kan i värsta fall bli så snäv att uppdateringen blir verkningslös även med stora mängder individspecifika data. Anledningen till detta är att en fördelning som är okänslig för variabilitet inom populationen koncentrerar sig kring populationens medelvärde. Stora mängder generiska data ger mycket precisa skattningar av medelvärdet vilket ger upphov till en "toppig" fördelning. Därmed inte sagt att detta medelvärde är representativt för den aktuella individen (eller undergruppen). Problemet är med andra ord att en eventuell avvikelse aldrig kan komma till uttryck. (Siu & Kelly s. 104)
- Även den s.k. dubbelräkningen leder till att a priorifördelningen blir snäv eller "overconfident". Carlin & Louis kallar PEB-metoder som inte uppmärksammar detta för "naive EB methods". (Carlin & Louis 2000, s. 32)

Det finns olika sätt att komma tillrätta med dessa problem. Vaurio (Vaurio & Jänkälä 2006) har t.ex. utvecklat en PEB-metod baserad på momentmetoden som "tolererar" att medelvärde och varians är nollvärda. Det finns vidare en mängd olika metoder för att väga in osäkerhet i efterhand (i steg två) för att på så sätt göra PEB mer lik hierarkisk Bayes. Sådana metoder behandlas kortfattat i Atwood et al. 2003 (kap. 8, s. 6) och mer ingående i Carlin & Louis (2000, ch. 3). Svårigheterna ovan kommenteras ytterligare i samband med Pörns metod (avsnitt 7.1.1.1) eftersom de är avgörande för hans motivering av tvåstegsmodellen.

7 Metoder i urval

I detta kapitel ska jag redogöra för Pörns metod och några tänkbara alternativ. Ett par av dessa har karaktären av "paketlösningar", dvs. de är i någon mening "hela metoder". Därtill kommer några tänkbara *modelleringsalternativ*, dvs. alternativ som är relevanta i begränsade aspekter av Pörns metod. Framställningen bygger huvudsakligen på kapitel 6. Fokus ligger såtillvida på metoder som föreslagits, använts eller används inom kärnkraftsområdet eller i liknande sammanhang.

I litteraturen ges modeller för tidsrelaterade fel ("λ-modeller") prioritet framför modeller för behovsrelaterade fel ("q-modeller"). Detta gäller både primär- och sekundärkällor. Ofta är detta försvarligt, dels eftersom modellerna är analoga, dels eftersom λ-modellerna är de mest betydelsefulla. Jag kommer att följa den rådande konventionen så länge jag anser att det inte leder till missförstånd.

7.1 Tvåstegs Bayes

Kaplan är än idag en av de teoretiker som flitigast refereras i litteratur som behandlar hierarkisk metodik i allmänhet och tvåstegsmetodik i synnerhet. Kaplans metod kan ses som ett slags urkund i sammanhanget, inte bara för att den är först i sin genre, utan också för att den är mycket generellt formulerad. Detta har emellertid också medfört att den utan vidare åthävor ansetts som bakgrund eller inspirationskälla till metoder som i vissa avseenden skiljer sig ganska mycket åt. Att två metoder hänförs till en gemensam utgångspunkt kan vara klargörande, det betyder ju att de i något avseende är lika. Om influenserna inte närmare preciseras kan å andra sidan sådana jämförelser vara förvirrande. Jag tar upp detta på förekommen anledning; litteraturen om bayesianska tvåstegsmetoder är full av exempel på ”diffusa” korsreferenser mellan olika metoder. Cooke et al. (2003) hänför t.ex. Pörns metod till Hora (1990) och ZEDB-metoden till Kaplan. ZEDB hänför å andra sidan den egna metoden till Pörn och Hora (ZEDB 2002, s. 33). Samtidigt menar Pörn att likheten mellan hans egen metod och Hora kommer av att de har Kaplan som gemensam utgångspunkt (Pörn 1996, s. 174). Även tvåstegsmodellen i R-DAT hänförs till Kaplan (Prediction Technologies, 2002).¹⁷

Den synbara förvirringen kommer sig av att dessa referenser avser olika saker utan att detta tydligt anges. Anledningen till att Cooke et al. (2003) refererar ZEDB-metoden till Kaplan torde ha att göra med att lognormalfördelningen används som prior.¹⁸ Att Pörn hänförs till Hora har på motsvarande sätt att göra med användningen av gammafördelningar. Pörn i sin tur tycks inte se denna likhet som fundamental utan tar istället fasta på strukturelligheter mellan sin metod och Kaplans vilket även förklarar hur ZEDB-metoden kan hänföras till Pörn och Hora.¹⁹ Även i R-DAT avses troligen den matematiska strukturen, dvs. tvåstegsmodellen i sig, eftersom lognormalfördelningen kan väljas som ett av många alternativ (Prediction Technologies, 2002).

Jag tar upp detta för att belysa att det är långtifrån självklart utifrån vilka kriterier olika tvåstegsmetoder ska jämföras. I själva verket är det inte uppenbart i vilka avseenden de kan kallas *olika*. Ett sätt att hantera situationen är att först och främst konstatera att tvåstegsmetodiken har blivit *allmängods* och sedan välja ut ett par ”helhetslösningar” som väl täcker in möjligheterna till variation. I en sådan jämförelse är Pörns metod redan given. Som kontrasterande exempel har jag valt den metod som används i ZEDB. Metoderna är inte bara intressanta för att de i olika teoretiska avseenden skiljer sig åt, utan också för att de är de mest kända metoderna i praktiskt bruk. Därmed kan även aspekter av tillämpningen jämföras. Till detta kommer att de redan jämförts med avseende på beräkningsresultat i en s.k. benchmark (T-book – ZEDB Benchmark, 2004).

7.1.1 Pörn

Pörns metod presenteras i avhandlingen *On Empirical Bayesian Inference Applied to Poisson Probability Models* (Pörn 1990) och användes första gången i samband med den tredje utgåvan av T-boken. Genom åren har vissa justeringar gjorts, dels till följd av teoretiska ställ-

¹⁷ R-DAT är en av de mer etablerade kommersiella programvarorna för bayesiansk analys.

¹⁸ I härledningen av sin metod använder Kaplan en lognormalfördelning som prior. Han påpekar emellertid att vilken kurvfamilj som helst, som kan antas innehålla *minst en god approximation* till den sanna variabilitetsfunktionen, kan användas (Kaplan 1983, s. 3). Därav framgår likväl att Kaplan tolkar priorn som en *population variability curve* (PVC).

¹⁹ Likheten mellan Kaplan och Hora understryks av Hofer (1999, s. 881).

ningstaganden, dels av praktiska hänsyn. Den framställning som görs här utgår i första hand från Pörns avhandling (Pörn 1990).²⁰ Viktiga förändringar tas upp fortlöpande.

Pörns metod bygger, enkelt uttryckt, på en tvåstegs bayesiansk modell med konjugerad prior (gammalfördelning) och icke-informativ hyperprior. För tidsrelaterade fel antas observationerna vara Poissonfördelade. Pörn föreslår även en *tvåparametrisk* modell, med en tids- och en behovsrelaterad parameter.²¹ Denna mer komplicerade modell är unik för Pörn. Även i valet av prior och hyperprior väljer Pörn att gå sin egen väg, åtminstone som metoden formulerades 1990.

7.1.1.1 Bakgrund

Pörns metod utvecklades som ett alternativ till de PEB-metoder som tidigare användes för härledning av T-bokens parametrar.

”Experiencing various difficulties to find reasonable point values in some cases and understanding the inconsistency hidden in the mixture of Bayesian and frequentistic reasoning became the seed of the development work in this study.” (Pörn 1990, s. 89)

Genom ansättandet av en hierarkisk bayesiansk modell kunde Pörn kringgå en rad förmenta svagheter hos den gamla metodiken:

- a) Problem relaterade till data, t.ex. liten variation och få fel.
- b) Problem relaterade till modellens robusthet.
- c) Begreppslig inkonsistens.

Problemen under (a) gäller generellt då klassiska skattningstekniker involveras. Pörn ägnar dem inte något större utrymme. I T-boken konstateras kort att "[...] skattning av de sekundära parametrarna (α, β) med klassiska metoder misslyckades" (T6 s. 31), och i inledningen till *On Empirical Bayesian Inference...*: "Among the more apparent consequences of this inconsistency one may mention the difficulty to find reasonable estimates of the secondary parameters." (Pörn 1990 s. 2)

Det andra citatet kan möjligen även syfta på problemen under (b). Dessa har att göra med att PEB-metoder inte antar någon *osäkerhet* i hyperparametrarna eftersom de *punktskattas*. PEB-metodikens lösning på vad Pörn kallar "the statistical inference problem for λ_{n+1} " beskrivs med den bayesianska ekvationen

$$p(\lambda_{n+1} | x_{n+1}) \propto p(x_{n+1} | \lambda_{n+1}) \cdot p^*(\lambda_{n+1}), \quad (7.1)$$

där $p^*(\lambda_{n+1})$ är en *skattning* av $p(\lambda_{n+1})$ på grundval av generiska data $\{x_1, \dots, x_n\}$. Den hierarkisk bayesianska tolkningen av $p(\lambda_{n+1})$ är att den är en "uncertain quantity in some space of distributions". Att ansätta $p^*(\lambda_{n+1})$ är därmed detsamma som att (efter att ha observerat $\{x_1, \dots, x_n\}$) tilldela en av medlemmarna i detta fördelningsrum sannolikheten 1. (Pörn 1990, s. 9)

²⁰ Jag har även haft god hjälp av en utvärdering gjord av prof. Roger Cooke på uppdrag av SKI (Cooke et al. 1995).

²¹ Jag vill påpeka att Cooke et al. (1995) i sin utvärdering av Pörns metod inte alls nämner den mer komplicerade modellen.

Svagheten med ett sådant förfarande är enligt Pörn att priorn inte representerar osäkerheten om förväntade observationer på ett rimligt sätt, osäkerheten är helt enkelt inte tillräckligt stor (fördelningen har för kort "svans"). Enligt Pörn är det ett grundläggande kriterium på robusthet att *a posteriorifördelningen så mycket som möjligt beror av data och så lite som möjligt av a priorifördelningen* i fall då *a priorifördelningen* är felaktig, dvs. i fall då *a priorifördelningen* och sannolikheten för observationen x_{n+1} (tagen som ML-skattningen x_{n+1}/T_{n+1}) har olika lokalitet. (Pörn 1990, s. 30) Detta kriterium torde enligt Pörn vara bättre uppfyllt med en modell där även hyperparametrar uttrycker osäkerhet. Pörn tar även upp en situation då likelihoodfunktionen med avseende på hyperparametrarna är mycket flack men likväl har ett maximum i (α^*, β^*) . Att då välja ut denna enda punkt framför alla andra *nästan lika sannolika* punkter (α, β) är enligt honom oförsvarligt. (Pörn 1990, s. 24-32)

Den ovan beskrivna problematiken rör i någon mening priorns täckningsgrad. Liten täckningsgrad gör modellen mindre flexibel, dvs. i ett avseende *mindre robust*. Därvidlag har PEB-modeller enligt Pörn en svaghet jämfört med hierarkiska modeller.

Anmärkning: Observera att Pörns resonemang förutsätter att data *inte dubbelräknas*, skattningen av x_{n+1} görs ju på grundval av observationerna $\{x_1, \dots, x_n\}$. Samtidigt tillstår Pörn att dubbelräkning kan vara ett sätt att komma tillrätta med robusthetsproblemen även om den principiellt sett är felaktig (Pörn 1990, s. 12). Pörns argument skulle därmed kunna tolkas som att *givet att data inte dubbelräknas* (vilket är den förment korrekta approachen) *så uppvisar PEB-metoder svagheter med avseende på robusthet*. Pörns argument drabbar alltså inte på något uppenbart sätt PEB-metoder där dubbelräkning faktiskt tillämpas, vilket likväl är regel snarare än undantag (se avsnitt 6.3).

□

Pörn menar att PEB-metoder är icke-bayesianska eftersom de involverar klassiska skattningsmetoder (Pörn 1990, s. 9). En möjlig slutsats utifrån problempunkt (c) är att PEB-metoder varken är bayesianska eller frekventistiska eftersom deras sannolikhetsbegrepp inte kan ges någon entydig tolkning. Pörns argumentation utgår från att λ (eller q i binomialfallet) i bayesiansk tradition är en *fiktiv* parameter snarare än en okänd fysikalisk storhet. Poängen med modellparametern ligger i att dess *fördelning* $p(\lambda)$ kan användas för att fånga upp ny information genom tillämpning av Bayes sats. Därmed finns ingen anledning att punktskatta λ . (Pörn 1990, s. 4) Ett analogt resonemang förs om $p(\lambda)$, dvs. i enlighet med bayesiansk tradition är även *fördelningen* för λ fiktiv. Osäkerheten beträffande $p(\lambda)$ överförs därmed till hyperparametrarna genom ansättandet av en fördelning $p(\alpha, \beta)$. (Pörn 1990 s. 9f) Motsvarande förfarande med PEB-metodik innebär att hyperparametrarna, och därmed $p(\lambda)$, istället "punktskattas" vilket kräver två olika sannolikhetsbegrepp, eller som Pörn själv uttrycker saken:

"[T]he PEB-approach is not logically consistent because we use sampling theory methods to estimate the prior distribution but otherwise interpret the densities as measures of knowledge rather than in a frequency sense." (Pörn 1990, s. 2)

Frågan är hur hårt ett sådant argument slår mot PEB-metodiken. Kan de två begreppen kombineras eller är PEB-metodikens utsagor rent av nonsens? Jag återkommer till dessa frågor i kapitel 11. För Pörns del tycks hur som helst de begreppsliga aspekterna ha varit väsentliga för valet av tvåstegsmetodik framför PEB. I T-boken säger han t.ex. att den egna modellen är "[---] fullständigt Bayesiansk och därmed logiskt konsistent" (T6, s. 36).

7.1.1.2 Tvåstegsmodellen

Pörn överför alltså osäkerheten i $p(\lambda)$ till hyperparametrarna genom att sätta

$$p(\lambda) = \int p(\lambda|\theta)p(\theta)d\theta, \quad (7.2)$$

där $\theta = (\alpha, \beta)$ (Pörn 1990, s. 47). A posteriorifördelningen $p(\lambda_{n+1}|x_1, \dots, x_{n+1})$ blir nu en funktion av $p(\theta)$ genom att (7.1) ersätts med

$$p(\lambda_{n+1}|x_1, \dots, x_{n+1}) \propto p(x_{n+1}|\lambda_{n+1}) \int p(\lambda_{n+1}|\theta)p(\theta|x_1, \dots, x_n)d\theta \quad (7.3)$$

(Pörn 1990, s. 12). Hyperposteriorifördelningen $p(\theta|x_1, \dots, x_n)$ erhålls i sin tur genom tillämpning av Bayes sats och ges av

$$p(\theta|x_1, \dots, x_n) \propto p(x_1, \dots, x_n|\theta)p(\theta) = \prod_{i=1}^n p(x_i|\theta)p(\theta) \quad (7.4)$$

(Pörn 1990, s. 11f, s. 48).²² Likelihoodfunktionen $p(x|\theta)$ i (7.4) är nu en *negativ binomialfördelning* (eftersom λ är gammafördelad) dvs.

$$p(x_i|\theta) = p(x_i|\alpha, \beta) = \text{NegBin}\left(\alpha, \frac{\beta}{\beta + T}\right) \quad (7.5)$$

(Pörn 1990, s. 19). Den negativa binomialfördelningen kallas även *gamma-Poissonfördelning*. Den känns igen från exempel 6.1 (avsnitt 6.3.1) och erhålls som fördelning för $p(x|\theta)$ då $p(\lambda|\theta) \in \Gamma(\alpha, \beta)$ och $p(x|\lambda) \in \text{Po}(\lambda T)$. Detta visas i appendix F.2. För en fullständig härledning av Pörns tvåstegsmodell hänvisar jag till Cooke et al. (1995, pt. 2, s. 7ff).

(7.3) kan sägas utgöra den generella matematiska strukturen för λ -modellen i Pörns metod. Pörns *prior*, motsvarande (7.2), kan därmed skrivas

$$p(\lambda_{n+1}|x_1, \dots, x_n) = \int p(\lambda_{n+1}|\theta)p(\theta|x_1, \dots, x_n)d\theta. \quad (7.6)$$

Pörn erhåller alltså sin prior genom att medelvärdesbilda över en mängd olika fördelningar $p(\lambda|\theta)$ som viktats med avseende på sannolikheten för θ givet observationerna $\{x_1, \dots, x_n\}$. Priorn är den s.k. *generiska fördelningen* (Pörn 1990, s. 23). Två saker är värda att notera:

- Integrationen över en mängd olika gammafördelningar i priorn (7.6) resulterar i en *blandning av gammafördelningar* vilken i sin tur *inte* är en gammafördelning.
- Med (7.3) visar Pörn tydligt att observationen x_{n+1} *inte* ska dubbelräknas.

7.1.1.3 Likelihoodfunktionen

I de tidigaste T-böckerna (T1 och T2) modellerades likelihoodfunktionen som i många andra metoder, dvs. antalet fel antogs vara observationer av antingen Poisson- eller binomialförde-

²² Detta kräver att x_i är oberoende av θ .

lade stokastiska variabler (med modellparametrarna λ och q). Från och med T-boken version 3 ersätts emellertid q i många fall av en s.k. standby-felintensitet λ_s (T6). Denna har ML-skattningen $\lambda_s^* = x/T$ där x är antalet *behovsrelaterade fel* och T är *tiden i standby*. λ_s hanteras i övrigt på samma sätt som λ [Pörn 1].²³ Bytet av modell förklaras med att felmekanismerna under standby-tid dominerar över dem vid aktivering samt att en helt och hållet behovsrelaterad parameter tenderar att återspegla skillnad i testintervall snarare än den verkliga felbenägenheten (Pörn 1990, s. 89f). Detta framgår i en rapport av Sven Olof Andersson vid Vattenfall Power Consultant AB (Andersson 1993). Ytterligare en fördel med en tidsrelaterad parameter är att information om den totala standby-tiden är betydligt mer lättåtkomlig än information om antalet aktiveringar (T6, s. 32). I själva verket förekommer emellertid behovs- och tidsrelaterade fel samtidigt (Pörn 1990, s. 91). Detta kräver en modell med två parametrar. En sådan, den s.k. " $q_0 + \lambda_s T$ -modellen", har utvecklats för några grupper av periodiskt testade komponenter där antalet aktiveringar är relativt välkänt, vilket framförallt gäller komponenter där antalet verkliga behov är litet jämfört med behov vid test (T6, s. 32. Pörn 1990, s. 90).²⁴ Enligt Pörn kräver modellen inte att data om de två ingående parametrarna särskiljs. Däremot måste komponenter med olika testintervall vara representerade i gruppen för att relationen mellan q_0 och λ_s ska bli meningsfull. (Pörn 1990, s. 89ff. T6, s. 32f) I $q_0 + \lambda_s T$ -modellen ersätts drifttiden av antalet aktiveringar och aktiveringsintervall (T6, s. 37).

Sammanfattningsvis kan sägas att någon "ren" q -modell, med bara behovsrelaterade fel, inte längre används [Pörn 1]. Det framstår också som att λ_s -modellen är ett slags provisorium och att det idealt sett är $q_0 + \lambda_s T$ -modellen som bör användas.

7.1.1.4 Prior

Inför valet av prior sätter Pörn upp tre minimikrav: Priorn bör vara

- Flexibel
- Robust
- Matematiskt hanterbar

Pörn utgår från det sista kravet och väljer att använda *konjugerade* fördelningar, dvs. för λ -modellen en klass av gammafördelningar $\Gamma(\cdot)$. För λ -modellen kan Pörns prior preliminärt skrivas

$$p(\lambda|\theta) = p_1 = \Gamma(\lambda|\alpha, \beta). \quad (7.7)$$

Konjugerade fördelningar tenderar emellertid att vara okänsliga för data "leading to posterior estimates which are too much influenced by the prior distribution also in cases when their prior seems to be less realistic" (Pörn 1990, s. 14). Detta är samma slags brist på robusthet som PEB-metodernas punktskattningar ger upphov till (se ovan). Pörn löser detta genom att till den informativa fördelningen $p_1 = \Gamma(\lambda|\alpha, \beta)$ addera en *icke informativ* fördelning $p_2 = \Gamma(\lambda|\frac{1}{2}, \beta_0)$. Den senare kallas för "the contamination distribution" och antas öka flexibiliteten och därmed robustheten [Pörn 2]. I p_2 har "hyperparametrarna" fixerats; $\alpha = \frac{1}{2}$ är konsistent

²³ För att skilja dessa parametrar åt används i T-boken (version 6) beteckningarna λ_s respektive λ_d (se t.ex. T6 s. 39). I tidigare versioner åtföljs de olika felintensiteterna istället av angivelser om "drift-" respektive "standbytid" och/eller angivelser om driftläge ("driftsatt" respektive "standby").

²⁴ Jag har indexerat λ för att tydliggöra att det är den s.k. standby-felintensiteten som används. I T-boken och Pörn (1990) anges modellen utan index.

med Jeffreys regel och Box & Tiaos rekommendationer för icke-informativa a priorifördelningar (se avsnitt 5.2.3). β_0 tilldelas ett litet positivt värde för att fördelningen ska bli äkta. (Pörn 1990, s. 14-18) Den fullständiga priorn skrivs nu

$$p(\lambda|\theta) = (1-c)p_1 + cp_2, \quad (7.8)$$

där c kan anta värden i intervallet $[0, 1]$ och således avgör hur mycket fördelningen ska "kontamineras" (Pörn 1990, s. 18).

I sin utvärdering av Pörns metod invänder Cooke et al. (1995) mot att den icke-informativa "kontamineringen" används tillsammans med en *icke-informativ hyperprior* (se nedan). Cooke et al. tycks mena att poängen med att ansätta en icke-informativ hyperprior går förlorad om denna inte har fullt inflytande över priorn (att p_2 har fixa parametrar gör att den inte påverkas av hyperpriorn). Cooke et al. säger vidare att: "[...] if the arguments for introducing non-informative priors on α and β carry any weight, then the same arguments would un motivate the introduction of [p_2]" (Cooke et al. 1995, pt.2, s. 7). Cooke et al. menar alltså att det som Pörn försöker åstadkomma med p_2 är hyperpriorns uppgift. Dessutom menar Cooke et al. att p_2 kan betraktas som högst *informativ* (Cooke et al. 1995, pt.2, s. 7). Roger Cooke har också tillstått att Pörns hyperprior är tillräckligt flexibel för att kontamineringen ska kunna betraktas som överflödig, därför har p_2 sedermera tagits bort [Pörn 2]. Aktuell prior är alltså p_1 enligt (7.7).

7.1.1.5 Hyperprior

Pörn väljer som sagt att använda en *icke-informativ hyperprior*. Han motiverar sitt val med att icke-informativa fördelningar "allows data to speak for themselves". Pörn ser dem därmed som ett slags "standardfördelningar" i situationer då inga tidigare observationer har gjorts. (Pörn 1990, s. 20) Eftersom det kan vara svårt att hitta "data translated likelihoods" i det flerparametriga fallet tillämpar Pörn sedan den enparametriska versionen av Jeffreys regel på hyperparametrarna var för sig, för att på så sätt åstadkomma en *approximativt icke-informativ hyperprior*. Denna version av Jeffreys regel kräver emellertid att parametrarna är oberoende fördelade vilket inte gäller α och β eftersom α inte är en lägesvariabel (Pörn 1990 s. 21). Pörn använder därför istället väntevärdet²⁵ $\mu = \alpha/\beta$ och variationskoefficienten $v = \sqrt{\alpha}$ för vilka han anser att fördelningarna är *a priori oberoende* i den meningen att "any prior idea about the location μ is not (much) influenced by one's idea about the relative spread" (Pörn 1990, s. 22). När dessa sedan transformeras tillbaka till α och β erhålls hyperpriorn

$$p(\theta) = p(\alpha, \beta) \propto \frac{1}{\beta \sqrt{\alpha \left(\alpha + \frac{\beta}{T} \right)}}. \quad (7.9)$$

För komponenter med noll fel (i T6) används istället en hyperprior proportionell mot

²⁵ I själva verket $\mu' = T\alpha/\beta$ som tolkas som det förväntade antalet fel vid tiden T givet α och β . (Cooke et al. 2003, s. 11)

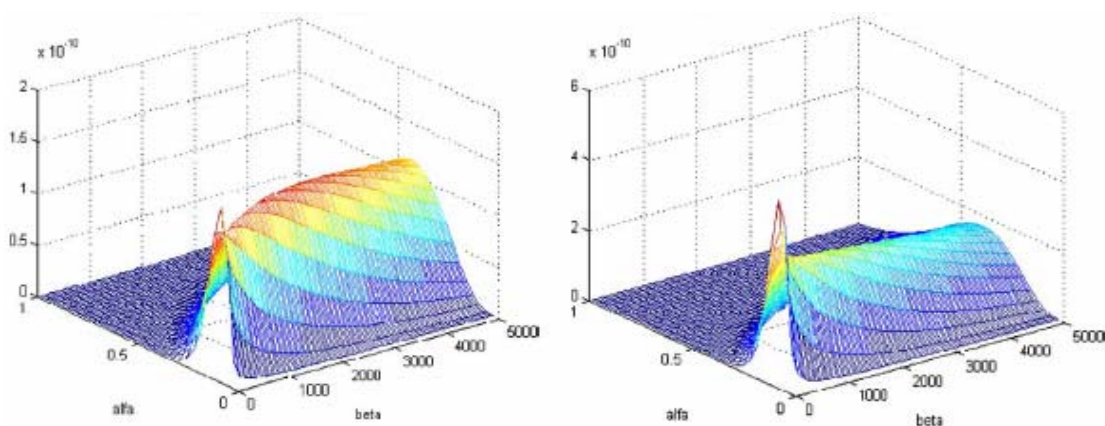
²⁶ En utförlig diskussion av Pörns hyperprior återfinns hos Meyer and Hennings (1999).

$$\frac{1}{\alpha\beta\sqrt{\alpha + \frac{\beta}{T}}} \cdot [\text{Pörn 2}] \quad (7.10)$$

Hyperprior (7.9) är oäkta vilket enligt Cooke et al. (1995) skapar vissa svårigheter: För den enkla modellen gäller att en oäkta fördelning blir äkta efter uppdatering med bara ett fel. Det samma gäller emellertid inte med en *oäkta hyperprior*, dvs. hyperposteriorifördelningen (7.4) blir inte nödvändigtvis äkta vid uppdatering. Detta beror på att hyperposteriorifördelningens uppförande framförallt bestäms av likelihoodfunktionen (7.5) som *inte har något maximum*. (Cooke et al. 1995, pt 2, s. 12) Detta är samma resultat som hos Cooke et al. (2003) (se avsnitt 6.2.2). Cooke et al. (1995) har gjort sina beräkningar med en likformig hyperprior men anser att resultaten är relevanta även med avseende på Pörns hyperprior:

”Improper uniform hyperpriors do not become proper after observing data; they stay improper and the posterior expectation for λ remains infinite. This means that the method of truncation essentially determines the posterior distribution of λ . If this conclusion does NOT hold for the non-informative hyperpriors of section 1.3 [dvs. Pörns hyperprior] then this fact is worthy of proof.” (Cooke et al. 1995 pt. 2, s. 15)

Detta betyder att Pörns hyperprior måste trunkeras och att modellen dessutom är känslig för hur trunkeringen görs. Cooke et al. (1995) visar att olika, men likväl försvarbara, val av trunkeringsgränser kan påverka medianen och 95%-percentilen med en faktor 5 (Cooke et al. 1995 pt. 2, s. 18). Även detta resultat återfinns hos Cooke et al. (2003). Hos Cooke et al. (2003) testas emellertid gammamodellen med två olika val av hyperprior; dels Pörns egen, dels en baserad på den enparametriska versionen av Jeffreys regel *utan transformation*. Resultatet visar att Pörns hyperprior bättre lyckas ”pressa ner” likelihoodfunktionen (Cooke et al. 2003, s. 17), vilket ger en hyperposteriorifördelning som rimligen är något lättare att trunkera (se fig. 5). Cooke et al. (2003) drar som tidigare nämnts ändå slutsatsen att lognormalfördelningen har betydande fördelar med avseende på trunkering (se avsnitt 6.2.2).



Figur 5. Hyperposteriorifördelningen i gammafallet med ”Jeffreys hyperprior” (utan transformation) respektive Pörns hyperprior (Cooke et al 2003, s. 17).

Slutligen ska påpekas att Pörn har gjort beräkningar som visar att valet av prior inte har någon signifikant betydelse (T-book – ZEDB Benchmark, 2004). Huruvida detta motsäger resultaten hos Cooke et al. (2003) är svårt att avgöra. Cooke et al. har för det första inte gjort sina beräkningar i T-Code. För det andra kan det fortfarande vara enklare att trunkera i lognormalfallet, även om Pörn har visat att gammafallet är hanterbart.

7.1.1.6 Tillämpning

Hyperposteriorifördelningen (7.4) beräknar Pörn utifrån generella data $\{x_1, \dots, x_n\}$, där x_i är antalet fel för respektive *enskild* komponent i den relevanta komponentgruppen (T-book – ZEDB Benchmark, 2004). Data för den ”aktuella komponenten” x_{n+1} räknas alltså inte då prior framställs (dvs. den tas inte med i integranden i (7.3)). Priorn uppdateras sedan med avseende på x_{n+1} . A posteriorifördelningen (7.3) är alltså *komponentspecifik*. En *anläggningsspecifik* fördelning bildas sedan som en linjärkombination av de komponentspecifika fördelningarna (Pörn 1996, s. 178).²⁷ I en mening har alltså Pörns metod *tre* steg där den bayesianska tvåstegsanalysen omfattar en generisk och en mängd komponentspecifika fördelningar. Den anläggningsspecifika fördelningen tas såtillvida fram i ett externt (i förhållande till tvåstegsmodellen) sista steg då alla komponentspecifika fördelningar är beräknade.

Att varje komponent tillskrivs en specifik felintensitet motiverar Pörn med att ”[o]m enskilda komponenters felbenägenhet vore kända med stor precision, skulle dessa inte vara sinsemellan lika” (T6, s. 32). Den anläggningsspecifika fördelningen återspeglar såtillvida ”den osäkerhet som gäller för en för anläggningen ifråga karaktäristisk komponents felbenägenhet” (T6, s. 34ff.) men också hur felintensiteten varierar mellan de enskilda komponenterna, dvs. både en osäkerhetsfördelning och en PVC. Pörns anläggningsspecifika fördelningar är mycket breda jämfört med de fördelningar som blir resultatet om komponenterna grupperas redan från början (jmf. ZEDB-metoden nedan). Pörns ”trestegsmodell” ger i vissa fall upphov till anläggningsspecifika fördelningar med 5%-percentiler i storleksordningen 10^{-40} . (T-book – ZEDB Benchmark, 2004)

Av ovanstående liksom av (7.3) framgår att Pörn inte dubbelräknar den aktuella komponenten x_{n+1} . Denna princip frångås emellertid i T-boken. Detta syns på att *varje komponentgrupp har en gemensam generisk fördelning* (se appendix C), medan (7.3) resulterar i att varje komponent får en egen generisk fördelning. Detta avsteg från metodens påbud torde ha rent praktiska orsaker; det är helt enkelt mycket lättare att beräkna *en* generisk fördelning än låt säga 730 (vilket är antalet komponenter i gruppen ”pneumatiska skalventiler”, T6, tab. 3.20.1, s. 205).²⁸ Den generiska fördelningen i (Pörn 1990) är alltså *inte* identisk med den generiska fördelningen i T-boken.

Vid beräkning av T-bokens parametrar gör alltså Pörn avkall på ett resultat som i avhandlingen framstår som högst väsentligt (Pörn 1990, s. 12). Till Pörns fördel hävdar Atwood et al. (2003) att dubbelräkning i praktiken inte ger någon signifikant skillnad med avseende på beräkningsresultat. Rimligen blir också skillnaden mindre ju fler komponenter som räknas i den generiska fördelningen. (Atwood et al. 2003 ch. 8, s. 3) Eftersom Pörn grupperar enligt principen ”en komponent – en individ”, dvs. med minsta möjliga enhet, har han alltså ett slags teoretiskt alibi för sin (enligt honom själv) felaktiga dubbelräkning.

²⁷ Som jag tolkar detta handlar det om en linjärkombination där alla fördelningar viktas lika, dvs. ett slags centralmått $1/n \cdot \sum p_i(\lambda)$. En sådan tolkning är konsistent med det resonemang som förs i T-boken (T6, s. 34ff.).

²⁸ Detta bekräftas också av Kurt Pörn själv [Pörn 1].

I T-boken anges de anläggningsspecifika fördelningarna för respektive komponentgrupp tillsammans med *deras gemensamma* generiska fördelning. Komponentspecifika fördelningar anges däremot inte. En adekvat fråga i sammanhanget är vad den generiska fördelningen används till. I T-boken sägs att ”den generiska fördelningen kan [...] användas som a priorifördelning i komponentspecifika analyser” (T6, s. 34). Detta överensstämmer med dess roll i den matematiska modellen. Att den också presenteras i T-boken förklaras med att det i vissa fall saknas anläggningsspecifika data. Det kan t.ex. vara så att en ny komponent introduceras på en anläggning men att den redan tidigare funnits på andra anläggningar. Den finns då visserligen representerad i T-boken, men bara för andra anläggningar än den aktuella. I så fall kan den generiska fördelningen anses mer representativ än någon godtyckligt vald anläggningsspecifik fördelning. Det händer också att T-boksdata efterfrågas på anläggningar utan koppling till TUD-systemet. Även då kan den generiska fördelningen anses mer relevant. Det kan vara värt att notera att i dessa fall kommer Pörns metod att tillämpas så som det var tänkt: Under förutsättning att *ingen dubbelräkning får göras* är nämligen den befintliga generiska fördelningen bara giltig för komponenter som inte finns med i T-boken.

Anmärkning: Ytterligare en adekvat fråga är följande: Om nu Pörn räknade ut 730 olika generiska fördelningar (t.ex. för gruppen av pneumatiska skalventiler), skulle dessa då kunna användas på det sätt som beskrivits ovan? I fallet med en ny komponent, vilken generisk fördelning ska anses mest representativ? Bör de generiska fördelningarna i så fall vägas samman?

□

7.1.2 ZEDB-metoden

ZEDB kan sägas vara Tysklands motsvarighet till TUD. ZEDB är alltså framförallt namnet på en databas (Zentrale Zuverlässigkeits- und Ereignisdatenbank)²⁹. ZEDB används emellertid också som benämning på det verktyg som används för beräkning av parametrar till ”ZEDB-boken” (ZEDB 2002), liksom den matematiska metoden (se t.ex. Cooke et al. 2003).

De matematiska modellerna är inte åtkomliga i detalj. Min studie har gjorts med utgångspunkt från ett kortare matematiskt avsnitt i ZEDB-boken samt den studie av ZEDB-metoden som görs hos Cooke et al. (2003). Kompletterande information har erhållits från den benchmark som gjorts för ZEDB-metoden och Pörns metod (T-book – ZEDB Benchmark, 2004) samt korrespondens med Günter Becker, företrädare för RISA Sicherheitsanalysen GmbH som ansvarar för programvaran.

Till att börja med finns inga skäl att tro att ZEDB-metoden skulle skilja sig stort från de allmänna drag som karaktäriserar Pörns metod. I ZEDB-boken hänförs metoden till ”The superpopulation technique – also commonly known as the two-stage Bayesian model [,] a standard method applied internationally”, med referenser till Pörn, Iman och Hora. Modellbeskrivningen samt utdatatabeller vittnar också om att principerna är desamma som hos Pörn.³⁰ (ZEDB 2002, s.33ff) Vad som ofta lyfts fram som den avgörande skillnaden är att felintensiteten antas vara lognormalfördelad i ZEDB-metoden istället för gammafördelad som hos Pörn. Detta är möjligen också anledningen till att ZEDB-metoden ibland hänförs till Kaplan (se t.ex. Cooke et al. 2003, s. 7). Vidare används lognormalfördelningen både för tids- och behovsrelaterade fel, huvudsakligen tillsammans med en Poissonfördelning. För ett fåtal komponenttyper används en modell med en trunkerad lognormalfördelning och binomialfördelning som likelihoodfunktion. Någon motsvarighet till Pörns $q_0 + \lambda_s T$ finns inte. [Becker]

²⁹ Översätts i Cooke et al. (2003) till ”The Centralised Reliability and Events Database”.

³⁰ Liksom hos Pörn dubbelräknas komponenterna i praktiken men inte i teorin.

I ZEDB-metoden görs likväl samma grundantaganden som hos Pörn (Cooke et al. 2003 s. 4ff) med ett viktigt undantag: Komponenterna grupperas *anläggningsvis*. De anläggningsspecifika fördelningarna beräknas alltså inte i efterhand som hos Pörn. Detta leder till långt snävare fördelningar än hos Pörn (olika mycket beroende på hur många komponenter som grupperas). I ZEDB-metoden antas alltså att det finns en *variabilitet mellan anläggningar* men att komponenter av en viss typ är identiska med avseende på felbenägenhet. (T-book – ZEDB Benchmark, 2004, s. 5)

Hyperpriorn $p(\mu, \sigma)$ framställs genom tillämpning av den enparametriska versionen av Jeffreys regel, allt enligt Box & Tiaos rekommendationer. Hyperparametrarna antas alltså vara oberoende och Jeffreys regel tillämpas på parametrarna var för sig. Eftersom antagandet om oberoende är rimligare i lognormal- än i gammafallet görs ingen transformering. Detta ger hyperpriorn

$$p(\mu, \sigma) \propto \frac{1}{\sigma}, \quad (7.11)$$

(Cooke et al. 2003, s. 20). μ antas alltså vara likformigt fördelad (jmf. Siu & Kelly 1998, s. 107).

Även om lognormalfördelningen rekommenderas och används så stöder ZEDB-verktyget även användning av gammafördelningar (i λ -fallet) (Cooke et al. 2003, s. 14).

7.1.3 Övriga

I litteraturen finns rikligt med kommentarer till tvåstegsmetodiken. Framförallt gäller detta framställningar av icke-informativa hyperpriors (jmf. Pörn 1990, s. 17). En utförlig diskussion av Pörns hyperprior återfinns t.ex. hos Meyer & Hennings (1999). Becker föreslår fyra olika hyperpriors (alla baserade på Jeffreys regel) för implementering i ZEDB (Cooke et al. 2003, s. 19ff). Hora och Iman tillämpar den multiparametriska versionen av Jeffreys regel i samma situation som Pörn (dvs. med gammafördelad felintensitet) och erhåller den approximativa (likväl oäkta) hyperpriorn $\alpha^{-1/2}\beta^{-1}$ (Hofer et al. 1997).

Den kritik som Cooke et al. (2003) riktar mot tvåstegsmetoden har redan berörts. Sammanfattningsvis kan sägas att de reserverar sig mot användning av icke-informativa hyperpriors som förblir oäkta vid uppdatering: "Improper hyperpriors should be avoided if propriety after observations cannot be demonstrated." (Cooke et al. 2003, s. 37)

Hofer (1999) hävdar att befintliga tvåstegsmetoder (Pörn, Kaplan, Hora & Iman) genom sitt sätt att ordna integralerna "refer to the wrong chance mechanism" och därmed inte lyckas göra reda för skillnaden mellan subjektiv och "klassisk" osäkerhet (där de senare härrör från likelihoodfunktionen) (Hofer 1999, s. 881). Cooke et al. invänder emellertid att Hofers alternativa modell är motsägelsefull och att en konsekvent tillämpning av gängse antaganden om oberoende gör att hans modell sammanfaller med den han själv kritiserar (Cooke et al. 2003, s. 37).

7.2 Parametric Empirical Bayes

PEB-metoder finns i en mängd varianter varav jag har tagit upp de vanligaste. Till detta kommer en hel arsenal av metoder för att åstadkomma approximationer till en "fully bayesian

analysis” (Siu & Kelly 1998, s. 101). En genomgripande redogörelse för PEB med ”tillbehör” återfinns hos Carlin & Louis (2000).

Inom kärnkraftsområdet var PEB-metoder tidigare standard. När det gäller händelser som är relevanta för T-boken har de emellertid ersatts av tvåstegsmetoder. Ett undantag som bekräftar regeln är Vaurios PREB.

7.2.1 Vaurios PREB

Vaurios PREB (*Parametric Robust Empirical Bayes*) presenterades första gången 1987 som ett enklare alternativ till de tvåstegsmetoder som utvecklats av Kaplan och Fröhner (Vaurio 1987, s. 329).³¹ Metoden används sedan slutet av 80-talet i probabilistiska säkerhetsanalyser på kärnkraftverket Loviisa (Fortum Power and Heat Oy, Finland), bl.a. för parameterskattningar av de slag som är relevanta för T-boken (Vaurio & Jänkälä 2006, 217f).

Som alternativ till tvåstegsmodellen föreslår Vaurio en s.k. ”prior moment matching method” baserad på Hill et al. (1984), men som till skillnad från denna antas kunna hantera fall med noll fel och identiska data (ingen varians) (Vaurio 1987, s. 329).³² Metoden har sedan dess kritiserats av Cooke et al. (2003) liksom Bunea et al. (2005) för att sakna just dessa egenskaper: ”Elegance and simplicity are its main advantages. Disadvantages are that it cannot be applied if all $X_i = 0$, or if the population consists of only 2 plants.” Metoden kritiseras också för att data dubbelräknas, vilket antas leda till snäva fördelningar, och för att den leder till icke-konservativa resultat vid extremt stora empiriska varianser. (Cooke et al. 2003, s. 9. Bunea et al. 2005, s. 125f)

Vaurio bemöter denna kritik i en publikation daterad 2006, där han likväl hävdar att ”[---] PREB works as well if not better than the alternative more complex methods, especially in demanding problems of small samples, identical data and zero failures” (Vaurio & Jänkälä 2006, s. 209). Ytterligare ett skäl till nypubliceringen är att metoden, trots att den utarbetades under 80-talet tycks vara okänd bland många praktiker (Vaurio & Jänkälä 2006, s. 210).

Vaurios metod är en PEB baserad på en *momentmetod* där hyperparametrarna skattas med hjälp av a priorifördelningens medelvärde och varians. Vidare används konjugerade fördelningar (gamma-Poisson, beta-binomial) vilket ger lösningar på slutan analytisk form enligt:

$$p(\lambda|x) = \Gamma(\hat{\alpha} + x_i, \hat{\beta} + T_i)$$

(jmf. (5.3)). För att hantera fall då momenten är lika med noll används s.k. ”additional bias terms” som är asymptotiskt = 0 (dvs. då $T \rightarrow \infty$). För medelvärdet utnyttjas den heuristik som presenterades i avsnitt 5.2.3. I ”nollfelsfallet” skattas alltså λ med $(0 + \frac{1}{2})/T$ vilket är konsistent med Jeffreys icke-informativa prior (ekvation (5.4)). Variansen skattas i sin tur med

$$V = v + \frac{m}{T^*} \quad (7.12)$$

³¹ Fröhners metod har vissa allvarliga begränsningar varför den inte behandlas i denna rapport. Metoden kritiserar framförallt av Kaplan (1985) men även av Cooke et al. (2003, s. 10).

³² Hill, J.R. et al. (1984) ”The Application of Stein and Related Parametric Empirical Bayes Estimators to the Nuclear Plant Reliability data System.”, Atwood et al. 2003/CR-3637, EGG-2295, university of Texas, Austin.

där m är det *empiriska populationsmedelvärdet* och v är motsvarande varians. I bias-termen m/T^* är

$$T^* = T - \max(T_i), \quad (7.13)$$

$$T = \sum_{k=1}^k T_i. \quad (7.14)$$

T^* är en "fri parameter" i den meningen att Vaurio ser sitt eget val som tillfälligt. I själva verket är den "kind of maximal" och kan modifieras eller omdefinieras av användaren. (Vaurio & Jänkälä 2006, s. 218)

Variansen V är således positiv om något $x_i > 0$. Vidare gäller att om alla x_i är lika och alla T_i är lika så kommer alla enheter att få en a posteriorifördelning som är konsistent med att data grupperas. Detta motiverar Vaurio med att: "With identical data one obviously expects the estimation method to yield results consistent with the assumption that the components are identical and the data can be pooled." (Vaurio 1987, s. 331) Denna a posteriorifördelning tilldelas likväl den enhet som har den längsta observationstiden i fall då alla x_i men inte alla T_i är lika (vilket är vanligt i små stickprov), medan enheter med kortare T_i får a posteriorifördelningar som uttrycker större osäkerhet "as expected" (Vaurio & Jänkälä 2006, s. 212).

Vaurios metod är *heuristisk* i den meningen att den till viss del är resultat av "inductive reasoning". Vaurio tillstår att vissa ekvationer kan förefalla märkliga men att metoden bevisligen har mycket goda egenskaper både asymptotiskt och i fall med små stickprov. Att den garanterar *realistiska* lösningar även då stickprovsvariansen är liten framhålls som metodens främsta attribut: "The key idea in PREB is to guarantee realistic solutions even in case the empirical variance happens to be small." (Vaurio & Jänkälä 2006, s. 212) Förmågan att hantera situationer med små eller nollvärda empiriska moment är också den huvudsakliga anledningen till att Vaurio kallar metoden för *robust*. Bias-termerna är konservativa i den meningen att de hindrar att momenten systematiskt underskattas. Vidare minimeras variansen för a posteriorifördelningarnas medelvärden genom utnyttjande av s.k. *optimala vikter*. Tack vare dessa kan skattade medelvärden (λ^*) uttryckas som *Steins shrinkage-estimators* vilket enkelt uttryckt betyder att skattningarnas varianser inte blir onödigt stora. (Vaurio & Jänkälä 2006, s. 210ff)

I sin kritik av gängse hierarkiska metoder betonar Vaurio de komplicerade multidimensionella integrationer som blir fallet oavsett vilket slags prior som används (Vaurio & Jänkälä 2006, s. 209). Utnyttjandet av konjugerade fördelningar betyder i Vaurios fall att beräkningar kan göras med hjälp av mycket små datorprogram eller miniräknare (Vaurio 1987, 329). Vaurio ser heller ingen anledning att "utesluta" den aktuella enheten ur populationsdata; hans egna resultat visar inte några tecken på att PREB skulle ge de snäva fördelningar som ofta förväntas (Vaurio & Jänkälä 2006, s. 218f). Jag har inte sett att Vaurio ger några explicita argument för dubbelräkningen, snarare tycks han se den som det enda rimliga alternativet; att utesluta den aktuella komponenten har enligt Vaurio den svårbegripliga implikationen att "[---] each of the components is a sample from a different source population" (Vaurio 1990b, s. 127). Pörn å sin sida (1990, s. 12) antar att Vaurio dubbelräknar "[---] in fear of otherwise losing some information". Hur det än är med den saken så kan dubbelräkningen i praktiken vara ett sätt att kringgå den *brist på robusthet* som Pörn anser att PEB-metoder har genom att de inte antar en tillräckligt flexibel prior (Pörn 1990, s. 12) (frågan togs upp i avsnitt 7.1.1.1 och behandlas ytterligare i kapitel 11).

Den kritik som går ut på att PEB-metoder är begreppsligt inkonsistenta (bl.a. framförd av Pörn) bemöter Vaurio genom att, enkelt uttryckt, hävda att resultat med PREB kan ges en *fullständigt frekventistisk tolkning*. Vaurio menar att om "a priorifördelningen" bestäms med klassisk metodik, så kommer likväl a posteriorifördelningen att kunna tolkas i klassiska termer. Utgångspunkten är att en familj av klassiska konfidensgränser definierar "[...] a complete distribution for an unknown parameter being estimated". (Vaurio 1990a, s. 55) Vaurio menar att detta faktum ofta förbises eftersom statistiker hellre pratar om enskilda konfidensintervall snarare än hela fördelningar (Vaurio 1990b, s. 122). Vidare är en fördelning framställd genom tillämpning av Bayes sats identisk med en sådan "klassisk fördelning" under förutsättning att den generiska fördelningen är känd: "When the distribution of the source population is known and used as a prior distribution, then the classical and Bayesian techniques coincide, and the posterior density also has the fractional meaning." (Vaurio 1990b, s. 121) När den generiska fördelningen är okänd är likväl det bästa alternativet att skatta den med klassisk metodik eftersom "[t]he best probability has a fractional meaning" (Vaurio 1990b, s. 122).

Det finns med andra ord ingen konflikt mellan bayesianska och klassiska tekniker så länge data utgörs av "direkta observationer" (Vaurio 1990b, s. 125ff). Det är först när underlaget till a priorifördelningen är diskutabelt ("disputable") som de två riktningarna divergerar. Om en icke-informativ a priorifördelning används har skillnaden sällan någon praktisk betydelse eftersom a posteriorifördelningen då framförallt kommer att influeras av data. Vaurio menar emellertid att en sådan fördelning logiskt sett är subjektiv så länge den inte bevisats konsistent med klassiska resultat. (Vaurio 1990b, s. 127)

Ett avgörande resultat hos Vaurio är att klassiska fördelningar kan användas på konventionellt sätt i PSA-sammanhang (Vaurio 1990b, s. 123). En klassisk fördelning kan nämligen tolkas helt analogt med en bayesiansk fördelning, dvs. *datas fördelning kan tolkas som en fördelning för parametern* (Vaurio 1990b, s. 121). Det enda som krävs därvidlag är "[a] change of reference" (Vaurio 1990b, s. 123), dvs. ett slags *perspektivskifte*. Enligt Vaurio finns emellertid ett problem: Vid samtidig sampling från klassiska och bayesianska fördelningar blir resultatet automatiskt subjektivt (Vaurio 1990b, s. 128).

PREB är definierad för både tids- och behovsrelaterade fel (gamma-Poisson respektive beta-binomial). Liksom Pörn använder Vaurio λ -modellen för de flesta standby-komponenter (dvs. motsvarande Pörns λ_s -modell). q -modellen används endast sparsamt. Angående möjligheten av en " $q+\lambda$ -modell" menar Vaurio att den grundläggande metodiken existerar om, och när, det kan avgöras huruvida ett godtyckligt fel har uppstått pga. tids- eller behovsrelaterade påfrestningar (Vaurio & Jänkälä 2006, s. 213). Vidare förordar Vaurio en grupperingsprincip motsvarande Pörns. [Holmberg, Vaurio]

Vaurio presenterar sin metod på "algoritmisk" form (Vaurio & Jänkälä 2006, s. 212). Algoritmen finns återgiven i appendix E.

7.2.2 Övriga

I de två första utgåvorna av T-boken användes olika varianter av PEB med konjugerade fördelningar (gamma-Poisson, beta-binomial). För skattning av hyperparametrarna användes företrädesvis ML-metoden. I fall då denna "kollapsade" användes istället viktade momentmetoder med utnyttjande av antingen marginal- eller a priorifördelningen. I T-bokens tabeller presenterades endast hyperparametrarna och användaren kunde sedan själv konstruera sin a posteriorifördelning genom att till α lägga antalet fel x för en specifik komponent och till β

motsvarande drifttid T (jmf. (5.3)). (T1, s. 25ff. T2, s. 15ff) Även om dessa metoder för T-bokens vidkommande har ersatts av tvåstegsmetodik så används de fortfarande för skattning av andra väsentliga "PSA-parametrar". [FKA]

Cooke et al. utvärderar en s.k. *Nonparametric Empirical Bayes* (NPEB) baserad på en Dirichletfördelning: "The Dirichlet two-stage model has the advantages of being analytically solvable with hyperparameters possessing a clear interpretation. Initial numerical results are encouraging." (Cooke et al. 2003, s. 37) Vaurio betraktar denna metod som mindre intressant än PEB eftersom den "requires a number of cells and boundary parameters to be subjectively selected" (Vaurio & Jänkälä 2006, s. 209f). Metoden är relevant för detta arbete men har av praktiska skäl lämnats utanför.

7.3 Ett klassiskt alternativ

Klassiska metoder används i kärnkraftssammanhang för hypotesprövning, test för oberoende, variansanalys, trendanalys mm. (Atwood ch. 6-7). Det finns mig veterligen bara ett exempel på att en helt och hållet klassisk metod har föreslagits som alternativ till gängse bayesianska metoder. Det gäller en metod utvecklad av Alm & Thedéen (1996) för skattning av frekvenser för inledande händelser. I den bayesianska analysen används vanligen a posteriorifördelningens medelvärde som punktskattning. Därmed spelar det ingen roll huruvida något fel har inträffat eller inte. Detta implicerar emellertid att de bayesianska skattningarna inte är väntevärdesriktiga. I gammafallet erhålls t.ex. skattningen

$$\hat{\lambda} = \frac{x + \delta}{T}$$

som är väntevärdesriktig bara då $\delta = 0$ (dvs. endast detta värde är konsistent med ML- och MK-skattningarna). I nollfelsfallet ger detta

$$\hat{\lambda} = 0$$

vilket inte är konsistent med någon subjektiv sannolikhet (se kapitel 4). Frågan är vilket värde på δ som ska rekommenderas? Vaurio ger en *intuitivt tilltalande* motivering till att välja $\delta = \frac{1}{2}$, vilket motsvarar "ett halvt fel" i nollfelsfallet (se avsnitt 5.2.3). I detta fall får intuitionen teoretiskt stöd av Jeffreys regel. Alm & Thedéen visar emellertid på möjligheten att göra ett teoretiskt motiverat val av δ helt oberoende av bayesiansk teori.

Alm & Thedéens förslag bygger på att ett 50%-igt övre konfidensintervall bildas för λ i nollfelsfallet. Tolkningen är då att λ med lika stor sannolikhet överskattas som underskattas. En sådan skattning skulle enligt Alm & Thedéen kunna kallas *medianriktig*. Vad som söks är alltså det λ som uppfyller $P(X \leq x) = \frac{1}{2}$. Detta ger $\delta = \ln(2)$ vilket inte är konsistent med Jeffreys regel utan snarare ger en a priorifördelning proportionell mot $\lambda^{\ln(2)-1}$. Alm & Thedéens skattning är emellertid inte tänkt att användas för framställning av a priorifördelningar. (Alm & Thedéen 1996)

Alm & Thedéens metod används bl.a. av Relcon som "snabbmetod" för skattning av frekvenser för inledande händelser i fall med noll fel. [Relcon]

8 Synpunkter från användarna

I detta kapitel redovisas i första hand de intervjuresultat som varit vägledande i mitt val av ”testmetod”.³³

Mina tidiga intervjuer gällde bland annat möjligheten av att frånga de bayesianska metoderna till förmån för klassisk metodik. Frågan var då framförallt huruvida osäkerheter kan representeras med klassiska konfidensintervall istället för hela fördelningar. Det stora problemet ansågs därvidlag vara att hela fördelningar behövs för Monte Carlo-simulering i RiskSpectrum® [FKA, OKG, Relcon]. I PSA-modellerna analyseras dessutom komponentfel simultant med t.ex. operatörsfel [Jung] vars parametrar anses kräva erfarenhetsbaserade skattningar [Holmberg]. Av praktiska skäl, och för att erhålla ett konsistent analysresultat, ansågs det lämpligt att utnyttja bayesianska tekniker över hela linjen [Holmberg].

På frågan om hur precis en metod för skattning av T-bokens parametrar måste vara fick jag svaret att en avvikelse på upp till en faktor 10 är acceptabel. Ytterligare precision tenderar att drunkna i den stora mängd antaganden som görs för övriga PSA-parametrar; komponentfelta är bara en aspekt av PSA. I det avseendet är det rimliga kriteriet på en metod att den ska vara *tillräckligt bra* snarare än exakt. [Holmberg] Detta ansågs också rimligt med tanke på att rådata ofta inte håller den kvalitet som krävs för att göra en exakt metod rättvisa. I verkligheten är det inte alltid känt *varför* en komponent har felat. Ibland händer det att ett fel hänförs till bristfällig teknik när den verkliga orsaken är mänskligt felhandlande. [Jung]

Ytterligare frågor gällde användningen av T-boken. Relcon framhöll problem med kurvanpassning till T-bokens percentiler. I RiskSpectrum® kan t.ex. ingen standardfördelning (gamma-, lognormal-, beta- etc.) på ett tillfredsställande sätt göra reda för T-bokens medelvärden vid exakt anpassning till percentilerna [Relcon].³⁴ En liknande invändning kommer från Forsmarks Kraftgrupp:

”Medelvärde, 5%, 50% och 95 %-percentiler från T-boken används hos FKA för att försöka återskapa den kontinuerliga (gamma)fördelning som ligger till grund för tabellvärdena i T-boken. Det vore bra om istället gammafördelningens formfaktorer även ingick i T-boken för att slippa detta extra steg.” [FKA]

Sanningen här är väl att några sådana formfaktorer i nuläget inte finns tillgängliga eftersom percentilerna i T-boken inte återspeglar någon gammafördelning, och ingen annan känd fördelning heller (jmf. avsnitt 7.1.1.2). Samma problem ur en något annorlunda aspekt framkom i en intervju med personal på OKG:

”T-boken behöver förbättras främst avseende dess kvalitetssäkring. Det skall exempelvis gå att reproducera parameterskattningarna utifrån givna indata utan att man är beroende av ”svarta lådor” av olika slag.” [OKG]

Den ”svarta lådan” i detta avseende, och likväl det enda verktyget för att återskapa T-bokens fördelningar, är Pörns metod tillsammans med T-Code.

³³ Ifrågasvarande intervjuer har gjorts med företrädare för Forsmarks Kraftgrupp (FKA), Oskarshamns kärnkraftverk (OKG), och Relcon AB som utvecklade och marknadsför PSA-verktyget RiskSpectrum®. Dessutom Gunnar Jung på SwedPower AB (SwedPower AB heter sedan den 24 april 2006 Vattenfall Power Consultant AB) och Jan-Erik Holmberg på VTT (VTT har sitt säte i Finland och är norra Europas största forskningsorganisation) (<http://www.vtt.fi/vtt/index.jsp?lang=sv>).

³⁴ Här avses i första hand percentilerna i CD-versionen av T-boken och i filen *Tbokrisk*. Fördelningen är där i allmänhet representerad med 50-100 punkter (T6, s. 36).

Vidare efterfrågade Relcon enklare metoder för kontinuerlig uppdatering. Exempelvis för komponenter som inte finns med i T-boken. En uppfattning var också att T-boken borde uppdateras oftare. [Relcon]

Övriga frågor som ansågs viktiga gällde:

- Datas representativitet: Är feldata fortfarande representativa efter 20 år? [Holmberg]
Kan data som samlats in under "lugn drift" generaliseras till extraordinära situationer? [FKA].
- Gruppering av data. Hur ska data grupperas med avseende på felintensitet? I vilken utsträckning kan data från "andra anläggningar" användas? Holmberg ansåg att strävan bör vara att skatta parametrar för så stora grupper som möjligt [Holmberg]. Å andra sidan ansågs det finnas skillnader inom dagens grupper som motiverar en finare uppdelning. [Relcon]
- Klassificering av data: Är ett "kritiskt fel" alltid kritiskt? [Holmberg]. Rapporteras allt, eller finns "luckor" som måste fyllas ut? [FKA]. Vidare ansågs det inte helt klart vad som menas med "funktionshinder"; en trasig packning behöver t.ex. inte vara funktionshinder förrän byte sker [Relcon].

Slutligen tycks de flesta användare vara nöjda med Pörns " $q+\lambda$ -modell", bortsett från att själva valet av vilka komponenter som ska omfattas av den mer komplicerade modellen uppfattades som svårt [OKG]. Günter Becker, företrädare för RISA Sicherheitsanalys GmbH och tillika en framstående PSA-teoretiker, är emellertid av en annan uppfattning:

"In evaluations, we calculate probabilities per demand for a few components and failure rates for most, but not by a combined model. To do this, you should have to classify your events in time related events and demand related events. Apparently, Pörn does not do this, but still, he ends up with a lambda, and a q. This is, what I do not understand." [Becker]

Sammanfattningsvis kan sägas att intervjuresultaten graviterar kring frågor om *datas kvalitet* samt *tolkning och användning av T-bokens fördelningar*. Det är framförallt de senare frågorna som kommer att vara betydelsefulla i fortsättningen av denna rapport.

9 Diskussion och val av "testmetod"

I valet av en metod för implementering, tester och vidare jämförelser med Pörns metod finns väsentligen två möjliga vägar: Antingen väljs en redan befintlig "helhetslösning" eller så konstrueras en ny lösning utifrån de olika byggstenar som beskrivits i tidigare kapitel. Ett grundläggande krav är vidare att "lösningen" i någon mening ska vara enklare än Pörns metod.

Det finns flera tänkbara sätt att modifiera Pörns metod. Annorlunda val skulle kunna göras både med avseende på hyperprior, prior, likelihoodfunktion och gruppering av data. En möjlighet som får starkt stöd i rapporten så här långt är att använda en lognormalfördelning som prior istället för en gammafördelning. Lognormalfördelningen har goda egenskaper både beträffande framställningen av en hyperprior och den trunkering som kan bli nödvändig om hyperpriori är oäkta; lognormalfördelningen tycks på det hela taget bidra till att metodiken *förenklas*. Att gammafördelningen är konjugerad med Poissonfördelningen är något som Pörn

bara kan utnyttja i begränsad utsträckning eftersom priorn, genom medelvärdesbildningen, blir en *blandad fördelning*.³⁵

Det nyss beskrivna alternativet är väsentligen likvärdigt med att föreslå ZEDB-metoden istället för Pörns metod.³⁶ Även om metoderna inte skiljer sig nämnvärt med avseende på beräkningsresultat framstår ZEDB-metoden som den enklare av de två. Förutom de fördelar som kommer av lognormalfördelningen har ZEDB en enklare grupperingsprincip i den meningen att den leder till färre beräkningar.³⁷ Emellertid är det inte avgjort vilken grupperingsprincip som är bäst. Ett test för att avgöra denna fråga skulle kunna vara intressant inom ramarna för detta arbete. Möjligen är det inte helt i linje med syftet (som kräver att en ”metod” implementeras).

Andra alternativ som får starkt stöd i litteraturstudien och i intervjuerna är olika varianter av PEB. Pörn har själv erfarenhet av sådana ”lösningar” från sitt arbete med de tidiga T-böckerna. Hans skäl att överge dessa metoder har redan behandlats. Flera av de problem som Pörn ansåg sig behöva lösa med hjälp av tvåstegsmetodik har emellertid Vaurio försökt lösa ”inom det gamla paradigmet”. Enligt egen utsago har han lyckats med detta. En av huvudpoängerna med Vaurios metod är dessutom att den ska vara just *ett enklare alternativ* till de gängse tvåstegsmetoderna. Bland metodens främsta förtjänster kan nämnas de enkla beräkningarna samt att a posteriorifördelningens form är känd. Faktorer som ytterligare höjer eventuella testers förklaringsvärde är för det första att metoden tycks vara relativt okänd i branschen; de jag intervjuat har i vissa fall känt till metodens existens men inte dess förmenta fördelar. Dessutom kritiserades metoden så sent som 2005 för att vara behäftad med de svagheter som Vaurio redan tidigare tillbakavisat (Bunea et al. 2005). Till detta kommer att Vaurio presenterar sin metod på ”algoritmisk” form vilket betyder att inga kompletterande antaganden eller tolkningar behöver göras; metoden är ”ett färdigt paket”.

Av de nyss anförda skälen väljer jag Vaurios metod (PREB) för tester och vidare jämförelser med Pörns metod.

10 Tester

I detta kapitel testas och analyseras Vaurios metod med avseende på beräkningsresultat. Val av testobjekt (komponenter) och indataset har gjorts utifrån den benchmark (T-book – ZEDB Benchmark, 2004) där Pörns metod jämförs med ZEDB-metoden (i fortsättningen kallad ”ZEDB”).³⁸ Det finns väsentligen två skäl till detta val:

- De komponenter som testas i benchmarken täcker på ett bra sätt in *olika* realistiska situationer med avseende på empiriskt underlag och komponentslag.
- I ett test av detta slag (dvs. ett test med avseende på output i realistiska situationer) är metoderna ”varandras facit”. Förklaringsvärdet kan då antas öka med antalet metoder som testas parallellt. Genom att utgå från benchmarken kan Vaurios metod jämföras med både Pörns metod och ZEDB.

³⁵ Fördelen med att använda en gammafördelning i tvåstegsfallet är väsentligen att beräkningen av likelihood-funktionen i steg ett, dvs. $p(x|\theta)$ i (7.4), förenklas (se även appendix F.2).

³⁶ Här ska påpekas att lognormalfördelningen kan väljas i T-Code, liksom gammafördelningen i ZEDB. Mitt resonemang utgår från vilka alternativ som rekommenderas.

³⁷ Vilket i sin tur gör det lättare att undvika den dubbelräkning som Pörn, mot sina egna principer, tvingas göra (se avsnitt 7.1.1.6).

³⁸ Filer med in- och utdata har erhållits från TUD-kansliet.

Avsaknaden av facit innebär att uttalanden om metodens *reliabilitet* inte kan göras i absoluta termer. Emellertid finns här en möjlig ”bakväg”: Med ZEDB som referensmetod blir uttalanden om metodernas *inbördes likhet* meningsfulla. Detta tillsammans med det faktum att Pörns metod och ZEDB bedöms vara ”in excellent agreement” (T-book – ZEDB Benchmark, 2004, s. 12) ger viss information om hur stora avvikelser som är acceptabla.³⁹

Kompletterande tester har gjorts för ett antal ”extremfall” för att möjliggöra uttalanden om metodens robusthet, samt för validering och verifiering.

Anmärkning: I benchmarken använder Pörn en modifierad version av sin hyperprior (7.9). Modifieringen är gjord utifrån Meyer och Hennings (1999) rekommendationer och ger en *äkt* a posteriorifördelning. Hyperpriorerna har följande utseende:

$$p(\alpha, \beta) \propto \left[\alpha^{-1.45} \beta^{-0.8} \left(1 + \frac{\beta}{\alpha T} \right)^{-0.6} \right]^{-1}. \quad (\text{T-book – ZEDB Benchmark, 2004, s. 18})$$

Denna hyperprior används endast i benchmarken. [Pörn 2]

□

10.1 Genomförande och resultat

I benchmarken mellan Pörns metod och ZEDB testas tre komponenttyper med data gällande för den femte versionen av T-boken:

- Centrifugalpump, horisontell (T5, tab. 1.1.1, s. 76). Komponenten är kontinuerligt driftsatt med felmoden *obefogat stopp* och modellparametern λ (felintensitet under drift). För 34 komponenter har 38 fel observerats. Den totala drifttiden (samtliga anläggningar) är ca $2 \cdot 10^6$ h.
- Skalventil, pneumatisk (T5, tab. 3.20.1, s. 186). Komponenten är i standby-läge med felmoden *obefogad lägesändring* och modellparametern λ_s (felintensitet i standby). För 718 komponenter har 10 fel observerats. Den totala standby-tiden är ca $2 \cdot 10^7$ h.
- Dieselaggregat (T5, tab. 7.1.1, s. 286). Komponenten är i standby-läge med felmoden *obefogat stopp* och modellparametern λ (felintensitet under ”verklig” aktivering eller testkörning). För 48 komponenter har 70 fel observerats. Den totala drifttiden är ca $3,5 \cdot 10^4$ h.

Sammanfattningsvis kan sägas att ventilen representerar ett extremt litet antal fel i förhållande till antalet komponenter och drifttid. Pumpen och dieselaggregatet skiljer sig genom att dieselaggregatet har dubbelt så många fel men många gånger kortare drifttid.

För att underlätta jämförelser har Vaurios metod testats med både Pörns och ZEDB’s grupperingsprinciper. Vilken princip som har använts anges med index (”P” för Pörn och ”Z” för ZEDB). Gruppering enligt ZEDB innebär att hyperparametrarna beräknas utifrån *anläggningsspecifika data* (data grupperas anläggningsvis). Gruppering enligt Pörn innebär att hyperparametrarna beräknas utifrån *individspecifika data*. Den anläggningsspecifika fördelningen är sedan medelvärde av komponenternas fördelningar på respektive anläggning. För venti-

³⁹ Den eventuella felaktigheten i att på detta sätt ge Pörn och ZEDB tolkningsföreträde diskuteras inte här.

lerna har jag av praktiska skäl endast tillämpat gruppering enligt ZEDB. Totalt sett bör ändå en hygglig jämförelse kunna göras.

Eftersom ZEDB inte har någon " $q+\lambda$ -modell" omfattar benchmarken endast de två λ -modellerna (λ och λ_s). Denna avgränsning är relevant även för Vaurios metod.

Anmärkning: I benchmarken (T-book – ZEDB Benchmark, 2004) jämförs metoderna med avseende på data från TUD såväl som ZEDB. Emellertid har samma grupperingsprincip använts endast för ZEDB-data. Dessvärre har utdatatabeller för dessa körningar varit svåra att skaffa fram (i benchmarken presenteras resultaten endast grafiskt). För att underlätta jämförelser har jag därför, ventilerarna undantaget, testat Vaurios metod med bägge grupperingsprinciperna.

□

Vaurios metod har också testats för ett antal extremfall (noll fel, extrema varianser och få komponenter). För "nollfelstestet" användes ett autentiskt fall från T5:

- Skrupvpump, diesel dagtank (T5, tab. 1.11.1, s. 106). Komponentens är i intermittent läge med felmoden *obefogat stopp* och modellparametern λ . För 38 komponenter har 0 fel observerats. Den totala drifttiden är $8,2 \cdot 10^5$ h.

Ett test har gjorts för samma komponent i T6 där 2 fel har tillkommit (T6, tab. 1.11.1, s. 116). Testet kan ge upplysningar om hur respektive metod "reagerar" på ny information i fall då dataunderlaget från början är extremt magert. Drifttiden är något längre i T5 än i T6 vilket beror på att ny information om den verkliga drifttiden har kunnat utnyttjas i T6.

10.1.1 Implementering

Vaurios metod har implementerats i MATLAB med utgångspunkt från den "algoritm" som finns återgiven i appendix E. Då olika valmöjligheter funnits har jag i första hand följt Vaurios egna rekommendationer. I andra hand har jag valt det alternativ som framstått som enklast. Programkoden återfinns i appendix D.

10.1.2 Redovisning av resultat

För benchmarken presenteras utdata i form av percentiler och medelvärde för *anläggnings-specifika* respektive *generiska* fördelningar. De kompletta tabellerna återfinns i appendix A. Motsvarande indatatabeller återfinns i appendix B. De generiska fördelningarna presenteras separat i tabell 1-3 nedan. Generiska fördelningar för skrupvpumpen återfinns i tabell 4. Resultaten för Pörns metod är hämtade direkt från T-boken (T5 och T6). För att öka möjligheterna till jämförelser presenteras ML-skattningen jämte medelvärdet för de testade metoderna. I utdatatabellerna är anläggningar med noll fel skuggade.

I resultattabellerna skrivs Vaurios metod med apostrof (PREB'). Detta tydliggör att testet inte är "auktorisat" och att de resultat som redovisas här inte ska förväxlas med resultat från andra tester med Vaurios metod. Resultaten för Pörns metod och ZEDB är emellertid fullt autentiska.

10.2 Analys

10.2.1 Pörn, Vaurio och ZEDB

Jag har valt att jämföra Vaurios och Pörns metoder med avseende på följande faktorer:

- Övergripande skillnader mellan utdata för de olika komponentslagen. Här är de generiska fördelningarna väsentliga.
- Median och medelvärde för anläggningsspecifika fördelningar.
- Snävhets, dvs. avstånd mellan 5- och 95%-percentiler respektive 95%-percentil och medelvärde.
- Fördelningar för anläggningar med noll fel.
- Skillnader mellan fördelningar för anläggningar med få respektive många fel i samma population.
- Motsvarande för drifttid.

I samtliga dessa jämförelser kan ZEDB tjäna som ”yttre referenspunkt”. Generellt betraktar jag avvikelser inom en faktor 5 som ”små”.⁴⁰ Den kvantitativa bedömningen görs emellertid huvudsakligen utifrån en parallell jämförelse mellan PREB och Pörn respektive ZEDB och Pörn.

PREB_p visar upp de största avstånden mellan 5%- och 95%-percentiler både för dieseln (Tab. 3) och pumpen (Tab.1). För ventilen (Tab. 2) haltar jämförelsen eftersom gruppering endast har gjorts enligt ZEDB. Dock ger PREB_z i samtliga fall en bredare fördelning än ZEDB. Dessa resultat är i överensstämmelse med Vaurios egna resultat (se Vaurio & Jänkälä 2006, s. 215). Skillnaden ligger framförallt i skattningen av 5%-percentilen. 95%-percentilerna är i mycket god överensstämmelse oavsett metod och komponentgrupp (i samtliga fall mindre än en faktor 5). För dieseln och pumpen är både median och medelvärde i mycket god överensstämmelse oavsett metod (mindre än en faktor 1,5).

Testerna tyder inte på att PREB skulle ge snävrare fördelningar än övriga metoder. 5%-percentilen skattas bara lägre med Pörns metod i ventilfallet (Tab. 2). Dessutom är förhållandet mellan 95%-percentil och medelvärde i god överensstämmelse för Pörns metod och PREB.

De största skillnaderna gäller för ventilen. Medelvärde och 95%-percentil är i god överensstämmelse medan median och 5%-percentil är kraftigt avvikande för Pörns metod. Detta har rimligen med grupperingen att göra. Emellertid ger PREB_z en betydligt lägre 5%-percentil än ZEDB, trots samma gruppering, liksom en lägre median (en faktor 40).

⁴⁰ Detta mått är baserat på de beräkningar som Cooke et al. (1995, 2003) gjort (se avsnitt 6.2.2 och 7.1.1.5) tillsammans med den bedömning som görs i benchmarken mellan Pörns metod och ZEDB (T-book – ZEDB Benchmark, 2004). Cooke et al. noterar skillnader mellan lognormal- och gammafallet på upp till en faktor 5, medan ZEDB och Pörns metod anses vara ”in excellent agreement” i benchmarken. En möjlig tolkning av detta är att skillnader på upp till en faktor 5 är acceptabla. Denna gränsdragning är konservativ i relation till de precisionskrav som framkom i intervjuerna (se kapitel 8). Eftersom gränsdragningen kan göras på många olika sätt lämnar jag emellertid det slutgiltiga avgörandet åt läsaren, dels genom att ange avvikelser i absoluta värden, dels genom att presentera kompletta utdatatabeller.

Det är inte uppenbart vilken av ZEDB och PREB som är mest lik Pörns metod med avseende på generiska fördelningar. Vad gäller medelvärden är PREB (oavsett gruppering) och Pörn mest lika för dieseln, medan ZEDB och Pörn är mest lika för pumpen. Skillnaderna i dessa fall är emellertid extremt små, maximalt en faktor 1,2 (gäller ZEDB och Pörn i dieselfallet). Nästan lika små är skillnaderna mellan 95%-percentilerna. För de nedre percentilerna är skillnaderna något större, men där har också grupperingen större betydelse. En tendens är att PREB skattar 5%-percentilen något lägre. Vad gäller medianen avviker PREB endast marginellt från de övriga metoderna.

Bortsett från 5%-percentilen är metoderna överlag lika även när det gäller specifika anläggningar. Skillnaden mellan $PREB_p$ och Pörns metod ligger inom en faktor 5 i samtliga fall (skillnaden är oftast väsentligt mindre). Ett undantag är ventilen där medianen skattas betydligt lägre med Pörns metod. Här hamnar $PREB_z$ mellan Pörn och ZEDB.

Vad gäller fördelningar för anläggningar med noll fel (pump och ventil) ger metoderna mycket lika resultat för pumpen. Den enda uppenbara tendensen är att $PREB_p$ skattar 5%-percentilen en faktor 10 *lägre* och ZEDB en faktor 10 *högre* än Pörns metod. Resultatet är något annorlunda för ventilen: 95%-percentilerna är i god överensstämmelse liksom i de flesta fall medelvärdena. För fyra anläggningar skattar emellertid $PREB_z$ medelvärdet en faktor 10 *lägre* än både Pörn och ZEDB. Gemensamt för dessa anläggningar är att de har ca 10 gånger längre drifttid än övriga anläggningar (detta gäller F3, O2, O3 och T1). Varken Pörns metod eller ZEDB gör denna åtskillnad. För övrigt avviker $PREB_z$ mer från ZEDB ju lägre percentiler som skattas.

Vid en jämförelse mellan två anläggningar som har stor skillnad i antal fel men ungefär samma drifttid (t.ex. O1 och R1 i pumpfallet) presterar alla metoderna mycket lika. Allra närmast varandra ligger Pörn och $PREB_p$.

Ett rimligt krav är vidare att av två komponenter med samma antal fel ska den med *längst drifttid* avvika mest från den generiska fördelningen (jmf. t.ex. F3 och R1 i dieselfallet eller O2 och R2 i pumpfallet); en helt ny anläggning bör ju vara identisk med PVC. (Jmf. T6, s. 41) Samtliga metoder uppfyller detta krav vad gäller medelvärde och 95%-percentil. PREB ger preliminärt något tveetydiga resultat för medianen och "kastar om" anläggningarna för 5%-percentilen.

Generic	Pörn	ZEDB	$PREB'_p$	$PREB'_z$
5%	5,8000E-07	3,4592E-06	4,0649E-08	7,1243E-07
50%	1,3040E-05	1,4330E-05	7,6228E-06	1,1571E-05
95%	5,9230E-05	4,8048E-05	7,3668E-05	5,3351E-05
Mean	1,9400E-05	1,9145E-05	1,8483E-05	1,7287E-05

$$ML = 1,7925E-05$$

Tabell 1. Generiska fördelningar, centrifugalpump.

Generic	Pörn	ZEDB	PREB _z
5%	4,8330E-39	7,0850E-08	4,3000E-18
50%	1,2820E-17	4,1223E-07	9,6640E-09
95%	1,8920E-06	2,0018E-06	6,3629E-06
Mean	5,6520E-07	8,0156E-07	1,1016E-06

$$ML = 4,7773E-07$$

Tabell 2. Generiska fördelningar, skalventil.

Generic	Pörn	ZEDB	PREB _p	PREB _z
5%	4,8890E-04	4,7047E-04	6,4386E-05	3,7517E-04
50%	1,8210E-03	1,9270E-03	1,4669E-03	2,0775E-03
95%	5,7520E-03	7,4438E-03	7,5812E-03	6,3118E-03
Mean	2,2590E-03	2,7899E-03	2,3386E-03	2,5449E-03

$$ML = 1,9836E-03$$

Tabell 3. Generiska fördelningar, dieselaggregat.

10.2.2 Extremfall

För skruvpumpen (Tab. 4) är resultatet något mer svårtolkat eftersom ZEDB inte finns med som referensmetod. I T5-fallet där samtliga anläggningar har noll fel ger Pörns metod ”noll-skattningar” utom för 95%-percentilen och medelvärde. För de tre percentilerna ger PREB_p skattningar i storleksordningen 10^{-9} (5%), 10^{-7} (50%) respektive 10^{-6} (95%). PREB_p skattar därmed 95%-percentilen en faktor 10 högre än Pörn medan det (knappt) omvända gäller för medelvärde. I T6-fallet, då två anläggningar har varsitt fel, är medelvärdesskattningen en dryg 5 högre för PREB_p än för Pörn. ML-skattningen är i detta fall $1 \cdot 10^{-5}$. Pörns metod ger $3 \cdot 10^{-4}$ och PREB_p $1,7 \cdot 10^{-3}$. Pörn ligger alltså närmre ML-skattningen. 95%-percentilen skattas en faktor 2,7 högre med PREB_p.

I bägge dessa fall är fördelningarna rimligen mycket sneda. Medelvärde blir då mindre intressant eftersom det inte säger så mycket om var någonstans sannolikhetsmassan är samlad. I T5-fallet (noll fel) ger Pörns metod det för bayesianska metoder typiska resultatet, nämligen ett medelvärde som är större än 95%-percentilen. Med ett medelvärde som är en knapp faktor 4 mindre än 95%-percentilen förefaller PREB_p vara mer robust. Med en större 95%-percentil än Pörns tycks PREB_p dessutom vara den mer konservativa metoden. I T6-fallet är metoderna mer jämbördiga. Detta tyder på en snabb konvergens mellan metoderna då ny information tillkommer. PREB visar upp ett resultat som kan tyckas överraskande: I T5-fallet är de nedre percentilerna många gånger högre än i T6-fallet, trots att inga fel har observerats. Detta kan tänkas bero på att inte bara medelvärde utan även *variansen* ökar med ett fåtal observerade fel. I så fall är en minskande 5%-percentil fullt realistisk.

Vad gäller övriga extremfall kan jag bara konstatera att PREB ”klarar av” de fall som Cooke et al. (2003, s. 9) invänder mot. Här finns emellertid bara ML-metoden att jämföra med. I fallet med en enda anläggning, ett fel och 100 drifttimmar ger ML-skattningen 0,01 medan PREB ger 0,015. En ensam anläggning utan fel och 100 drifttimmar får medelvärde 0,005 vilket är Jeffreys δ dividerat med 100.

Generic	T5		T6	
	Pörn	PREB _p	Pörn	PREB _p
5%	0	2,5110E-09	0	0,0000E+00
50%	0	2,9052E-07	0	0,0000E+00
95%	2,4E-07	2,4531E-06	9,341E-04	2,6331E-03
Mean	9,2E-06	6,3859E-07	3,057E-04	1,7170E-03

ML = 1,1125E-05

Tabell 4. Generiska fördelningar, skruvpumpar.

10.2.3 Sammanfattning

Omdömet utifrån de generiska fördelningarna är att om Pörns metod och ZEDB är ”in excellent agreement” så gäller detta även Pörns metod och PREB. Om 5%-percentilen räknas bort är skillnaderna mycket små med avseende på både generiska och anläggningsspecifika fördelningar, även för anläggningar med noll fel. Att räkna bort 5%-percentilen är adekvat eftersom den sällan eller aldrig används som beslutsunderlag i PSA (Vaurio & Jänkälä 2006, s. 215). På det hela taget ger inte PREB snävare fördelningar än Pörns metod.

I ett avseende skiljer sig PREB markant från de övriga två metoderna: För ventilen, där de allra flesta komponenterna är felfria, gör PREB en större skillnad mellan anläggningar med lång respektive kort drifttid.

I extremfallen talar resultaten preliminärt för PREB även om ingen uttömmande jämförelse kan göras med mindre än att den verkliga fördelningen är känd. Övriga extremfallstester visar att PREB inte kollapsar då momenten är noll (vilket tillika inses genom en titt på Vaurios algoritm, alternativt min egen programkod).

Huruvida ZEDB eller PREB är mest lik Pörn är fortfarande en ganska öppen fråga. Inget i de tester som utförts här tyder på att Pörn och ZEDB skulle vara mer lika än Pörn och PREB. Med de mått på *likhet* som står till buds här tycks alltså PREB ge resultat som är i mycket god överensstämmelse med Pörns metod.

11 Avslutande diskussion

I kapitel 9 förde jag på teoretiska grunder fram Vaurios PREB som det intressantaste alternativet till Pörns metod, framförallt eftersom metoden uppvisar flera goda egenskaper med avseende på enkelhet. Mina tester visar att det finns skäl att tro att Vaurios PREB inte bara är ett enklare alternativ, utan dessutom *ett bra alternativ*.

De fördelar PREB har gentemot Pörns metod kommer huvudsakligen av att den bygger på konventionell PEB-metodik. Sammantaget ger rapporten följande argument för dylika metoder:

- Genom att hyperparametrarna punktskattas kan fördelarna med *konjugerade fördelningar* utnyttjas maximalt. I PREB är t.ex. a posteriorifördelningen en gammafördelning med kända parametrar vilket gör att den kan replikeras exakt. Detta underlättar vid överföring till användaren och vid eventuell kontroll av beräkningsresultat.
- PEB ger generellt sett enklare beräkningar än tvåstegsmetoder. Särskilt enkla blir beräkningarna med konjugerade fördelningar.
- PEB har en lång tradition inom kärnkraftsområdet och är relativt välkänd i PSA-kretsar.
- PEB-metodiken är teoretiskt enkel: De två ”stegen” är tydligt separerade (*punktskattning* respektive *Bayes sats*) vilket generellt ger mer transparenta matematiska uttryck, i synnerhet med användning av konjugerade fördelningar. Dessutom undviks den begreppsligt svårgenomskådliga *hyperprior*.

Samtliga dessa aspekter tillkommer PREB, med viss reservation för den sista; I PREB tycks en viss trade-off-effekt finnas mellan teoretisk enkelhet och robusthet, t.ex. med avseende på T^* dvs. den ”fria parametern” i variansskattningen. Vad som däremot talar starkt för PREB är att den saknar konventionella PEB-metoders mest uppenbara svaghet; PREB kollapsar inte pga. bristfälligt dataunderlag. Dessutom tycks PREB ge bredare fördelningar än tvåstegsalternativen, inte snävare såsom ofta hävdas. Likväl är de väsentliga percentilerna i god överensstämmelse med resultaten för tvåstegsalternativen.

I samband med PEB har jag även tagit upp kritik som rör *dubbelräkning* av data liksom den förment begreppsliga *inkonsistens* som är ett resultat av att klassisk och bayesiansk metodik blandas. Denna kritik är relevant även för PREB.

Vad jag förstår är debatten om dubbelräkningen till stor del en filosofisk debatt som inte avgörs i brådrasket. I mångt och mycket sammanfaller den med den traditionella kontroversen mellan frekventister och bayesianer (se kapitel 4): För en frekventist är det självklart att en individ ska räknas till ”sin egen grupp” och därmed ingå i superpopulationen. Denna ståndpunkt är vanlig inom PEB. För en bayesian är begreppet ”superpopulation” snarare *epistemiskt*: Superpopulationen ger den relevanta *förhandskunskapen* (a priori), medan uttalanden om en specifik individ bara kan göras i *efterhand* (a posteriori), dvs. efter uppdatering. Därmed kan individen inte heller räknas in i superpopulationen. Jag ska inte försöka fälla något avgörande här, och är inte heller säker på att det behövs. Oavsett vilken sida som har rätt gäller att dubbelräkningen ger mindre utslag ju större populationen är, och eftersom Vaurio grupperar komponentvis har han samma slags ”alibi” som Pörn i fallet med T-bokens generiska fördelningar (se avsnitt 7.1.1.6). Dessutom kan dubbelräkningen vara ett sätt att kringgå den

brist på robusthet som Pörn anser att PEB-metoder har i och med att de inte antar en tillräckligt flexibel prior (se avsnitt 7.1.1.1): Pörn utgår från att data inte dubbelräknas, vilket ger en fördelning som inte kan "fånga upp" information från den aktuella komponenten om den är starkt avvikande. Om data från den avvikande komponenten tas med vid framställning av den generiska fördelningen (dubbelräkning) elimineras detta problem, åtminstone i praktiken.

När det gäller den begreppsliga inkonsistensen argumenterar Vaurio själv för sin sak: Vaurio tycks mena att PREB är begreppsligt *konsistent* eftersom dess resultat kan ges en frekventistisk tolkning; att Bayes sats används för uppdatering gör inte metoden bayesiansk så länge skattningen utgår från "direkta observationer" (se avsnitt 7.2.1). PREB skulle såtillvida snarast vara en *klassisk metod*. Något som eventuellt ställer till det för Vaurio är den medelvärdesbias som används; om bias-termen motiveras med hjälp av Jeffreys regel är metoden inte uppenbart *icke-bayesiansk*. Å andra sidan: Om Vaurios metod kan betraktas som klassisk har han definitivt löst problemet med den begreppsliga inkonsistensen.⁴¹ Vad som då kvarstår är att ta ställning till om *datas fördelning kan betraktas som en fördelning för parametern*, dvs. om Vaurios perspektivskifte är möjligt. Detta avgör i sin tur huruvida klassiska fördelningar kan användas i konventionell PSA (med eller utan sällskap av "subjektivistiska" bayesianska fördelningar).

Så långt finns det goda skäl att tro att Vaurio har löst de problem som fick Pörn att lämna PEB-metodiken till förmån för tvåstegs bayesiansk metodik (dvs. punkt a-c i avsnitt 7.1.1.1). Vad som kan förefalla besvärligt är att några av frågorna kräver ett *ställningstagande*.

Ett pragmatiskt alternativ är att avfärda allt tal om dubbelräkning och begreppslig inkonsistens som överflödigt i den mån metoden ger rimliga resultat. Sådantillvida får PREB stöd av de tester som utförts här. Mina resultat är emellertid mycket preliminära. För ett slutligt avgörande krävs ytterligare studier av PREB, liksom jämförelser med flertalet alternativa metoder. I en sådan studie kan det visa sig att ingen metod är bäst i alla avseenden. T.ex. kunde det visa sig att Pörns " $q+\lambda$ -modell" är överlägsen alla alternativ (q och λ_s) som har behandlats här. I så fall kanske Pörns metod ska användas med avseende på behovsrelaterade fel, jämte en enklare metod för tidsrelaterade fel, t.ex. PREB. Vidare erbjuder PREB i sig flera möjligheter till variation. T.ex. kan bias-termerna varieras utan att detta gör våld på Vaurios principer. Ett intressant alternativ till Vaurios δ är den "nollpunktsskattning" som föreslås av Alm & Thedéen. Bägge approacherna är förövrigt intressanta som alternativ till den medelvärdesskattning som används i T-boken, särskilt för mycket sneda fördelningar.

Att Vaurio inte har någon " $q+\lambda$ -modell" kan i sig ses som en svaghet. Här har Vaurio emellertid ett motargument, nämligen att en sådan modell är omöjlig så länge inte de två typerna av fel kan separeras, och någon sådan separation gör inte Pörn. Detta liknar den kritik som Becker framför (se kapitel 8). Min behandling av " $q+\lambda$ -modellen" är emellertid alltför knapphändig för att några slutsatser ska kunna dras.

Slutligen gäller att ju större tillgången på data är, desto enklare metoder kan användas. Därvidlag kan det visa sig att även Vaurios metod är onödigt avancerad. Vad som först måste undersökas är emellertid hur länge en observation kan anses vara relevant, dvs. hur gamla data som får användas, liksom hur stora populationer som kan skapas utan att viktig information går förlorad. Likväl är det viktigt att data håller erforderlig kvalitet. En så pass sofistike-

⁴¹ Sådantillvida som dubbelräkning "får göras" i klassiska metoder kan även det problemet anses vara löst, även om bayesianen fortfarande kan invända att *alla klassiska metoder är felaktiga*.

rad och genomtänkt metod som Pörns tenderar att bli en ”eld för kråkor” om data håller låg kvalitet. Även om mina resultat skulle visa sig vara ohållbara kan alltså PREB komma ifråga som ett *tillräckligt bra alternativ*.

12 Slutsatser

Vaurios PREB är *enklare* än Pörns eftersom den bygger på konventionell PEB-metodik, vilket väsentligen innebär att *parametrarna i den generiska fördelningen punktskattas*. Genom utnyttjande av konjugerade fördelningar (gamma-Poisson) kan uppdateringen sedan göras på enklaste sätt. Vidare är både generiska och specifika fördelningar fullständigt kända vilket gör att de kan replikeras exakt (vid användning eller kontroll av beräkningsresultat). Jämfört med konventionella PEB-metoder har PREB fördelen av att fungera även för nollvärda moment (dvs. då medelvärde och/eller varians är noll).

De tester som gjorts i denna rapport pekar dessutom mot att PREB är ett *bra* alternativ till Pörns metod i den meningen att metoderna ger liknande resultat. För att uttömmande besvara frågor om PREB’s precision och robusthet krävs emellertid ytterligare tester och teoretiska jämförelser.

13 Förslag till fortsatt arbete

Denna rapport lämnar en hel del till övers för fortsatta efterforskningar. Förutom prövning och utvärdering av resultaten måste vissa kompletterande studier göras, i första hand beträffande PREB’s robusthet och precision. Vissa intressanta frågor, ibland av mycket grundläggande karaktär, har likväl bara kunnat behandlas i förbigående. Dessa frågor är i många fall intressanta oavsett om en alternativ metod väljs eller inte.

- För att avgöra vilken metod som uppvisar de bästa egenskaperna med avseende på robusthet och precision bör Vaurios PREB testas tillsammans med flera andra metoder. Urvalet kan göras utifrån denna rapport. Därtill bör metoderna ytterligare jämföras med avseende på teoretisk och praktisk enkelhet. Om ingen metod visar sig vara bäst i alla avseenden måste möjligheten att använda flera metoder parallellt utvärderas. Likväl är det relevant med en undersökning av möjligheterna att modifiera Vaurios metod. En sådan undersökning utgår lämpligen ifrån en utvärdering av andra moderna PEB-metoder.
- Pörns ” $q+\lambda$ -modell” har i denna rapport bara behandlats mycket översiktligt. Viss kritik har redovisats vad gäller möjligheten att separera tids- och behovsrelaterade fel. Denna kritik bör bemötas med reda bevis så länge Pörns metod används. Om en alternativ metod är aktuell bör det utredas huruvida ” $q+\lambda$ -modellen” är bättre än motsvarande modell hos den alternativa metoden och, om så är fallet, huruvida ” $q+\lambda$ -modellen” kan adopteras.
- Litteraturstudien bör kompletteras med en utvärdering av olika NPEB-metoder (Non-parametric Empirical Bayes). En intressant NPEB, baserad på en Dirichletfördelning, utvärderas hos Cooke et al. (2003) (se avsnitt 7.2.2).
- En viktig fråga gäller huruvida det största osäkerhetsbidraget kommer från variation mellan komponenter (variation i λ) eller från den stokastiska processen (variation i x). Om det visar sig att variationen väsentligen kommer av att *observationerna uppträder slumpmässigt* kan komponenterna betraktas som homogena (med avseende på λ) vil-

ket i sin tur motiverar större populationer. Att skapa så stora populationer som möjligt är alltid eftersträvansvärt, men särskilt avgörande vid övervägande av *enkla* metoder.

- Ett annat sätt att skapa större dataunderlag är att använda gamla data. Först måste det emellertid nog undersökas *hur gamla* data som kan användas. Utvärderingen kan göras med något slags trendanalys; en möjlighet är att anpassa en tidsberoende modell till synbara förändringar. Det kan också visa sig att utvecklingen tar ”språng” pga. att större förändringar skett på kort tid (t.ex. vid införande av ny teknik). I så fall bör data bortom en viss tidsgräns klassificeras som ”alltför gamla” och således uteslutas.
- Vikten av minutiös insamling och klassificering av data kan inte nog understrykas. Problemen med bristande datakvalitet minskar inte nödvändigtvis med användning av bayesianska metoder. Om en alltför stor tilltro sätts till möjligheterna att kompensera för bristande datakvalitet, eller om metoderna används slentrianmässigt, kan problemen istället fördjupas. Den ledande principen bör vara att *mätbara storheter ska mätas i den mån det är praktiskt möjligt*. Likväl bör compensation för bristande datakvalitet med hjälp av a prioriantaganden betraktas som en högst tillfällig lösning. I de fall då subjektiva antaganden kompletterar rådata bör dessa modelleras explicit. Att Pörns metod är bayesiansk betyder inte att skattningar kan göras ”på en höft”, då blir antaganden icke spårbara och skattningsprocessen degenererar.

14 Referenser

14.1 Tryckta referenser

Alm, S. E. & Thedéen, T. (1996). *Prediktering av frekvenser för inledande händelser*, Centrum för säkerhetsforskning, KTH.

Apostolakis, G., Kaplan, S., Garrick, B. J., Dughy, R. J. (1980). ”Data specialization for plant specific risk studies”, *Nuclear Engineering and Design*, Vol. 56.

Atwood, C. L., LaChance, J. L., Martz, H. F., Anderson, D. J., Engelhardt, M., Whitehead, D. et al. (2003). *Handbook of Parameter Estimation for Probabilistic Risk Assessment*, NUREG/CR-6823, SAND2003-3348P.

Atwood, C. L. (1986). ”Correct Bayesian and frequentist intervals are similar”, *Nuclear Engineering and Design*, Vol. 93.

Andersson, S. (1993). *Analys bakgrundsdata T-boken m a p tidsberoende*, Vattenfallrapport PT-28/93, PTS-SOA/KA50A.

Blom, G. (1984). *Sannolikhets teori med tillämpningar*, 2:a uppl., ISBN 91-44-00323-4, Studentlitteratur, Lund.

Blom, G. & Holmquist, B. (1998). *Statistik teori med tillämpningar*, 3:e uppl., ISBN 91-44-04372-4, Studentlitteratur, Lund.

Cooke, R. M., Bunea, C., Charitos, T., Mazzuchi, T. A. (2003). *Mathematical Review of ZEDB Two-Stage Bayesian Model*, VGB PowerTech e.V., ISSN 1439-7498, Essen.

- Bunea, C., Charitos, T., Cooke, R. M., Becker, G. (2005). "Two-stage Bayesian models – application to ZEDB project", *Reliability Engineering and System Safety*, Vol. 90.
- Carlin, B. P. & Louis, T. A. (2000). *Bayes and Empirical Bayes Methods for Data Analysis*, 2nd ed., ISBN 58488-170-4, Chapman and Hall, New York.
- Cooke, R. M., Dorrepaal, J., Bedford, T. (1995). *Review of SKI Data Processing Methodology*, SKI Report 95:2, ISSN 1104-1374, ISRN SKI-R—95/2—SE.
- Englund, G. (2000). *Introduktion till datorintensiva metoder*, Matematiska institutionen, Avdelningen för matematisk statistik, KTH.
- Heising, C. D. & Shwayri, N. (1987). "An application of the two-stage Bayesian method to estimate the outage rate of oil-fired boilers", *Reliability Engineering*, Vol. 18.
- Hofer, E. (1999). "On two-stage Bayesian modelling of initiating event frequencies and failure rates", *Reliability Engineering and System Safety*, Vol. 66.
- Hofer, E., Hora, S. C., Iman, R. L., Peschke, J. (1997). "On the solution approach for Bayesian modelling of initiating event frequencies and failure rates", *Risk Analysis*, Vol. 17, No. 2.
- Kaplan, S. (1983). "On a 'two-stage' Bayesian procedure for determining failure rates from experimental data", *IEEE Transactions on Power Apparatus and Systems*, Vol. PAS-102.
- Kaplan, S. (1985). "The two-stage Poisson-type problem in probabilistic risk analysis", *Risk Analysis*, Vol. 5, No. 3.
- Kaplan, S. (1986). "On the use of data and judgement in probabilistic risk and safety analysis", *Nuclear Engineering and Design*, Vol. 93.
- Leonard, T. & Hsu, J. S. J. (1999). *Bayesian Methods – An Analysis for Statisticians and Interdisciplinary Researchers*, ISBN 0-521-59417-0, Cambridge University Press, Cambridge.
- Meyer, W. & Hennings, W. (1999). "Prior distributions in two-stage Bayesian estimation of failure rates", *Safety and Reliability*, Schueller & Kafka (eds), Balkema, Rotterdam.
- Nyman, R. (1995). *The T-book seminar 1995-01-27 in Stockholm*, SKI/RA-001/95.
- Prediction Technologies (2002). *R-DAT and R-DAT Plus 1.5 – User's Manual*, Prediction Technologies Inc. (www.prediction-technologies.com).
- Pörn, K. (1990). *On Empirical Bayesian inference Applied to Poisson Probability Models*, Linköping Studies in Science and technology, Dissertation No. 234, ISBN 91-7870-696-3, Linköping University, Linköping.
- Pörn, K. (1996). "The two-stage Bayesian method used for the T-Book application", *Reliability Engineering and System Safety*, vol. 51, No. 2.
- Råde, L. & Westergren, B. (1998). *Mathematics Handbook for Science and Engineering*, 4:e uppl., ISBN 91-44-00839-2, Studentlitteratur, Lund.

Siu, N. O. & Kelly, D. L. (1998). "Bayesian parameter estimation in probabilistic risk assessment", *Reliability Engineering and System Safety*, Vol. 62.

T1: T-boken (1982). *Tillförlitlighetsdata för komponenter i svenska kokvattenreaktorer*, Version 1, Utarbetad av Rådet för Kärnkraftssäkerhet (RKS), Asea-Atom AB och Studsvik Energiteknik AB, Rapport RKS 1985-05.

T2: T-boken (1985). *Tillförlitlighetsdata för komponenter i svenska kraftreaktorer*, Version 2, Utarbetad av Rådet för Kärnkraftssäkerhet (RKS), Asea-Atom AB och Studsvik Energiteknik AB, Rapport RKS 1982-07.

T5: T-boken (2000). *Tillförlitlighetsdata för komponenter i nordiska kraftreaktorer*, Version 5, Utarbetad av TUD-kansliet och Pörn Consulting, ISBN 91-630-9862-8.

T6: T-boken (2005). *Tillförlitlighetsdata för komponenter i nordiska kraftreaktorer*, Version 6, Utarbetad av TUD-kansliet och Pörn Consulting, ISBN 91-631-7231-3.

T-book – ZEDB Benchmark (2004). *Calculation of Reliability Data using Two-Stage Bayesian Methods*, VGB PowerTech e.V., ISSN 1439-7498.

Vaurio, J. K. (1987). "On analytic empirical Bayes estimation of failure rates", *Risk Analysis*, Vol. 7, No. 3.

Vaurio, J. K. (1990).

- a) "Objective prior distributions and Bayesian updating", *Reliability Engineering and System Safety*, Vol. 35.
- b) "Technical Note – On the Meaning of Probability and Frequency", *Reliability Engineering and System Safety*, Vol. 28.

Vaurio, J. K. & Jänkälä, K. (2006). "Evaluation and comparison of estimation methods for failure rates and probabilities", *Reliability Engineering and System Safety*, Vol. 91.

ZEDB (2002). *Reliability Data for Nuclear Power Plant Components: Analysis for 2002*, VGB PowerTech e.V., ISSN 1439-7498, Essen.

14.2 Intervjuer och korrespondens

För muntliga intervjuer anges det datum då intervjun genomfördes. För skriftliga intervjuer anges det datum då svar mottogs.

[Becker] Korrespondens per e-post med Günter Becker (RISA Sicherheitsanalysen GmbH, Tyskland) (1/3 2006)

[FKA] Skriftlig intervju med personal från Forsmarks Kraftgrupp AB, företrädare av Johan Andersson (13/12 2005).

[Holmberg] Muntlig intervju med Jan-Erik Holmberg (VTT) (25/11 2005).

[Jung] Muntlig intervju med Gunnar Jung (SwedPower AB) (5/1 2006).⁴²

[Relcon] Muntlig intervju med Ola Bäckström och Per Hellström (Relcon AB) (17/11 2005).

[OKG] Skriftlig intervju med personal från OKG AB (Oskarshamns kärnkraftverk) företrädare av Pär Lindahl (23/11 2005).

[Pörn] Skriftliga intervjuer med Kurt Pörn (Pörn Consulting):

1. (31/1 2006).
2. (1/3 2006).

⁴² SwedPower AB heter sedan den 24 april 2006 Vattenfall Power Consultant AB.

15 Appendix

Appendix A: Jämförelsetabeller

I tabellerna nedan anges 5-, 50- och 95%-percentilerna i a posteriorifördelningen $p(\lambda|x)$ för de olika anläggningarna i TUD-systemet, samt motsvarande generiska fördelningar. Beräknade värden ligger till grund för den jämförelse av metoder som görs i avsnitt 10.2.1.

Utdata för Pörns metod och ZEDB är hämtade från *T-book – ZEDB Benchmark* (2004). Indexering av Vaurios metod ($PREB_p$ respektive $PREB_z$) anger gruppering enligt Pörn respektive ZEDB (komponent- respektive anläggningsvis gruppering). Skuggade rader avser anläggningar med noll fel.

Motsvarande indata presenteras i appendix B och gäller för T-boken version 5 (T5). Fullständiga in- och utdatafiler har erhållits från TUD-kansliet.

A.1

Centrifugalpump (T5, 1.1.1)

Antal komponenter: 34

Antal fel: 38

Total drifttid: 2119985 h

5%	Pörn	ZEDB	PREB _p	PREB _z
B1	5,7100E-07	5,4940E-06	2,9049E-06	4,2700E-06
B2	1,9940E-07	2,2809E-06	1,3441E-08	2,3000E-07
F1	8,9030E-06	1,4253E-05	1,1569E-05	1,6750E-05
F2	3,9400E-06	1,1834E-05	1,0993E-05	1,3180E-05
F3	2,1440E-07	2,3342E-06	1,3926E-08	2,5000E-07
O1	2,3210E-07	3,0098E-06	1,0991E-06	1,2900E-06
O2	5,7040E-06	8,2068E-06	4,6533E-06	7,9500E-06
O3	2,1120E-07	2,3231E-06	1,3642E-08	2,4000E-07
R1	3,8740E-06	2,1294E-05	2,3251E-05	2,6730E-05
R2	1,1840E-06	6,1710E-06	2,0845E-06	5,3600E-06
R3	5,8320E-07	3,8481E-06	1,0400E-06	2,0100E-06
R4	2,1410E-07	2,3334E-06	1,3798E-08	2,5000E-07
T1	7,4510E-07	6,8691E-06	3,5169E-06	6,1500E-06
T2	6,6830E-07	4,9360E-06	1,5704E-06	3,6000E-06
Generic	5,8000E-07	3,4592E-06	4,0649E-08	7,1243E-07

50%	Pörn	ZEDB	PREB _p	PREB _z
B1	1,3270E-05	1,4406E-05	1,2338E-05	1,4250E-05
B2	6,3480E-06	9,0016E-06	2,4563E-06	3,7800E-06
F1	2,8890E-05	2,8830E-05	3,3135E-05	3,4170E-05
F2	2,4720E-05	2,4493E-05	3,1051E-05	2,8770E-05
F3	6,6300E-06	9,2521E-06	2,5448E-06	4,0200E-06
O1	8,0500E-06	9,4311E-06	8,2542E-06	6,3600E-06
O2	1,8040E-05	1,7598E-05	1,7979E-05	1,9010E-05
O3	6,5540E-06	9,1997E-06	2,4929E-06	3,9700E-06
R1	3,6640E-05	4,0287E-05	4,5959E-05	4,6330E-05
R2	1,2700E-05	1,3242E-05	1,1101E-05	1,2810E-05
R3	1,0820E-05	1,2062E-05	8,4728E-06	9,9400E-06
R4	6,6140E-06	9,2481E-06	2,5214E-06	4,0100E-06
T1	1,4200E-05	1,5984E-05	1,6812E-05	1,6740E-05
T2	1,2540E-05	1,2970E-05	1,2513E-05	1,2020E-05
Generic	1,3040E-05	1,4330E-05	7,6228E-06	1,1571E-05

95%	Pörn	ZEDB	PREB _p	PREB _z
B1	4,3210E-05	3,0749E-05	3,966E-05	3,3930E-05
B2	2,7940E-05	2,1835E-05	2,357E-05	1,7440E-05
F1	6,8890E-05	5,6200E-05	7,346E-05	6,0850E-05
F2	7,9060E-05	4,8588E-05	7,005E-05	5,3550E-05
F3	2,8600E-05	2,2555E-05	2,442E-05	1,8520E-05
O1	3,7330E-05	2,0488E-05	3,116E-05	1,8350E-05
O2	4,1960E-05	3,3957E-05	4,613E-05	3,7450E-05
O3	2,8180E-05	2,2411E-05	2,392E-05	1,8290E-05
R1	1,1270E-04	7,0718E-05	8,432E-05	7,3780E-05
R2	3,7680E-05	2,4294E-05	3,657E-05	2,5230E-05
R3	3,7740E-05	2,7322E-05	3,6E-05	2,8700E-05
R4	2,8410E-05	2,2544E-05	2,419E-05	1,8500E-05
T1	6,5180E-05	3,2200E-05	5,856E-05	3,5610E-05
T2	5,1350E-05	2,6931E-05	5,16E-05	2,8600E-05
Generic	5,9230E-05	4,8048E-05	7,3668E-05	5,3351E-05

Mean	Pörn	ZEDB	PREB _p	PREB _z	ML
B1	1,6360E-05	1,6121E-05	1,5650E-05	1,6040E-05	1,5530E-05
B2	9,2000E-06	1,0521E-05	5,9244E-06	5,6500E-06	0,0000E+00
F1	3,2540E-05	3,1832E-05	3,6599E-05	3,5880E-05	4,2959E-05
F2	3,1040E-05	2,6902E-05	3,4549E-05	3,0470E-05	3,5437E-05
F3	9,4960E-06	1,0856E-05	6,1379E-06	6,0000E-06	0,0000E+00
O1	1,2060E-05	1,0677E-05	1,1166E-05	7,6400E-06	5,0463E-06
O2	2,0200E-05	1,9322E-05	2,0716E-05	2,0380E-05	2,1248E-05
O3	9,3700E-06	1,1032E-05	6,0128E-06	5,9300E-06	0,0000E+00
R1	4,5820E-05	4,2977E-05	4,8849E-05	4,7780E-05	5,7036E-05
R2	1,5180E-05	1,4141E-05	1,4143E-05	1,3720E-05	1,3103E-05
R3	1,3870E-05	1,3794E-05	1,2193E-05	1,1940E-05	9,2975E-06
R4	9,4490E-06	1,0910E-05	6,0815E-06	5,9900E-06	0,0000E+00
T1	2,0920E-05	1,7605E-05	2,2080E-05	1,8270E-05	1,8598E-05
T2	1,7570E-05	1,4392E-05	1,7723E-05	1,3520E-05	1,2294E-05
Generic	1,9400E-05	1,9145E-05	1,8581E-05	1,8483E-05	1,7287E-05

A.2

Skalventil (T5, 3.20.1)

Antal komponenter: 718

Antal fel: 10

Total drifttid: 20932445 h

5%	Pörn	ZEDB	PREB _z
B1	4,7770E-39	5,9450E-08	0,0000E+00
B2	4,9510E-39	1,4499E-07	8,6006E-08
F1	4,8390E-39	1,0388E-07	3,0166E-08
F2	4,9110E-39	2,7657E-07	3,7221E-07
F3	4,7700E-39	4,5014E-08	0,0000E+00
O1	4,8190E-39	6,5053E-08	1,0000E-18
O2	4,7790E-39	5,3754E-08	0,0000E+00
O3	4,7720E-39	5,1318E-08	0,0000E+00
R1	4,8430E-39	1,0772E-07	3,3161E-08
R2	5,9750E-39	3,0325E-07	1,7466E-06
R3	4,7730E-39	6,6732E-08	1,0000E-18
R4	4,7720E-39	6,6114E-08	1,0000E-18
T1	4,7690E-39	5,1721E-08	0,0000E+00
T2	4,8500E-39	1,9819E-07	1,9131E-07
Generic	4,8330E-39	7,0850E-08	4,3000E-18

50%	Pörn	ZEDB	PREB _z
B1	7,7070E-18	3,3828E-07	1,0400E-09
B2	3,4900E-17	5,4697E-07	9,4919E-07
F1	1,3320E-17	3,9775E-07	3,3292E-07
F2	1,9900E-17	7,4398E-07	1,1856E-06
F3	7,0270E-18	2,3282E-07	2,2237E-10
O1	1,1490E-17	3,7441E-07	2,2384E-09
O2	7,8400E-18	2,9838E-07	5,4972E-10
O3	7,2070E-18	2,8069E-07	4,2666E-10
R1	1,3740E-17	4,1152E-07	3,6597E-07
R2	1,3200E-13	1,2072E-06	7,8100E-06
R3	7,3520E-18	3,8527E-07	3,0338E-09
R4	7,2320E-18	3,8141E-07	2,7071E-09
T1	6,9120E-18	2,8361E-07	4,4415E-10
T2	1,4360E-17	5,8930E-07	8,5547E-07
Generic	1,2820E-17	4,1223E-07	9,6640E-09

95%	Pörn	ZEDB	PREB _z
B1	1,0980E-06	1,0997E-06	6,8474E-07
B2	3,3970E-06	1,9536E-06	3,8104E-06
F1	1,9220E-06	1,0655E-06	1,3365E-06
F2	2,9790E-06	1,9168E-06	2,7523E-06
F3	1,0500E-06	6,3435E-07	1,4641E-07
O1	1,6880E-06	1,3856E-06	1,4738E-06
O2	1,1190E-06	8,8616E-07	3,6194E-07
O3	1,0600E-06	8,0918E-07	2,8092E-07
R1	1,9860E-06	1,1237E-06	1,4691E-06
R2	4,5530E-05	1,0319E-05	2,1515E-05
R3	1,0680E-06	1,5041E-06	1,9975E-06
R4	1,0620E-06	1,4601E-06	1,7824E-06
T1	1,0440E-06	8,2108E-07	2,9244E-07
T2	2,0940E-06	1,6028E-06	2,3566E-06
Generic	1,8920E-06	2,0018E-06	6,3629E-06

Mean	Pörn	ZEDB	PREB _z	ML
B1	3,2560E-07	4,4281E-07	1,1854E-07	0,0000E+00
B2	8,1560E-07	7,5466E-07	1,3182E-06	1,3466E-06
F1	4,7940E-07	4,8007E-07	4,6236E-07	4,3533E-07
F2	9,0430E-07	8,8590E-07	1,3247E-06	1,3343E-06
F3	3,0570E-07	2,8386E-07	2,5347E-08	0,0000E+00
O1	4,8220E-07	5,3205E-07	2,5514E-07	0,0000E+00
O2	3,2970E-07	3,7027E-07	6,2660E-08	0,0000E+00
O3	3,1080E-07	3,4855E-07	4,8634E-08	0,0000E+00
R1	5,0120E-07	5,0001E-07	5,0826E-07	4,8056E-07
R2	7,0440E-06	2,7523E-06	9,2206E-06	1,5223E-05
R3	3,1500E-07	5,7141E-07	3,4581E-07	0,0000E+00
R4	3,1150E-07	5,5845E-07	3,0857E-07	0,0000E+00
T1	3,0240E-07	3,5132E-07	5,0627E-08	0,0000E+00
T2	5,9020E-07	7,0401E-07	1,0100E-06	1,0055E-06
Generic	5,6520E-07	8,0156E-07	1,1016E-06	4,7773E-07

A.3

Dieselaggregat (T5, 7.1.1)

Antal komponenter: 48

Antal fel: 70

Total drifttid: 35290 h

5%	Pörn	ZEDB	PREB _p	PREB _z
B1	6,7360E-04	1,1780E-03	7,9533E-04	1,2700E-03
B2	6,1810E-04	9,0180E-04	5,2655E-04	9,4271E-04
F1	1,5190E-04	3,3073E-04	1,1088E-04	2,5558E-04
F2	2,7070E-04	6,1341E-04	2,6935E-04	5,8443E-04
F3	2,3760E-04	6,5864E-04	3,3796E-04	6,3731E-04
O1	1,7200E-03	2,3311E-03	1,8378E-03	2,4910E-03
O2	1,3400E-03	1,9017E-03	1,4209E-03	2,0657E-03
O3	6,6340E-04	1,2818E-03	5,8325E-04	1,3616E-03
R1	4,4680E-04	1,2636E-03	5,7919E-04	1,3557E-03
R2	5,4520E-04	2,9038E-03	1,5852E-03	3,1003E-03
R3	7,3650E-04	1,4918E-03	6,7559E-04	1,5941E-03
R4	2,8610E-04	1,0484E-03	5,2918E-04	1,1093E-03
T1	9,8940E-05	2,1975E-04	5,8516E-05	1,2954E-04
T2	1,3880E-04	3,3589E-04	1,3571E-04	2,6105E-04
Generic	4,8890E-04	4,7047E-04	6,4386E-05	3,7517E-04

50%	Pörn	ZEDB	PREB _p	PREB _z
B1	2,3690E-03	2,5962E-03	2,5729E-03	2,8117E-03
B2	2,0630E-03	2,1225E-03	2,0776E-03	2,2977E-03
F1	1,0650E-03	8,2448E-04	7,2938E-04	7,1580E-04
F2	1,4600E-03	1,2965E-03	1,2292E-03	1,2939E-03
F3	1,4720E-03	1,3127E-03	1,3193E-03	1,3139E-03
O1	3,5240E-03	4,0993E-03	4,0496E-03	4,2241E-03
O2	3,1420E-03	3,6765E-03	3,6409E-03	3,8527E-03
O3	2,2010E-03	2,4974E-03	2,3023E-03	2,6556E-03
R1	2,2080E-03	2,6056E-03	2,3215E-03	2,7951E-03
R2	3,2390E-03	5,2505E-03	4,1946E-03	5,2574E-03
R3	2,3470E-03	2,8083E-03	2,4856E-03	2,9732E-03
R4	1,9350E-03	2,2808E-03	2,0362E-03	2,4559E-03
T1	8,4170E-04	6,1918E-04	5,0228E-04	4,5446E-04
T2	1,0870E-03	8,3784E-04	8,4647E-04	7,3111E-04
Generic	1,8210E-03	1,9270E-03	1,4669E-03	2,0775E-03

95%	Pörn	ZEDB	PREB _p	PREB _z
B1	5,6890E-03	5,1736E-03	6,1710E-03	5,2803E-03
B2	4,8620E-03	4,3872E-03	5,4490E-03	4,5791E-03
F1	2,9120E-03	1,6508E-03	2,5751E-03	1,5476E-03
F2	3,6990E-03	2,3807E-03	3,5963E-03	2,4300E-03
F3	3,7790E-03	2,3115E-03	3,5529E-03	2,3568E-03
O1	6,5520E-03	6,7537E-03	7,5821E-03	6,6216E-03
O2	6,4070E-03	6,5921E-03	7,4878E-03	6,4600E-03
O3	5,2670E-03	4,4614E-03	6,0393E-03	4,5907E-03
R1	5,6860E-03	4,8900E-03	6,3708E-03	5,0135E-03
R2	9,6940E-03	8,8232E-03	9,2365E-03	8,2412E-03
R3	5,4010E-03	4,8841E-03	6,3543E-03	4,9853E-03
R4	5,9200E-03	4,4483E-03	5,8799E-03	4,6121E-03
T1	2,5100E-03	1,3216E-03	2,0424E-03	1,1102E-03
T2	3,2930E-03	1,6772E-03	2,8824E-03	1,5807E-03
Generic	5,7520E-03	7,4438E-03	7,5812E-03	6,3118E-03

Mean	Pörn	ZEDB	PREB _p	PREB _z	ML
B1	2,6640E-03	2,8710E-03	2,9090E-03	2,9829E-03	3,2206E-03
B2	2,3140E-03	2,3503E-03	2,4137E-03	2,4688E-03	2,4233E-03
F1	1,2370E-03	9,0273E-04	9,5614E-04	7,8441E-04	4,8900E-04
F2	1,6540E-03	1,3840E-03	1,4891E-03	1,3727E-03	1,1425E-03
F3	1,6690E-03	1,3773E-03	1,5506E-03	1,3816E-03	1,1910E-03
O1	3,7540E-03	4,2872E-03	4,2936E-03	4,3469E-03	4,9554E-03
O2	3,4150E-03	3,9365E-03	3,9414E-03	4,0042E-03	4,6667E-03
O3	2,4820E-03	2,6719E-03	2,6748E-03	2,7740E-03	2,8463E-03
R1	2,5210E-03	2,8329E-03	2,7476E-03	2,9390E-03	3,1037E-03
R2	3,9310E-03	5,4824E-03	4,6440E-03	5,4101E-03	6,7249E-03
R3	2,6170E-03	2,9942E-03	2,8657E-03	3,0901E-03	3,2619E-03
R4	2,3620E-03	2,4940E-03	2,4680E-03	2,6054E-03	2,6298E-03
T1	1,0120E-03	7,0803E-04	7,0497E-04	5,1554E-04	2,1552E-04
T2	1,3190E-03	9,2689E-04	1,0913E-03	8,0119E-04	5,0125E-04
Generic	2,2590E-03	2,7899E-03	2,3386E-03	2,5449E-03	1,9836E-03

Appendix B: Indata benchmark

Indata till de tester där Vaurios PREB jämförs med Pörns metod och ZEDB. Dessa tester analyseras i avsnitt 10.2.1.

Centrifugalpump (T5, 1.1.1)		
Anlägg.	Antal fel	Drifttid
B1	2	128785
B2	0	109576
F1	6	139667
F2	5	141096
F3	0	100072
O1	1	198164
O2	4	188252
O3	0	101999
R1	10	175329
R2	4	305281
R3	1	107556
R4	0	100220
T1	3	161306
T2	2	162682
S:a	38	2119985

Skalventil (T5, 3.20.1)		
Anlägg.	Antal fel	Drifttid
B1	0	805446
B2	1	742632
F1	1	2297120
F2	3	2248330
F3	0	4123980
O1	0	322221
O2	0	1610400
O3	0	2102862
R1	1	2080890
R2	2	131380
R3	0	212268
R4	0	249616
T1	0	2016240
T2	2	1989060
S:a	10	20932445

Dieselaggregat (T5, 7.1.1)		
Anlägg.	Antal fel	Drifttid
B1	4	1242
B2	3	1238
F1	2	4090
F2	4	3501
F3	5	4198
O1	10	2018
O2	7	1500
O3	6	2108
R1	5	1611
R2	10	1487
R3	7	2146
R4	4	1521
T1	1	4640
T2	2	3990
S:a	70	35290

Appendix C: T-bokstabell

Tabellen är hämtad från T-boken version 6 (T6), s. 116.

Skruvpump, diesel dagtank					Tabell 1.11.1
Flöde:	<75 kg/s				Antal komponenter: 38
Tryckuppsättning:	0.2 - 2 MPa				Antal fel: 2
Driftläge:	Intermittent				Drifttid: 0.179780E6
Felmod:	Obefogat stopp				
Felintensitet:	$(\lambda_d \cdot 10^{-6}/h)$				Effektiv medelreptid
Anläggning	5%	50%	95%	medelv.	(h)
Barsebäck 1	0.0	0.1	1615.6	288.3	-
Barsebäck 2	0.0	0.0	1412.6	255.5	-
Forsmark 1	0.0	0.0	1215.7	231.7	-
Forsmark 2	0.0	4.1	8438.2	1673.1	10
Forsmark 3	0.0	0.0	726.8	133.8	-
Oskarshamn 1	0.0	0.0	1477.5	265.4	-
Oskarshamn 2					
Oskarshamn 3	0.0	0.0	726.8	133.8	-
Ringhals 1	0.0	0.0	1365.4	247.1	-
Ringhals 2	0.0	0.0	1143.0	208.7	-
Ringhals 3	0.0	0.0	8.3	1.5	-
Ringhals 4	0.0	0.0	8.3	1.5	-
Olkiluoto 1	0.0	2.2	4122.3	706.5	3
Olkiluoto 2	0.0	0.0	765.1	140.2	-
Generisk	0.0	0.0	934.1	305.7	7

Appendix D: MATLAB-kod

Koden är baserad på Vaurios ”algorithm” för PREB (appendix E). I den undre MATLAB-funktionen (*PREB_gamma_param*) härleds parametervektorn $\theta = (\alpha, \beta)$ för den generiska fördelningen. Denna uppdateras sedan med komponent- eller anläggningsspecifika data $\{k_i, t_i\}$ beroende på grupperingsprincip. Funktionen anropas från den övre funktionen (*PREB_lambda*) där fördelningarnas percentiler och medelvärden beräknas och returneras till promptern.

Den övre funktionen anropas från promptern med argumentet (k_i, t_i) dvs. två vektorer innehållande antal fel respektive drifttid. Indata ges alltså på formen

`PREB_lambda([k1 k2 ... kn], [t1 t2 ... tn]),`

där k_i och t_i väljs som komponentspecifika eller anläggningsspecifika data av användaren (gruppering görs alltså före exekvering).

```
function distdata = PREB_lambda(k,t)

theta = PREB_gamma_param(k,t); % Ask function PREB_gamma_param() for gamma parameters.

a = theta(1,:);
b = 1./theta(2,:);

% Compute distribution data.
quant_5 = gaminv(0.05,a,b);
quant_50 = gaminv(0.5,a,b);
quant_95 = gaminv(0.95,a,b);
mean = a.*b; % Corresponds to a./beta, where beta is theta(2,:).

distdata = [quant_5; quant_50; quant_95; mean]';
```

```
function theta = PREB_gamma_param(k,t)

delta = 0.5; % Compromise "initial alpha".

n = length(k); % Number of plants.

if n == 1 % Data available only from one plant.
    x = delta;
    y = 0;
else % Generic data available.
    T = sum(t);
    T_star = T - max(t);
    w = (t.*n + T)/(2*n*T); % Normalized initial values (recommended).
    % Simpler alternatives: w(i) = 1/n, w(i) = t./T

    m = sum(w.*k./t);

    if m == 0
        x_0 = 0;
        y_0 = T_star;
    else
        v = 1/(1 - sum(w.^2))*sum(w.*(k./t - m).^2) + m/T_star;

        r = 0; % Check t(i) for equality.
        for i = 1:n
            if t(1) ~= t(i)
                r = 1;
            end
        end
    end
end
```



```

        end
    end

    if r == 1          % All t(i) not equal.

        m_0 = 0;      % Reference variables for convergence.
        v_0 = 0;

        while ((abs(m - m_0) > eps) || (abs(v - v_0) > eps)) % Iterate until m and v
                                                                converge
            m_0 = m;
            v_0 = v;
            u = t./(t + m/v);
            w = u./sum(u);
            m = sum(w.*k./t);
            v = 1/(1 - sum(w.^2))*sum(w.*(k./t - m).^2) + m/T_star;
        end
    end

    y_0 = m/v;
    x_0 = m*y_0;
end

% Compute the prior parameters
x = x_0 + delta*y_0/T_star;
y = y_0;
end

% Compute the posterior parameters
alpha = k + x;
beta = t + y;

% Parameter vector theta
theta = [alpha ; beta];

```

Appendix E: Vaurios "algorithm" (PREB)

Nedan återges Vaurios metod på "algoritmisk" form enligt Vaurio & Jänkälä (2006, s. 212).

Anmärkning: Ett "corrigendum" avseende tryckfel i nämnda publikation har erhållits från Jussi Vaurio: I steg 4 ska $(K_i/T_i - m)$ ändras till $(K_i/T_i - m)^2$. I steg 11 ska $\delta/y_0 T^*$ ändras till $\delta/(y_0 T^*)$.

PREB

Följande procedur ger skattningarna x_c , y_c , M_c , och V_c för x , y , M respektive V :

1. Om endast data från den aktuella anläggningen finns tillgängliga ($n = 1$), välj $y_c = 0$ och $x_c = \delta$, där $0 \leq \delta \leq 1$; rekommenderat värde är $\delta = 1/2$. Vid feltrunkering välj $\delta = 0$. Gå till steg 12.
Om $n > 1$, sätt

$$T = \sum_{i=1}^n T_i$$
 och välj alla initiala vikter $w_i = 1/n$, eller $w_i = T_i/T$, $i = 1, 2, \dots, n$. Rekommenderade initiala värden är $w_i = (T + nT_i)/(2nT)$.
2. $T^* = T - \max(T_i)$.
3. $m = \sum_{i=1}^n w_i \frac{K_i}{T_i}$. Om $m = 0$, välj ett litet positivt ε och sätt $m = \varepsilon/T^*$, alternativt sätt $v = 0$, $y_0 = T^*$, $x_0 = 0$ och gå till steg 9.
4.
$$v = \frac{1}{1 - \sum_{i=1}^n w_i^2} \sum_{i=1}^n w_i \left(\frac{K_i}{T_i} - m \right)^2 + \frac{m}{T^*}.$$
5.
$$u_i = \frac{T_i}{(T_i + (m/v))}, i = 1, 2, \dots, n.$$
6.
$$w_i = \frac{u_i}{\sum_{j=1}^n u_j}, i = 1, 2, \dots, n.$$
7. Iterera steg 3 – 6 (om inte alla T_i är lika) tills m och v konvergerar.
8. $y_0 = m/v$, $x_0 = (m^2/v) = my_0$.
9. Välj $\delta = 0$ ("optimistic"), $\delta = 1/2$ ("compromise") eller $\delta = 1$ ("conservative"). Vid feltrunkering välj $\delta = 0$.
10. $x_c = x_0 + \delta(y_0/T^*)$, $y_c = y_c$.
11. A priorimoment: $M_c = m + (\delta/T^*)$ och $V_c = v + (\delta/(y_0 T^*))$.
12. A posteriorifördelningarna är nu $\Gamma(K_i + x_c, T_i + y_c)$.

Appendix F: Matematik

F.1 Fördelningar

Nomenklaturen utgår ifrån Råde & Westergren (1998) men är delvis anpassad till konventioner i övrig studerad litteratur.

Betafördelningen

$$\beta(p, q) = B^{-1} x^{p-1} (1-x)^{q-1}. \quad \{0 \leq x \leq 1\}$$

B är *betafunktionen* som definieras av

$$B(p, q) = \frac{\Gamma(p)\Gamma(q)}{\Gamma(p+q)},$$

där $\Gamma(\cdot)$ är *gammafunktionen* (se nedan).

Binomialfördelningen

Låt x vara antalet händelser vid n oberoende försök, där sannolikheten i varje försök är p . Sannolikheten $p(x)$ att observera x är då binomialfördelad dvs.

$$p(x) \in \text{Bin}(n, p) \Leftrightarrow$$

$$p(x) = \binom{n}{x} p^x (1-p)^{n-x}. \quad \{x = 0, 1, \dots, n\}$$

Exponentialfördelningen

$$p(x) \in \text{Exp}(\lambda) \Leftrightarrow$$

$$p(x) = \lambda e^{-\lambda x}. \quad \{x \geq 0\}$$

$$x < 0 \Rightarrow p(x) = 0.$$

Gammafördelningen

$$\Gamma(\alpha, \beta) = \frac{\beta^\alpha}{\Gamma(\alpha)} \lambda^{\alpha-1} e^{-\beta\lambda}.$$

Speciellt gäller att $\Gamma(1, \lambda) = \text{Exp}(\lambda)$, dvs. en exponentialfördelning med intensiteten λ .

$\Gamma(\cdot)$ är *gammafunktionen*:

$$\Gamma(p) = \int_0^{\infty} u^{p-1} e^{-u} dt = (p-1)!. \quad \{\text{Då } p \text{ är ett positivt heltal}\}$$

Lognormalfördelningen

$$p(\lambda) = \frac{1}{\sqrt{2\pi}\lambda\sigma} \cdot \exp\left[-\frac{\ln(\lambda/\mu)^2}{2\sigma^2}\right].$$

Poissonfördelningen

Låt x vara antalet oberoende händelser med exponentialfördelade tidsavstånd och felintensiteten λ . Sannolikheten att observera x under tiden T är då Poissonfördelad dvs.

$$p(x) \in Po(\lambda T) \Leftrightarrow$$

$$p(x) = \frac{(\lambda T)^x}{x!} e^{-\lambda T}. \quad \{x = 1, 2, \dots\}$$

F.2 Bevis

Härledning av likelihoodfunktionen (7.5) i Pörns hyperposteriorifördelning (7.4)

Påstående:

$$\left. \begin{array}{l} p(x|\lambda) \in Po(\lambda T) \\ p(\lambda|\theta) \in \Gamma(\alpha, \beta) \end{array} \right\} \Rightarrow p(x|\theta) \in NegBin\left(\alpha, \frac{\beta}{\beta+T}\right),$$

där $\theta = (\alpha, \beta)$

Bevis:

$$\begin{aligned} p(x|\theta) &= \int_0^{\infty} p(x|\lambda) p(\lambda|\theta) d\lambda = \int_0^{\infty} \frac{(\lambda T)^x}{x!} e^{-\lambda T} \cdot \frac{\beta^\alpha \lambda^{\alpha-1}}{\Gamma(\alpha)} e^{-\lambda \beta} d\lambda = \\ &= \frac{\beta^\alpha T^x}{x! \Gamma(\alpha)} \int_0^{\infty} \lambda^{\alpha+x-1} e^{-\lambda(\beta+T)} d\lambda = \left\{ \begin{array}{l} u = \lambda(\beta+T) \\ du = d\lambda(\beta+T) \end{array} \right\} = \\ &= \frac{\beta^\alpha T^x}{x! \Gamma(\alpha)} \cdot \frac{1}{(\beta+T)^{\alpha+x}} \int_0^{\infty} u^{\alpha+x-1} e^{-u} du = \left\{ \int_0^{\infty} u^{\alpha+x-1} e^{-u} du = \Gamma(\alpha+x) \right\} = \end{aligned}$$

$$= \frac{\Gamma(\alpha + x)}{\Gamma(\alpha)\Gamma(x+1)} \left(\frac{\beta}{\beta + T} \right)^\alpha \left(\frac{T}{\beta + T} \right)^x.$$

□

Uttrycket på sista raden kan också skrivas

$$\frac{(x + \alpha - 1)!}{x!(\alpha - 1)!} \left(\frac{\beta}{\beta + T} \right)^\alpha \left(\frac{T}{\beta + T} \right)^x,$$

eftersom $\Gamma(x+1) = x!$ (om x är ett positivt heltal). (Blom 1984, s. 71. Cooke et al. 1995 pt 2, s 12).