

**Weiterentwicklung
des Quantifizierungs-
verfahrens für GVA
zur Vermeidung von
Schätzfehlern aufgrund
vereinfachender
Modellannahmen**

Weiterentwicklung des Quantifizierungs- verfahrens für GVA zur Vermeidung von Schätzfehlern aufgrund vereinfachender Modellannahmen

Jan Stiller
Stefanie Blum
Albert Kreuser
Moritz Leberecht

Juni 2014

Anmerkung:

Dieser Bericht wurde im Rahmen des Vorhabens RS1198 mit Mitteln des Bundesministeriums für Wirtschaft und Energie (BMWi) erstellt.

Die Arbeiten wurden von der Gesellschaft für Anlagen- und Reaktorsicherheit (GRS) gGmbH durchgeführt.

Die Verantwortung für den Inhalt dieser Veröffentlichung liegt beim Autor. Die hierin geäußerten Meinungen müssen nicht mit der Meinung des Auftraggebers übereinstimmen.

Deskriptoren

GVA, Modell, Quantifizierung, Unsicherheiten, Konvergenz, Mapping

Kurzfassung

Die Sicherheit von Kernkraftwerken kann durch solche Ereignisse erheblich beeinflusst werden, bei denen aufgrund einer gemeinsamen Ursache Nicht-Verfügbarkeiten von mehreren Redundanten eines Systems auftreten. Solche Ereignisse werden als gemeinsam verursachte Ausfälle (GVA) bezeichnet. Probabilistische Sicherheitsanalysen (PSA) moderner Kraftwerke haben gezeigt, dass insbesondere bei höher redundanten Systemen Systemfunktionsausfälle aufgrund von GVA dominierend gegenüber Ausfällen von Systemfunktionen aufgrund mehrerer unabhängiger Ausfälle sein können, obwohl GVA-Ereignisse in der Betriebserfahrung relativ selten auftreten.

In dem vom Bundesministerium für Wirtschaft und Energie geförderten Forschungs- und Entwicklungsvorhaben RS1198 wurden Möglichkeiten zur Weiterentwicklung der Quantifizierung von GVA durch verbesserte GVA-Modelle entwickelt und diskutiert. Zuerst wurde der aktuelle Stand der GVA-Quantifizierung mit dem Kopplungsmodell unter Berücksichtigung aller Weiterentwicklungen, die seit der Veröffentlichungen der „Methoden zur probabilistischen Sicherheitsanalyse für Kernkraftwerke“ vorgenommen wurden, geschlossen dargestellt. Die Charakteristika des Modells wurden diskutiert. Diese beinhalten einerseits die umfassende Berücksichtigung von Schätzunsicherheiten und die Möglichkeit, Betriebserfahrung von Komponentengruppen abweichender Größe zu verwenden, als auch unerwünschte Konvergenzeigenschaften: Bei Anwachsen der Anzahl von GVA-Ereignissen nimmt die Unsicherheit der Schätzungen der GVA-Wahrscheinlichkeiten nicht in dem Maße ab, wie es der abnehmenden statistischen Unsicherheit entspricht. Diese Eigenschaft ist mit der grundlegenden Annahme des Kopplungsmodells verbunden, dass die Komponenten bei Auftritt eines GVA-Phänomens mit einer gewissen Wahrscheinlichkeit ausfallen (dem Kopplungsparameter), und dieser bei verschiedenen GVA-Phänomenen im Allgemeinen divergierend ist und deshalb für alle GVA-Ereignisse unabhängig geschätzt wird. Deshalb muss, um die unerwünschten Konvergenzeigenschaften zu vermeiden, auf diese zentrale Modellannahme verzichtet werden.

Um einen neuen Modellansatz zu gewinnen, wurden zunächst die international üblichen Vorgehensweisen zum Schätzen von GVA (u. a. das Alpha-Faktor- und das Beta-Faktor-Modell) beschrieben. Verfahren zur Schätzung der Modellparameter unter Verwendung statistischer Methoden von Bayes wurden in einer einheitlichen Form dargestellt. Diese Verfahren lassen sich nicht unmittelbar auf die deutsche Betriebserfahrung übertragen, da einerseits für die Schätzung benötigte Informationen nicht vorliegen,

andererseits bei der weiterentwickelten Modellierung die umfassende Einbeziehung der verschiedenen Unsicherheitsquellen, die die bisherige Vorgehensweise kennzeichnet, erhalten bleiben soll. Deshalb wurden Kriterien entwickelt, die der Entwicklung eines für die deutsche Betriebserfahrung geeigneten Modells zugrunde liegen sollten. Basierend auf diesen Kriterien wurden drei Modellansätze entwickelt. Im ersten Modellansatz (Modell A) werden GVA mit verschiedenen Ausfallkombinationen als unabhängige Basisereignisse angesehen. Die Raten dieser Ereignisse stellen die Modellparameter dar. In den beiden weiteren Modellansätzen (Modell B und C) werden generische Zustände „GVA“ bzw. „GVA-Phänomen“ postuliert, die mit einer Rate auftreten. Aus diesem Zustand geht das Modell in Endzustände über, die den verschiedenen Ausfallkombinationen entsprechen. Die entsprechenden bedingten Wahrscheinlichkeiten sind die weiteren Modellparameter. Diese Modellvorstellung ähnelt dem Alpha-Faktor-Modell. Ein wesentlicher Unterschied ist allerdings, dass nur Ausfälle mit systematischer Ursache und keine Einzelfehler beschrieben werden. Deshalb sind zum Schätzen der Modellparameter aus der Betriebserfahrung auch keine Einzelfehler erforderlich. Bayes'sche Schätzverfahren wurden für die drei Modelle hergeleitet. Bei ihnen werden die verschiedenen Unsicherheitsquellen in gleicher Qualität wie beim Kopplungsmodell berücksichtigt. Untersuchungen der Konvergenz der Modellparameter zeigen, dass Modelle B und C eine starke Unterschätzung der Wahrscheinlichkeiten von einzelnen GVA-Kombinationen (z. B. komplette GVA) zeigen können. Als Grund wurde die langsame Konvergenz der Modellparameter erkannt, die die Verteilung der Ereignisse auf die verschiedenen Ausfallkombinationen beschreiben, während die Konvergenz der Gesamtrate schneller ist. Demgegenüber treten bei Modell A nur Überschätzungen auf. Deshalb ist nur mit Modell A eine konservative Schätzung der GVA-Wahrscheinlichkeiten möglich. Die entwickelten GVA-Modelle sind komponentengruppengrößenspezifisch mit der Folge, dass zur Quantifizierung unmittelbar nur GVA-Ereignisse verwendet werden können, die in Komponentengruppen derselben Größe aufgetreten sind wie der GVA-Gruppe, für die GVA quantifiziert werden sollen. Da nicht für alle Gruppengrößen ausreichend Betriebserfahrung vorliegt, sind separate Algorithmen erforderlich, um die Ereignisse zwischen Komponentengruppen verschiedener Größe zu übertragen. Für dieses Mapping wurden verschiedene, teilweise neu entwickelte Ansätze, aufgeführt. Dabei ist besonders ein neuer Ansatz für das Mapping Up hervorzuheben, der nur auf der Annahme basiert, dass eine kleine Komponentengruppe als zufällige Untermenge der Komponenten einer großen angesehen werden kann, die nicht vollständig beobachtet wird. Mittels Bayes'scher statistischer Methoden können GVA-Wahrscheinlichkeiten in der großen Komponentengruppe berechnet werden. Die mathematischen Beziehungen lassen sich analytisch ausdrücken, d. h. man

ist nicht auf Monte-Carlo-Verfahren zur Implementation angewiesen. Für einen Spezialfall wurde die Konvergenz der geschätzten Parameter gegen ihre wahren Werte demonstriert.

Im Zusammenhang mit diesem Verfahren wurde auch die Kompatibilität der Annahme, dass eine kleine Komponentengruppe als Teil einer großen angesehen werden kann, mit den für die Schätzalgorithmen verwendeten a priori-Verteilungen untersucht mit dem Ergebnis, dass sie nicht kompatibel sind. Es wurden denkbare Lösungsmöglichkeiten diskutiert; jedoch ist nicht erkennbar, wie ein kompatibler a priori konstruiert werden könnte. Diese Inkompatibilität betrifft nicht nur das neu entwickelte Verfahren, sondern auch weitere Mappingverfahren, die auf dieser Grundannahme basieren (probabilistisch-kombinatorisches Mapping Down) und international häufig zusammen mit dem Alpha-Faktor-Modell angewandt werden. Bei dieser Vorgehensweise existiert ebenfalls die genannte Inkompatibilität, so dass eine solche Vorgehensweise eine innere Widersprüchlichkeit aufweist. Das Kopplungsmodell ist nicht von diesem Problem betroffen, da es keine komponentengruppengrößenspezifischen Parameter aufweist. Es wurde diskutiert, wie – abgesehen von der Forderung nach innerer Widerspruchsfreiheit – die verschiedenen Mappingalgorithmen bewertet werden können. Hierbei ist zu berücksichtigen, dass keine ausreichende empirische Evidenz vorhanden ist, wie sich GVA-Phänomene in Komponentengruppen unterschiedlicher Größe tatsächlich auswirken. Deshalb ist keine fundierte Bewertung im Hinblick darauf möglich, inwieweit sie ein realistisches Mapping gewährleisten. Darum wurden Kriterien für die Konservativität von Mappingverfahren entwickelt. Diese werden vollständig von den Verfahren „konservatives Mapping Down“ und „konservatives Mapping Up“ erfüllt, die sich als „weglassen“ der am schwächsten geschädigten Komponenten bzw. „duplizieren“ der am stärksten geschädigten Komponente charakterisieren lassen. Wegen der Konservativität wird die mit dem Mapping verbundene Unsicherheit in diesen Verfahren nicht explizit ausgewiesen.

Die entwickelten GVA-Modelle wurden anhand der deutschen Betriebserfahrung aus Kernkraftwerken erprobt. Dafür wurden zwei Datensätze zusammengestellt, die Populationen mit sehr wenigen beobachteten GVA-Ereignissen und Populationen mit vielen GVA-Ereignissen repräsentieren. In diesen Datensätzen sind nur Ereignisse an GVA-Gruppen einer Größe (Redundanzgrad 4) enthalten, um die Schätzverfahren unabhängig vom Mapping vergleichen zu können. Vergleiche der Schätzergebnisse der Modelle A und B mit dem Kopplungsmodell zeigen, dass die Ergebnisse sehr ähnlich sind. Die Mittelwerte der Ergebnisverteilungen liegen jeweils innerhalb der 95 %-Konfidenzinter-

valle aller anderen Verfahren. Dies gilt sowohl vor als auch nach der Einbeziehung der verbleibenden Unsicherheiten. Die Schätzungen mit dem Kopplungsmodell sind nicht signifikant verschieden von denjenigen mit den neuen Modellen.

Zusätzlich wurde das Verfahren zum konservativen Mapping in Verbindung mit Modell A angewandt, um die GVA-Wahrscheinlichkeiten in Komponentengruppen der Größe 3 konservativ zu schätzen. Die Ergebnisse wurden mit Schätzungen anhand der Betriebserfahrung in Komponentengruppen nur der Größe 3 verglichen, die nur ein Ereignis beinhaltet. Es zeigt sich, dass trotz der Konservativität die Schätzungen unter Verwendung des Mapping deutlich geringere Werte ergaben, da die Schätzungenauigkeit aufgrund der geringen Ereigniszahl bei Verwendung der Betriebserfahrung in Komponentengruppen nur der Größe 3 sehr hoch ist. Die Ergebnisse der konservativen Vorgehensweise sind auch vergleichbar zu den mit dem Kopplungsmodell erzielten Ergebnissen.

Modell A erlaubt es somit, konservative Schätzungen von GVA-Wahrscheinlichkeiten zu berechnen. In Verbindung mit dem konservativen Mapping kann auch Betriebserfahrung in Komponentengruppen abweichender Größe einbezogen werden.

Aus den erzielten Ergebnissen resultiert weiterer Forschungsbedarf. Dies betrifft einerseits die vertiefte Untersuchung von GVA-Entstehung und Entdeckung, um Erkenntnisse zu erlangen, die eine empirisch fundierte Bewertung des Mapping erlauben. Andererseits sollte die Kompatibilität von a priori-Verteilungen mit den dem Mapping zugrunde liegenden Annahmen mathematisch weiterführend betrachtet werden, um herauszufinden, in wieweit ein innerer Widerspruch von Mapping und der den Schätzverfahren zugrunde liegenden a priori-Verteilungen vermieden werden kann.

Abstract

The safety of nuclear power plants can be significantly affected by events with more than one redundant components being unavailable due to a common cause. Such events are called common cause failures (CCF). Although CCF are rare events, Probabilistic safety analyses (PSA) of modern nuclear power plants have shown that, particularly for systems with a high degree of redundancy, the unavailability due to CCF may be dominant in comparison to those due to independent failures.

In the research and development project RS1198 sponsored by the German federal Ministry for Economic Affairs and Energy (BMWi), possible ways of further developing the quantification of CCF by applying improved CCF models have been researched. Firstly a self-contained comprehensive description of the current procedures for CCF quantification with the coupling model has been developed. It includes all improvements introduced since the publication of the technical document on PSA methods supplementing the German PSA Guide. The characteristics of the coupling model have been discussed. On the one hand these include a comprehensive consideration of different uncertainties and the possibility of using the operating experience of component groups of differing sizes. On the other hand they include undesired convergence characteristics: With a growing number of CCF events, the estimation uncertainty does not appropriately reflect the decreasing statistical uncertainty. This property is linked to the central assumption of the coupling model that the components fail with a certain probability (the coupling parameter) when a CCF phenomenon occurs and that this parameter will generally be different for different CCF phenomena, which is why it is estimated independently for all CCF events. Hence this central model assumption needs to be dispensed with in order to avoid the undesired convergence characteristics.

In order to develop a new model approach, international methods for estimating CCFs have been described (e. g. the alpha factor and beta factor models). Methods for estimating the model parameters using Bayes' statistical methods have been shown in a consistent form. These methods cannot be directly applied to German operating experience as information that is needed for the estimation is not readily available and because the comprehensive consideration of the different sources of uncertainty that has been established with the current quantification method needs to be preserved. Therefore, criteria have been developed on which the development of a model that is suitable for German operating experience was based on. Based on these criteria, three model

approaches have been developed. In the first model approach (Model A), CCFs with different failure combinations are considered as independent basic events. The rates of these events are the model parameters. In the other two model approaches (Models B and C), generic states "CCF" or "CCF phenomenon", respectively, are postulated that occur at a specific rate. From this state, the model proceeds to final states that correspond to the different failure combinations. The corresponding conditional probabilities are the additional model parameters. This model structure is similar to the alpha factor model. One major difference, however, is that only failures with a systematic cause and no independent single failures are described. Hence no single failures are necessary for estimating the model parameters from operating experience. Bayes' estimation methods have been derived for the three models. In all three cases, the different sources of uncertainty are considered in the same quality as in the coupling model. Numerical studies of the convergence properties of the model parameters demonstrate that Models B and C can show a strong underestimation of the probabilities of individual CCF combinations (e. g. complete CCFs). The reason for this has been found to be the slow convergence of the model parameters that describe the distribution of the results over the different failure combinations, while the convergence of the overall rate is faster. In contrast, Model A can only show overestimations. Hence a conservative estimation of the CCF probabilities is only possible with Model A. The CCF models that have been developed are component-group-size-specific, with the consequence that for quantification, only those CCF events from operating can be directly used that occurred in groups of the same size as the CCF group for which the CCFs are to be quantified. As in many cases there is not sufficient operating experience available for all group sizes, separate so called mapping algorithms are necessary to apply the results to component groups of different sizes. For this mapping, various approaches – some of them newly developed ones – have been compiled. One of them that particularly noteworthy is a new approach for the mapping-up that is solely based on the assumption that a small component group can be considered as a random subset of the components of a larger group that is not fully observed. Applying Bayes' statistical methods, CCF probabilities in the large component group can be calculated. The mathematical relations can be expressed analytically, i. e. there is no need to rely on Monte-Carlo methods for implementing that method. For one special case, the convergence of the estimated parameters against their true values has been demonstrated.

In the course of the development of this approach, the compatibility of the assumption that a small component group can be regarded as part of a large group with the non-

informative a priori assumptions the Bayesian estimators are based on has been investigated, yielding the result that they are not indeed compatible. Possible solutions have been discussed; however, it could not be conceived how a compatible a priori could be constructed. This incompatibility concerns not only the newly developed method but also other mapping methods that are based on this fundamental assumption (probabilistic-combinatorial mapping-down) and which are frequently used together with the alpha factor model. When these two methods are combined the resulting procedure suffers from internal inconsistency. The coupling model is not affected by this problem as it has no component group size specific parameters.

It has been discussed how – apart from the demand for internal consistency – the different mapping algorithms can be assessed. There is no sufficient empirical evidence on what effect CCF phenomena in component groups of different sizes actually have. Hence it is not possible to conduct well-founded assessment regarding the realism of different mapping algorithms. This is why criteria for conservatism have been developed. These are fulfilled entirely by the "conservative mapping-down" and the "conservative mapping-up" methods, which can be characterised as "leaving out" the component with the least impairment and "duplicating" the most affected component. Owing to the conservatism, the uncertainty associated with the mapping is not explicitly represented in this method.

The methods developed have been tested using German operating experience. For this purpose, two data sets have been compiled that represent populations with very few observed CCF events and populations with many CCF events. These data sets only contain events of CCF groups of a single size (degree of redundancy 4) in order to be able to compare the estimation methods independently of the mapping. Comparisons of the estimation results of models A and B with the results obtained with the coupling model show that the results are very similar. The expected values of the distributions each lie within the symmetrical 95 %-confidence intervals of all other methods. This is true before as well as after the consideration of the remaining uncertainties. Thus the estimates made with the coupling model are not significantly dissimilar from those made with the new models. In so far the results do not imply a necessity to modify the current German CCF quantification methods for the present available German operating experience.

In addition, the method for conservative mapping has been applied together with Model A to provide conservative estimates CCF probabilities in component groups of size 3.

The results have been compared with estimates based on operating experience of component groups of size 3 only comprising only a single event. It turned out that despite the conservatism the estimates were significantly lower when mapping than when only using operating experience of component groups solely of size 3. This is due to the resulting very large statistical uncertainty due to the low number of events. The results of the conservative method are also similar to the ones obtained with the coupling model.

Model A hence allows the calculation of conservative estimates of CCF probabilities. In combination with conservative mapping, operating experience with component groups of differing sizes can also be taken into account.

The results achieved have revealed further need for research. This concerns on the one hand the in-depth study of CCF origin and detection in order to obtain knowledge that will allow an empirically well-founded assessment of different mapping approaches. On the other hand, additional research should be devoted to the question of the compatibility of a priori distributions and the assumptions on which the mapping is based in order to find out to what extent an incompatibility between the mapping and the a priori distributions can be avoided.

Inhaltsverzeichnis

1	Einführung	1
2	Bisherige Modellierung von gemeinsam verursachten Ausfällen.....	3
2.1	Beschreibung des Kopplungsmodells	3
2.1.1	Grundlagen des Modells	4
2.2	Gleichungen zur Berechnung von GVA-Wahrscheinlichkeiten	7
2.2.1	Berücksichtigung von Unsicherheiten	9
2.2.2	Berücksichtigung der verbleibenden Unsicherheitsquellen	18
2.3	Diskussion der Modelleigenschaften.....	20
3	Internationale Vorgehensweisen zur Modellierung gemeinsam verursachter Ausfälle	25
3.1	Notation	25
3.2	Basic-Parameter-Modell	27
3.3	Beta-Faktor-Modell	27
3.3.1	Parameter.....	28
3.3.2	Berechnung von GVA-Wahrscheinlichkeiten.....	28
3.4	Multiple-Greek-Letter-Modell	29
3.4.1	Parameter.....	29
3.4.2	Berechnung von GVA-Wahrscheinlichkeiten.....	31
3.5	Alpha-Faktor-Modell	31
3.5.1	Parameter.....	31
3.5.2	Berechnung von GVA-Wahrscheinlichkeiten.....	32
3.6	Binomial-Failure-Rate-Modell.....	32
3.6.1	Parameter.....	33
3.6.2	Berechnung von GVA-Wahrscheinlichkeiten.....	33
3.7	Übersicht über die verschiedenen Modelle und ihre Parameter	33

4	Schätzung der Modellparameter und GVA-Wahrscheinlichkeiten	35
4.1	Likelihood-Funktion.....	36
4.2	Direkte Schätzung der $q_{k \setminus r}$	37
4.3	Basic-Parameter-Modell	39
4.3.1	Schätzung unter Verwendung der Likelihood-Funktion	39
4.3.2	Vereinfachte Schätzung.....	40
4.4	Alpha-Faktor-Modell	42
4.4.1	Schätzung der Wahrscheinlichkeit von Ereignissen q_r^{BE}	43
4.4.2	Schätzung der Alpha-Faktoren	44
4.4.3	Erwartungswerte der GVA-Wahrscheinlichkeiten.....	45
4.5	Beta-Faktor-Modell	47
4.6	Multiple-Greek-Letter-Modell	47
5	Entwicklung von Modellen zur Quantifizierung von GVA aus der deutschen Betriebserfahrung	49
5.1	Randbedingungen der Modelle	49
5.2	Beschreibung der Modelle	51
5.2.1	Modell A	51
5.2.2	Modell B	53
5.2.3	Modell C	55
5.3	Schätzung der Modellparameter aus Ereignisanzahlen	57
5.3.1	Modell A	57
5.3.2	Modell B	59
5.3.3	Modell C	61
5.3.4	Konvergenzeigenschaften der Modelle	64
5.4	Schätzung der Modellparameter aus der Betriebserfahrung	73
5.4.1	Bestimmung der bedingten Wahrscheinlichkeit $p(\mathfrak{N} \mathfrak{E})$	74
5.4.2	Schätzung der Modellparameter aus der Betriebserfahrung	76
5.4.3	Implementation der Schätzalgorithmen mit Monte-Carlo-Verfahren	78
5.4.4	Berücksichtigung der verbleibenden Unsicherheitsquellen	80

5.5	Vergleich der Modelle	80
6	Mapping.....	83
6.1	Phänomene mit Ausfall aller Komponenten	83
6.1.1	Deutsche Betriebserfahrung	84
6.2	Phänomene ohne notwendige Ausfall aller Komponenten	86
6.2.1	Modellbasierter Ansatz	86
6.2.2	Probabilistisch-kombinatorischer Ansatz.....	87
6.2.3	Konservative Abwandlung des Probabilistisch-kombinatorischen Ansatzes.....	109
6.2.4	Heuristisches Mapping Up	111
6.2.5	Modellbasiertes Mapping Up	114
6.3	Kompatibilität des Mappings mit a priori-Annahmen	114
6.3.1	Diskussion der Lösungsmöglichkeiten und Konsequenzen.....	125
6.4	Vergleich der Mappingansätze.....	126
7	Anwendung auf die Betriebserfahrung	131
7.1	Beispieldatensätze.....	131
7.2	Vergleich der GVA-Modelle und Schätzalgorithmen mit dem Kopplungsmodell	133
7.3	Mapping.....	138
8	Zusammenfassung	143
	Literaturverzeichnis.....	147
	Abbildungsverzeichnis.....	149
	Tabellenverzeichnis.....	153
A	Anhang: Kombinatorische Formeln für das Mapping	155

1 Einführung

Die Sicherheit und Zuverlässigkeit von Kernkraftwerken können durch solche Ereignisse erheblich beeinflusst werden, bei denen aufgrund einer gemeinsamen Ursache Nicht-Verfügbarkeiten von mehreren Redundanten eines Systems auftreten. Solche Ereignisse werden als gemeinsam verursachte Ausfälle (GVA) bezeichnet.

Probabilistische Sicherheitsanalysen (PSA) für moderne Kraftwerke haben gezeigt, dass insbesondere bei höher redundanten Systemen Systemfunktionsausfälle aufgrund von GVA dominierend gegenüber Systemfunktionsausfällen aufgrund mehrerer unabhängiger Ausfälle sein können, obwohl GVA-Ereignisse in der Betriebserfahrung relativ selten auftreten. Deshalb kommt im Rahmen einer PSA der sachgerechten Quantifizierung von GVA-Ausfallwahrscheinlichkeiten von redundanten, sicherheitstechnisch wichtigen Systemen eine hohe Bedeutung zu.

Im Fachband zu PSA-Methoden /FAK 05/, der den deutschen PSA-Leitfaden /BMU 05/ ergänzt, werden verschiedene anerkannte Modelle zur Quantifizierung von GVA dargestellt. Diese lassen sich in Modelle mit direkten Parameterschätzungen (Modelle mit komponentenbasierten bzw. systembasierten Parameterschätzungen) und Modelle mit postulierten modellbezogenen Parametern (Schock-Modelle) unterteilen. Das bekannteste Modell mit modellbezogenen Parametern ist das BFR-Modell (Binomial Failure Rate Model). Dieses Modell soll gemäß Methodenband zum PSA-Leitfaden nur in einer weiterentwickelten Form angewendet werden. Als eine solche Weiterentwicklung hat die GRS das Kopplungsmodell /KRE 01/, /KRE 06/ entwickelt. Es erfüllt die im Leitfaden enthaltenen Anforderungen, die die Modellparameter aus der Betriebserfahrung bestimmen und dass die Unsicherheiten berücksichtigt werden. Grundannahme des Modells ist, dass bei Auftreten eines GVA-Phänomens die einzelnen Komponenten unabhängig mit einer Phänomen-spezifischen Wahrscheinlichkeit, dem sogenannten Kopplungsparameter, unverfügbar werden. Diese dem Modell zugrunde liegende vereinfachende Modellannahme war in der Vergangenheit erforderlich, da nur wenig Betriebserfahrung zur Schätzung von Nicht-Verfügbarkeiten aufgrund von GVA vorlag. Sie führt aber im Fall von einer großen Anzahl von GVA-Ereignissen zu unerwünschten Konvergenzeigenschaften. In den letzten Jahren hat die bezüglich gemeinsam verursachter Ausfälle ausgewertete Betriebserfahrung erheblich zugenommen. Um diese wachsende Betriebserfahrung besser nutzen zu können und zu genaueren Schätzungen von Nicht-Verfügbarkeiten aufgrund von GVA zu kommen, wurden verschiedene Ansätze für Schätzverfahren und für Verfahren zur Übertragung von GVA-Ereignissen

zwischen Komponentengruppen verschiedener Größe (so genanntes 'Mapping') entwickelt und erprobt. Mapping ist in sehr vielen Fällen erforderlich, da Komponentengruppen verschiedener Größe in einer Population zusammengefasst werden müssen, weil typischerweise nicht für jede Komponentengruppengröße ausreichend Betriebserfahrung vorliegt.

Zunächst wird die aktuelle Vorgehensweise zur Schätzung von GVA-Unverfügbarkeiten mit dem Kopplungsmodell unter Berücksichtigung der in den letzten Jahren vorgenommenen Verbesserungen und Weiterentwicklungen geschlossen dargestellt (Abschnitt 2). Dann werden internationale Vorgehensweisen zur Schätzung von GVA-Wahrscheinlichkeiten diskutiert (Abschnitt 3). In Abschnitt 4 wird die grundsätzliche Vorgehensweise zur Schätzung von Modellparametern und GVA-Wahrscheinlichkeiten aus der Betriebserfahrung mit Bayes'schen statistischen Methoden diskutiert und auf die in Abschnitt 3 beschriebenen Modelle angewandt. In Abschnitt 5 wird die Entwicklung von Modellen und Methoden zur Quantifizierung von GVA aus der deutschen Betriebserfahrung dargestellt. Zunächst werden die speziellen Randbedingungen diskutiert (Abschnitt 5.1). Diese schließen eine einfache Übertragung internationaler Vorgehensweisen aus. Dann werden drei Modellansätze für GVA entwickelt und die Schätzung der Modellparameter aus der Betriebserfahrung unter Berücksichtigung der verschiedenen Unsicherheitsquellen dargestellt (Abschnitte 5.3 und 5.4). Hierbei werden insbesondere die Konvergenzeigenschaften untersucht (Abschnitt 5.3.4). Anschließend werden in Abschnitt 6 Ansätze für die Übertragung von Ereignissen zwischen Komponentengruppen verschiedener Größe (Mapping) diskutiert. Hierbei wird insbesondere ein neues Verfahren zur Übertragung von Ereignissen auf größere Komponentengruppen entwickelt, das nur auf der Annahme basiert, dass sich eine Komponentengruppe als zufällige Untermenge der Komponenten einer größeren Gruppe auffassen lässt. Die Kompatibilität dieser Annahme mit in den Schätzverfahren verwendeten a priori-Verteilungen wird untersucht (Abschnitt 6.3). Die verschiedenen Ansätze zum Mapping werden vergleichend diskutiert. Dazu werden Kriterien für Konservativität entwickelt und angewandt (Abschnitt 6.4). In Abschnitt 7 werden die entwickelten Vorgehensweisen auf Beispieldatensätze aus der deutschen Betriebserfahrung angewandt. Die Ergebnisse werden in Abschnitt 8 zusammengefasst.

2 Bisherige Modellierung von gemeinsam verursachten Ausfällen

In diesem Abschnitt wird die bisherige Vorgehensweise der GRS zur Schätzung von GVA-Wahrscheinlichkeiten aus Ereignissen der Betriebserfahrung mit dem Kopplungsmodell dargestellt. Die Darstellung stellt eine Weiterentwicklung der Beschreibung des Kopplungsmodells im Methodenband zur probabilistischen Sicherheitsanalyse für Kernkraftwerke /FAK 05/ an. Die jüngsten Weiterentwicklungen des Modells zur konsistenten Berücksichtigung aller Unsicherheitsquellen /STI 08/, /STI 09/ sind in die Darstellung einbezogen worden, wobei die Unterschiede zur Vorgehensweise in /FAK 05/ aufgeführt wurden.

2.1 Beschreibung des Kopplungsmodells

Nachfolgend wird das von der GRS entwickelte mathematische Modell, das sogenannte Kopplungsmodell, beschrieben. Das Kopplungsmodell wurde entwickelt, um auch bei wenig vorliegender Betriebserfahrung (z. B. nur bei einem einzelnen beobachteten Ereignis) zu Schätzungen von Nicht-Verfügbarkeiten durch GVA zu kommen und die verschiedenen Quellen von Schätzunsicherheit umfassend zu berücksichtigen. Deshalb hat es folgende wesentliche Eigenschaften:

- Durch die getroffenen Modellannahmen kann eine Schätzung der GVA-Wahrscheinlichkeiten der verschiedenen GVA-Ausfallkombinationen¹ auch dann erfolgen, wenn
 - nur GVA-Ereignisse in Komponentengruppen aufgetreten sind, die eine andere Größe haben als die Zielkomponentengruppe² und/oder
 - bestimmte Ausfallkombinationen nicht beobachtet wurden.

¹ Als Ausfallkombination wird die Anzahl ausgefallener Komponenten bezogen auf die Gesamtzahl der Komponenten einer GVA-Komponentengruppe bezeichnet. Z. B. wird ein Ausfall von 3 von 4 redundanten Komponenten als (3 von 4)-Ausfall bezeichnet. Mathematische Größen, die sich auf einen (3 von 4)-Ausfall beziehen, werden mit dem Index $3 \setminus 4$ gekennzeichnet.

² Als Zielkomponentengruppe wird die der PSA modellierte Komponentengruppe bezeichnet, für die GVA-Wahrscheinlichkeiten geschätzt werden sollen.

- Das Modell berücksichtigt umfassend quantitativ die verschiedenen Schätzunsicherheiten. Dies umfasst:
 - statistische Unsicherheiten, die sich aus dem beschränkten Umfang der Betriebserfahrung ergeben,
 - Unsicherheiten von Expertenbewertungen der Komponentenschädigungen bei Ereignissen und der Übertragbarkeit von GVA-Phänomenen,
 - eine mögliche Inhomogenität von beobachteten Populationen, d. h. ein nicht vollständig gleiches Ausfallverhalten aller in der Population enthaltenen Komponentengruppen über die gesamte Beobachtungszeit.

2.1.1 Grundlagen des Modells

Die Betriebserfahrung hat gezeigt, dass eine Komponentengruppe verschiedenen GVA-Phänomenen, wie z. B. GVA aufgrund von Korrosion, Fertigungs- oder Auslegungsfehlern, ausgesetzt sein kann, deren Auswirkungen auf die Komponentengruppe unterschiedlich stark sein können. Das Auftreten eines GVA-Phänomens stellt ein „Schockereignis“ dar, dessen Auswirkungen auf das Ausfallverhalten der Komponenten durch einen Kopplungsparameter³ η beschrieben werden. Es wird angenommen, dass die Komponenten unabhängig voneinander und mit der gleichen Wahrscheinlichkeit η ausfallen, wenn ein bestimmtes Schockereignis eingetreten ist, während sie mit der Wahrscheinlichkeit $(1 - \eta)$ verfügbar sind. Der Kopplungsparameter η ist somit die bedingte Wahrscheinlichkeit dafür, dass die Komponente ausfällt, wenn ein bestimmtes Schockereignis aufgetreten ist. Schockeeinwirkungen auf eine Komponentengruppe können wegen der Verschiedenheit der GVA-Phänomene von Schockereignis zu Schockereignis mit deutlich unterschiedlichen Ausfallwahrscheinlichkeiten der Komponenten verbunden sein. Wie die Betriebserfahrung zeigt, hätte die Schätzung eines einheitlichen Kopplungsparameters η für alle GVA-Ereignisse zur Folge, dass GVA-Wahrscheinlichkeiten für hohe Ausfallkombinationen systematisch deutlich unterschätzt würden. Aus diesem Grund wird im vorliegenden Modell für jedes beobachtete GVA-Ereignis ein separater Kopplungsparameter η bestimmt. Dieser beschreibt die

³ In alten Veröffentlichungen wurde der Kopplungsparameter meist mit „p“ bezeichnet. Um Verwechslungen mit Wahrscheinlichkeiten bzw. Wahrscheinlichkeitsdichten zu vermeiden, wird jetzt „ η “ verwendet.

bedingte Ausfallwahrscheinlichkeit jeder Komponente der Gruppe unter der Bedingung, dass der GVA-Mechanismus des beobachteten Ereignisses auf die Komponentengruppe einwirkt.

Populationsbildung

Das Kopplungsmodell ist ein sogenanntes absolutes GVA-Modell. Bei diesen Modellen werden die GVA-Wahrscheinlichkeiten direkt geschätzt und nicht auf die Wahrscheinlichkeiten unabhängiger Ausfälle bezogen. Deshalb muss die betrachtete Population aus definierten Untersuchungseinheiten bestehen, deren individuelle Zusammensetzung über den gesamten Beobachtungszeitraum konstant bleibt. Bei der GVA-Bewertung werden die Untersuchungseinheiten einer Population durch Komponentengruppen dargestellt. Zu einer Population werden Komponentengruppen aus Komponenten gleicher Komponententart (z. B. Kreislumpen oder Absperrschieber) zusammengefasst. In den meisten Fällen besteht eine Komponentengruppe aus den redundanten Komponenten eines mehrsträngigen Systems. Bei einer teilweise diversitären Komponentengruppe muss festgelegt werden, ob die nicht-diversitären Teil-Komponentengruppen einzeln bewertet werden. Dann sind GVA-Phänomene, die in mehreren Teil-Komponentengruppen aufgetreten sind, entsprechend mehrfach zu zählen. Unabhängig davon kann eine Kopplung zweier zueinander teilweise diversitärer Komponentengruppen im Fehlerbaum modelliert werden.

Eingangsgrößen

Die Berechnung der GVA-Wahrscheinlichkeiten auf der Basis eines beobachteten GVA-Ereignisses j verwendet folgende Informationen:

- Beobachtungszeit T ,
- Wert des Kopplungsparameters η_j ,
- Fehlerentdeckungszeit t_j für die Zielkomponentengruppe,
- Übertragbarkeitsfaktor f_j .

Diese Größen werden im Folgenden diskutiert.

Beobachtungszeit

Zur Bestimmung der Beobachtungszeit T einer Population werden die Beobachtungszeiten aller Komponentengruppen, die die Population bilden, addiert. Liegen die dafür erforderlichen detaillierten Informationen über die einzelnen Komponentengruppen einer Population nicht vor, kann die Beobachtungszeit als Produkt der mittleren Anzahl der beobachteten Komponentengruppen pro Anlage und der Gesamtbeobachtungszeit aller betrachteten Anlagen abgeschätzt werden.

Kopplungsparameter

Die Wahrscheinlichkeitsverteilung des Kopplungsparameters η_j wird aus den von Experten bestimmten Schädigungen der Komponenten der Komponentengruppe, die vom betrachteten GVA-Ereignis betroffen wurde, ermittelt (siehe Abschnitt 2.2).

Fehlerentdeckungszeit

Die Fehlerentdeckungszeit t_j wird durch die Instandhaltungsstrategie für die zu analysierende Komponentengruppe oder durch deren Betriebsweise bestimmt.

Bei Komponenten, die während des Betriebs im Stand-by-Zustand sind, werden im Allgemeinen die Zeitintervalle der wiederkehrenden Prüfungen (WKP) oder die jährlichen Funktionsprüfungen vor Anfahren der Anlage berücksichtigt. Hierbei ist zu beachten, dass nur die Prüfungen zu berücksichtigen sind, die nach Prüfumfang und Prüfmart geeignet sind, das entsprechende GVA-Phänomen zu entdecken. Daher ist die Fehlerentdeckungszeit einzeln für jedes Ereignis zu ermitteln.

Bei Komponenten, die während des Anlagenbetriebs zeitweise in Betrieb sind, werden als Fehlerentdeckungszeit die zwischen den Anforderungen liegenden Zeitintervalle der Komponenten gewählt. Wenn die Zeitintervalle stark schwanken, wird der Erwartungswert der Verteilung der Zeitintervalle als Fehlerentdeckungszeit verwendet.

Bei bestimmten Komponenten, die während des Anlagenbetriebs ständig in Betrieb sind, können bei unveränderten Anforderungen an diese Komponenten (z. B. Füllstandsmessungen oder Regelventile bei konstantem Leistungsbetrieb) mögliche GVA-Phänomene nicht erkannt werden. Als Fehlerentdeckungszeit wird in diesem Fall der zeitliche Abstand zwischen den jährlichen Funktionsprüfungen beim Anfahren der Anlage genommen.

Bei versetzter Testweise ist zu berücksichtigen, dass die Erkennung eines Fehlers der ersten Komponente in vielen Fällen nicht zur Erkennung eines GVA führt. Je nach Anzahl der vom GVA betroffenen Komponenten und der auf die Entdeckung des ersten Fehlers folgenden Instandhaltungsmaßnahmen können unterschiedlich lange Zeitintervalle bis zur Erkennung des GVA vergehen. Daher wird die nach der Betriebserfahrung typische Fehlerentdeckungszeit, das Doppelte des zeitlichen Abstandes aufeinander folgender Tests innerhalb der Komponentengruppe, angenommen. Beispielsweise wird im Falle einer vierwöchentlichen Prüfung von vier Komponenten mit versetzter Testweise die Fehlerentdeckungszeit gleich zwei Wochen gesetzt. Bei nicht versetztem Testen ist die Fehlerentdeckungszeit gleich dem Zeitintervall zwischen aufeinander folgenden Tests.

Übertragbarkeitsfaktor

Mit dem Übertragbarkeitsfaktor f_j hat der Experte die Möglichkeit zu bewerten, ob das dem Ereignis j zugrunde liegende GVA-Phänomen in der Zielkomponentengruppe mit kleinerer, gleicher oder größerer Wahrscheinlichkeit als in den übrigen betrachteten Komponentengruppen auftreten kann. Im Allgemeinen ist der Übertragbarkeitsfaktor gleich eins, da im Beobachtungsumfang nur vergleichbare Komponentengruppen zusammengefasst worden sind. Hiervon kann abgewichen werden, wenn grundlegende technische oder administrative Randbedingungen in der zu analysierenden Anlage vorliegen, die eine andere Wahrscheinlichkeit des Auftretens des beobachteten GVA-Phänomens in der zu analysierenden Anlage erwarten lassen.

2.2 Gleichungen zur Berechnung von GVA-Wahrscheinlichkeiten

Für jedes beobachtete GVA-Ereignis j wird für jede zu bewertende Ausfallkombination (k von r), wobei die Zielkomponentengruppe r Komponenten aufweist und $k \in \{0, 1 \dots r\}$ ist, ein anteiliger Beitrag $q_{k \setminus r; j}$ an der GVA-Wahrscheinlichkeit $q_{k \setminus r}$ der Zielkomponentengruppe berechnet als:

$$q_{k \setminus r; j} = \varphi_j p(k \setminus r | \eta_j) \quad (2.1)$$

Hierbei bezeichnet φ_j die Wahrscheinlichkeit, dass ein GVA durch Phänomen j in der Zielkomponentengruppe auftritt. Gleichung (2.1) hat die Form eines Produktes der Wahrscheinlichkeit, dass ein GVA durch Phänomen j in der Zielkomponentengruppe

auftritt, mit der bedingten Wahrscheinlichkeit $p(k|r|\eta_j)$, dass k von r Komponenten ausfallen, gegeben dass das GVA-Phänomen j aufgetreten ist. Unter den oben genannten Bedingungen genügt die Anzahl ausgefallener Komponenten einer Binomialverteilung mit dem Parameter η_j , wenn in der Zielkomponentengruppe das GVA-Phänomen j aufgetreten ist:

$$p(k|r|\eta_j) = \binom{r}{k} \eta_j^k (1 - \eta_j)^{r-k} \quad (2.2)$$

Die Wahrscheinlichkeit φ_j , dass das GVA-Phänomen j in der Zielkomponentengruppe auftritt, wird berechnet als

$$\varphi_j = f_j t_j \lambda_j \quad (2.3)$$

Hierbei ist, wie oben dargestellt, f_j der Übertragbarkeitsfaktor und t_j die Fehlerentdeckungszeit. λ_j bezeichnet die Rate des GVA-Phänomens j in der beobachteten Population. Gleichung 2.3 ist nur gültig für den Fall, dass λ_j klein ist, d. h. dass

$$f_j t_j \lambda_j \ll 1 \quad (2.4)$$

gilt. Dies ist in praktischen Anwendungen der Fall.

Die GVA-Wahrscheinlichkeit für die Ausfallkombination (k von r) der Zielkomponentengruppe $q_{k\backslash r}$ ist die Summe der anteiligen GVA-Wahrscheinlichkeiten $q_{k\backslash r;j}$ über alle relevanten GVA-Ereignisse:

$$q_{k\backslash r} = \sum_{j=1}^N q_{k\backslash r;j} \quad (2.5)$$

Dabei gibt N die Zahl der relevanten beobachteten GVA-Ereignisse in der betrachteten Population von Komponentengruppen an.

2.2.1 Berücksichtigung von Unsicherheiten

Im Folgenden wird die Berücksichtigung von Unsicherheiten dargestellt.

Statistische Unsicherheiten

Die Schätzung des Kopplungsparameters η_j wird mittels Bayes'scher statistischer Verfahren durchgeführt. Dabei wird angenommen, dass man sich bei der Schätzung von η_j nur auf die vorhandene Beobachtung stützen kann und keine zusätzlichen Vorinformationen zur Verfügung stehen. Zur Bestimmung der nichtinformativen a priori-Verteilung wird das Verfahren von Jeffreys /BOX 73/ angewandt.

$$\pi(\eta_j) \propto \frac{1}{\sqrt{\eta_j(1 - \eta_j)}} \quad (2.6)$$

Hierbei bezeichnet $\propto$ die Proportionalität. Mit dieser nichtinformativen a priori-Verteilung und einem beobachteten GVA-Ereignis mit $(k \text{ von } m)$ -Ausfällen erhält man über den Satz von Bayes eine Betaverteilung mit den Parametern $k + 1/2$ und $m - k + 1/2$ als a posteriori-Verteilung des Kopplungsparameters η_j . Für die Dichte dieser Verteilung gilt:

$$p(\eta_j) = \frac{\Gamma(m + 1)}{\Gamma(k + 1/2) \Gamma(m - k + 1/2)} \eta_j^{k-1/2} (1 - \eta_j)^{m-k-1/2} \quad (2.7)$$

Zur konsistenten Berücksichtigung der statistischen Unsicherheit der Rate des Auftretens von GVA-Ereignissen wird analog zur Berücksichtigung der Schätzunsicherheit des Kopplungsfaktors vorgegangen. Es wird mittels Bayes'scher Verfahren eine a posteriori-Verteilung der Rate bestimmt.

Dabei wird von der a priori-Verteilung der Rate λ_j , mit welcher Ereignisse mit dem GVA-Phänomen von Ereignis j auftreten, ausgegangen. Diese a priori-Verteilung wird analog zur a priori-Verteilung des Kopplungsparameters als nichtinformativ a priori über die Jeffreys'sche Regel /BOX 73/ hergeleitet. Dies entspricht auch der Vorgehensweise bei der Schätzung von Verteilungen für Ausfallraten unabhängiger Ausfälle (siehe Abschnitt 3.3 in /FAK 05a/). Die nichtinformativ a priori-Verteilung nach Jeffreys lautet:

$$\pi(\lambda_j) \propto \frac{1}{\sqrt{\lambda_j}} \quad (2.8)$$

Der a priori kann nur bis auf eine Konstante angegeben werden, da es sich um einen nicht normierbaren a priori, einen sogenannten ‘improper prior’ /BER 80/ handelt.

Wie oben dargestellt, wird davon ausgegangen, dass bei den verschiedenen GVA-Ereignissen verschiedene Phänomene wirksam geworden sind. Nach dem Satz von Bayes folgt für die a posteriori-Verteilung der Rate des Auftretens von GVA-Phänomen j, da in der Gesamtbeobachtungszeit aller Komponentengruppen des Beobachtungskollektivs T ein Ereignis des GVA-Phänomens j aufgetreten ist:

$$p(\lambda_j) = \frac{T^{3/2}}{\Gamma(3/2)} \lambda_j^{1/2} e^{-\lambda_j T} = \frac{2}{\sqrt{\pi}} \sqrt{T^3 \lambda_j} e^{-\lambda_j T} \quad (2.9)$$

Diese Verteilung ist für alle GVA-Phänomene $j = 1, 2 \dots N$ identisch und entspricht einer Gamma-Verteilung mit den Parametern 3/2 und 1/T.

Beim Entwicklungsstand des Kopplungsmodells, wie es im Methodenband zur probabilistischen Sicherheitsanalyse für Kernkraftwerke /FAK 05/ beschrieben ist, wurde die statistische Schätzunsicherheit nicht in der oben beschriebenen Form berücksichtigt, sondern eine Punktschätzung $\lambda_j = 1/T$ verwendet und die Schätzunsicherheit erst während der abschließenden „Verbreiterung“ (siehe Abschnitt 2.2.1.5 in /STI 09/) berücksichtigt.

Interpretationsunsicherheiten

Eine wesentliche Unsicherheitsquelle, die Einfluss auf die Schätzung des Kopplungsparameters η_j hat, ist die Einschätzung des Experten, wie das beobachtete GVA-Ereignis zu bewerten ist.

Bei der GVA-Bewertung von in der Betriebserfahrung aufgetretenen Ereignissen ist zu entscheiden, ob ein GVA vorliegt und wenn ja, wie viele Komponenten der betroffenen Gruppe durch das GVA-Phänomen ausgefallen sind oder aber geschädigt wurden, ohne dass es zu einem Ausfall kam. Dabei sind auch solche Ereignisse als GVA zu werten, bei denen zwar keine oder nur eine Komponente vollständig ausgefallen ist, bei

denen aber davon ausgegangen werden kann, dass die Komponentengruppe von einem GVA-Phänomen betroffen wurde, da an anderen Komponenten der betroffenen Gruppe ein für den GVA typisches (evtl. erst beginnendes) Schadensbild beobachtet wurde oder Fehler wie z. B. der Einsatz ungeeigneter Betriebsstoffe vorliegen, die ein entsprechendes GVA-Phänomen verursachen können.

In engem Zusammenhang damit steht die Entscheidung des Experten, wie das beobachtete Ereignis zu bewerten ist. Sind beispielsweise zu einem Testzeitpunkt durch ein GVA-Phänomen eine Komponente aus einer Gruppe von vier Komponenten als ausgefallen und zwei weitere Komponenten aus dieser Gruppe als geschädigt festgestellt worden, stellt sich für den Experten die Frage, wie diese Beobachtung zu bewerten ist. Dazu wird betrachtet, ob die Komponenten über eine durch die ermittelten Mindestanforderungen festgelegte Einsatzdauer ihre Funktion erfüllen würden, wenn sie in dem beobachteten Schadenszustand angefordert würden.

Meist ist eine sichere Bewertung, ob zu der bereits ausgefallenen Komponente entweder keine, eine oder sogar beide der geschädigten Komponenten bezüglich des zugrunde gelegten Anforderungsfalls zusätzlich als ausgefallen zu bewerten sind, nicht möglich. Die Beurteilung des GVA-Ereignisses durch den Experten ist folglich mit Unsicherheiten behaftet, da er das beobachtete GVA-Ereignis nicht mit Sicherheit eindeutig klassifizieren kann. Diese Art von Unsicherheit wird als Interpretationsunsicherheit bezeichnet.

Da diese Art von Unsicherheiten in der Praxis von GVA-Bewertungen häufig vorkommen und einen nicht unerheblichen Einfluss auf die Schätzung von GVA-Wahrscheinlichkeiten haben können, ergibt sich die Notwendigkeit, Interpretationsunsicherheiten in der Auswertung des Modells direkt zu berücksichtigen.

Das GVA-Modell bietet dem Experten die Möglichkeit, seine Interpretationsunsicherheiten bezüglich des beobachteten GVA-Ereignisses zu spezifizieren. Dies erfolgt dadurch, dass er alle in Frage kommenden Möglichkeiten der Beurteilung (Interpretationshypothesen bzw. Interpretationsalternativen) mit subjektiven Wahrscheinlichkei-

ten⁴ belegt, die seinen Grad an Sicherheit bezüglich des Zutreffens der jeweiligen Alternative ausdrückt.

vorliegt⁵. Die subjektiven Wahrscheinlichkeiten $w_{0\setminus r}, w_{1\setminus r}, \dots, w_{r\setminus r}$ beschreiben den jeweiligen Grad an Vertrauen, den der Experte zu den einzelnen Alternativen hat. Es ist zu beachten, dass folgende Bedingungen erfüllt sind:

$$\sum_{i=0}^r w_{i\setminus r} = 1 \text{ und } \forall_{i=0,1,\dots,r}: w_{i\setminus r} \in [0,1] \quad (2.10)$$

Im Fall $w_k = 1$ und $w_{i\setminus r} = 0$ für alle $i \neq k$ ist sich der Experte absolut sicher, dass ein (k von r)-Ausfall vorliegt. Mit $w_{k\setminus r} = 0$ wird ausgedrückt, dass mit Sicherheit kein (k von r)-Ausfall eingetreten ist.

Die obige Ausdrucksweise der Unsicherheiten durch Interpretationshypothesen mit den zugehörigen subjektiven Wahrscheinlichkeiten lässt sich übersichtlich durch einen so genannten Interpretationsvektor $W = (w_{0\setminus r}; w_{1\setminus r}; \dots; w_{r\setminus r})$ darstellen. Die Unsicherheit bezüglich der Ereignisinterpretation in der beobachteten Anlage wird somit durch verschiedene Alternativen mit dazu spezifizierten subjektiven Wahrscheinlichkeiten ausgedrückt.

Sind die Interpretationshypothesen mit den dazugehörigen subjektiven Wahrscheinlichkeiten festgelegt, wird durch eine Mischung von Betaverteilungen die Kenntnis über den Kopplungsparameter η_j beschrieben. Dazu wird zu jeder Interpretationshypothese die entsprechende Betaverteilung erzeugt, die dann mit der jeweiligen subjektiven Wahrscheinlichkeit als Gewicht in die Mischung eingeht.

Die aus den einzelnen Alternativen gewonnenen Verteilungen werden dann mit den subjektiven Wahrscheinlichkeiten $w_{k\setminus r}$, $k = 0, 1, \dots, r$ als Gewichte gemittelt. Als Ergebnis erhält man damit die Wahrscheinlichkeitsdichte $p(\eta_j)$ für den Kopplungsparameter η_j des GVA-Ereignisses j :

⁴ Eine ausführliche Darstellung Bayes'scher statistischer Verfahren und der dort verwendeten Begriffe ist in /BER 80/ zu finden.

⁵ Gegenüber vorigen Arbeiten wurde die Notation dahingehend erweitert, dass die Ausfallkombination $k\setminus r$ nun vollständig ausgeschrieben wird.

$$p(\eta_j) = \sum_{i=0}^r w_{i \setminus r} \frac{\Gamma(r+1)}{\Gamma(i+1/2) \Gamma(r-i+1/2)} \eta_j^{i-1/2} (1-\eta_j)^{r-i-1/2} \quad (2.11)$$

Es ist offensichtlich, dass nur diejenigen Alternativen in die Berechnung eingehen, deren subjektive Wahrscheinlichkeiten größer als 0 sind. Die so erhaltene Mischverteilung spiegelt den Kenntnisstand für den Kopplungsparameter η_j des beobachteten GVA-Ereignisses j unter Einbeziehung der statistischen Unsicherheit sowie der Interpretationsunsicherheiten wieder. Mit Hilfe der Gleichungen 2.1, 2.2, 2.3, 2.9 und 2.11 lassen sich die resultierenden Verteilungen der GVA-Unverfügbarkeiten $q_{k \setminus r}$ bestimmen.

Voreinstellung der subjektiven Wahrscheinlichkeiten zu den Interpretationsalternativen

Die Aufgabe, direkt subjektive Wahrscheinlichkeiten der verschiedenen (k von r)-Interpretationen anzugeben, stellt sich auch für Experten als sehr schwierig dar. Aus diesem Grund wurde ein Verfahren entwickelt, das anhand von leichter zu spezifizierenden Angaben durch den Experten eine automatische Voreinstellung der Interpretationshypothesen und der zugehörigen subjektiven Wahrscheinlichkeiten erlaubt. Diese Voreinstellung der subjektiven Wahrscheinlichkeiten kann in Abhängigkeit der Gegebenheiten und des Kenntnisstandes des Experten verändert werden.

Die notwendigen Angaben zur Erzeugung der Voreinstellung sind die Anzahl der ausgefallenen und die Anzahl der geschädigten Komponenten. Die geschädigten Komponenten können in verschiedene Schädigungskategorien, wie z. B. in stark geschädigte, schwach geschädigte und gering (sehr schwach) geschädigte Komponenten, eingeteilt werden. Jeder Schädigungskategorie wird ein entsprechender Schädigungswert zugeordnet. In Tab. 2.1 sind die Schädigungskategorien und zugehörigen Schädigungswerte der Voreinstellung aufgeführt.

Der Schädigungswert kann als der Grad an Vertrauen des Experten interpretiert werden, dass eine geschädigte Komponente bei der nächsten Anforderung ausfallen würde. Bei einer stark geschädigten Komponente beschreibt z. B. ein Schädigungswert von 0.5, dass die geschädigte Komponenten mit subjektiver Wahrscheinlichkeit 0.5 bei ihrer nächsten Anforderung ausfallen würde.

Tab. 2.1 Schädigungskategorien und Schädigungswerte

Schädigungskategorie	Schädigungswert
Ausfall	1
Starke Schädigung	0.5
Schwache Schädigung	0.1
Geringe (d. h. sehr schwache) Schädigung	0.01
Keine Schädigung	0

Für alle Komponenten $i = 1, 2, \dots, r$ einer von einem GVA-Ereignis betroffenen Komponentengruppe der Größe r wird das Ausmaß d_i der Schädigung der Komponente i bestimmt. Dazu gibt der Experte eine verbale Beschreibung des Schädigungszustandes der einzelnen Komponenten ab und klassifiziert diese in die Kategorien „ausgefallen“, „stark geschädigt“, „schwach geschädigt“, „gering (d. h. sehr schwach) geschädigt“ und „keine Schädigung“. Den Beschreibungen und darauf folgenden Klassifizierungen werden die oben angegebenen standardisierten Schädigungswerte zugeordnet. Dies entspricht dem Vorgehen der US-amerikanischen Aufsichtsbehörde NRC (Nuclear Regulatory Commission) bei der Bewertung von GVA-Ereignissen [WIE 07]. Das Verfahren ist in den USA und inzwischen auch in Deutschland erprobt und hat sich bewährt.

In bestimmten Fällen können sich Besonderheiten bei der Beurteilung von GVA-Ereignissen ergeben, so dass von den Standardwerten abweichende Schädigungswerte angegeben werden können.

Die Werte der Elemente des Interpretationsvektors werden auf der Basis der möglichen Fehlerkombinationen berechnet. Als Beispiel sind für eine Komponentengruppe mit drei Komponenten die Berechnungsvorschriften in Tab. 2.2 dargestellt.

Tab. 2.2 Berechnungsvorschriften für eine Komponentengruppe mit drei Komponenten

Ausfallkombination	Berechnung des zugehörigen Wertes der Komponente des Interpretationsvektors
0 von 3	$w_{0\setminus 3} = (1 - d_1)(1 - d_2)(1 - d_3)$
1 von 3	$w_{1\setminus 3} = d_1 (1 - d_2)(1 - d_3) + (1 - d_1)d_2(1 - d_3) + (1 - d_1)(1 - d_2)d_3$
2 von 3	$w_{2\setminus 3} = d_1 d_2(1 - d_3) + d_1(1 - d_2) d_3 + (1 - d_1)d_2 d_3$
3 von 3	$w_{3\setminus 3} = d_1 d_2 d_3$

Wurde zum Beispiel bei einem GVA-Ereignis in einer Komponentengruppe mit drei Komponenten beobachtet, dass eine Komponente ausgefallen, eine stark geschädigt und eine nicht geschädigt war, ergibt sich aus dem Schädigungsvektor $D = (d_1; d_2; d_3) = (1, 0.5, 0)$ und aus der obigen Berechnungsvorschrift der voreingestellte Interpretationsvektor $W = (w_{0\setminus 3}, w_{1\setminus 3}, w_{2\setminus 3}, w_{3\setminus 3}) = (0, 0.5, 0.5, 0)$. Das Ereignis wird dann aufgrund der Schädigungsbeurteilung der Komponenten so interpretiert, dass mit 50 % subjektiver Wahrscheinlichkeit ein (1 von 3)-Ausfall und mit 50 % subjektiver Wahrscheinlichkeit ein (2 von 3)-Ausfall vorliegt.

Unter Verwendung des ermittelten Interpretationsvektors W wird über die oben beschriebene Vorgehensweise (siehe Gleichung (2.9)) eine subjektive Wahrscheinlichkeitsverteilung für den Kopplungsparameter bezüglich des zugrunde liegenden GVA-Ereignisses bestimmt.

Berücksichtigung unterschiedlicher Expertenschätzungen

Aufgrund der oftmals nur unvollständig vorliegenden Beschreibungen der GVA-Ereignisse und der sich damit als schwierig erweisenden qualitativen Beurteilungen der beobachteten GVA-Ereignisse und der Übertragbarkeit auf die Zielkomponentengruppe kann die Bewertung von GVA-Ereignissen im Allgemeinen nicht als exakt betrachtet werden. Vielmehr ist die Bewertung von der subjektiven Einschätzung des Experten abhängig.

Damit der Ermittlung der GVA-Wahrscheinlichkeiten eine möglichst realistische Beurteilungsbasis zugrunde liegt, sollten die verschiedenen Beurteilungsalternativen, die

sich aufgrund der verschiedenen Sichtweisen der Experten ergeben, berücksichtigt werden. Deshalb ist es sinnvoll, mehrere Experten in die Beurteilung und Bewertung der vorliegenden GVA-Ereignisse einzubeziehen.

Deshalb werden die Bewertungen mehrerer Experten und die damit verbundenen Unterschiede bei der Ermittlung der interessierenden GVA-Wahrscheinlichkeiten berücksichtigt. Die angewandte Methodik zur Expertenbeurteilung von GVA-Ereignissen besteht aus zwei Teilen:

1. Diskussion der GVA-Ereignisse und der Übertragbarkeit auf die Komponentengruppen in der Zielanlage unter den Experten,
2. Beurteilung der GVA-Ereignisse und der Übertragbarkeit durch die Experten.

Im Teil 1 wird den teilnehmenden Experten jeweils die Fallbeschreibung eines relevanten GVA-Ereignisses vorgelegt. Nach Durchsicht der Fallbeschreibung erfolgt die Diskussion unter den Experten, wobei verschiedene Sichtweisen dargelegt und argumentativ begründet sowie Mehrdeutigkeiten bzw. Unklarheiten in der Beschreibung erörtert und so weit wie möglich geklärt werden sollen. Diskussionsziel der Expertenrunde ist in jedem Fall eine einheitliche Festlegung der Größe der betroffenen Komponentengruppen. Außerdem sollte möglichst ein Konsens bzgl. der qualitativen Beurteilung des vorliegenden GVA-Ereignisses erzielt werden. Die qualitative Beurteilung wird schriftlich festgehalten.

Trotz eingehender Diskussion werden jedoch oftmals unterschiedliche Meinungen bzgl. des GVA-Ereignisses zwischen den Experten bestehen bleiben, die folglich zu unterschiedlichen qualitativen Beurteilungen und den darauf basierenden quantitativen Bewertungen führen. Durch die unterschiedlichen Ereignisbeurteilungen der Experten ergibt sich somit ein vollständigeres Bild von der Unsicherheit bzgl. der Beurteilung des zu bewertenden GVA-Ereignisses und der Übertragbarkeit als im Falle nur eines Experten.

Nach der Diskussion des zu bewertenden GVA-Ereignisses erfolgt im nächsten Schritt (Teil 2 der Vorgehensweise) die eigentliche Expertenbewertung. Dazu führt jeder einzelne Experte für sich die quantitative Bewertung des zugrunde liegenden GVA-Ereignisses durch, indem er – wie bereits beschrieben – jeder Komponente der Gruppe eine der möglichen Schädigungsklassen und den damit verbundenen Schädigungswert zuordnet und die Übertragbarkeit des dem GVA-Ereignisses zugrunde liegenden GVA-

Phänomens auf die Zielkomponentengruppe bestimmt. Diese quantitativen Bewertungen jedes Experten werden schriftlich festgehalten.

Für jeden Experten ergibt sich somit eine subjektive Wahrscheinlichkeitsverteilung des Kopplungsparameters sowie ein Wert des Übertragbarkeitsfaktors. Um die Bewertungsunsicherheit einzubeziehen und dabei die Angaben eines jeden teilnehmenden Experten gleichermaßen zu berücksichtigen, wird die Mischverteilung aus den jeweiligen Verteilungen der GVA-Ausfallwahrscheinlichkeiten gebildet, die sich anhand der Ereignisbewertungen der einzelnen Experten ergeben haben.

Nehmen NE Experten an der Beurteilung der GVA-Ereignisse teil, wird für jeden einzelnen Experten L ($L = 1, \dots, NE$) seiner Beurteilung gemäß eine subjektive Verteilung $p_L(q_{k \setminus r; j})$ der anteiligen GVA-Ausfallwahrscheinlichkeiten nach der im vorigen Abschnitt beschriebenen Vorgehensweise ermittelt. Der gegenüber der Notation in vorigen Abschnitt hinzugefügte Index ' L ' drückt aus, dass es sich um die jeweilige subjektive Wahrscheinlichkeitsverteilung des Experten ' L ' handelt.

Da alle teilnehmenden Experten als gleichermaßen kompetent betrachtet werden, ist jede einzelne Expertenbeurteilung als gleich bedeutend und demzufolge mit gleicher Gewichtung in die Berechnung der resultierenden Verteilung der anteiligen GVA-Ausfallwahrscheinlichkeiten einzubeziehen.

Die resultierende Wahrscheinlichkeitsverteilung der anteiligen GVA-Ausfallwahrscheinlichkeiten ergibt sich somit als Mischung der Verteilungen $p_L(q_{k \setminus r; j})$ über alle Experten, wobei jede der Verteilungen mit dem gleichen Gewicht $1/NE$ in die Mischung eingeht. Damit erhält man für die Wahrscheinlichkeitsverteilung der anteiligen GVA-Ausfallwahrscheinlichkeiten:

$$p(q_{k \setminus r; j}) = \frac{1}{NE} \sum_{L=1}^{NE} p_L(q_{k \setminus r; j}) \quad (2.12)$$

Die gesuchten Verteilungen der GVA-Ausfallwahrscheinlichkeiten $p(q_{k \setminus r})$ lassen sich daraus gemäß Gleichung (2.5) berechnen.

2.2.2 Berücksichtigung der verbleibenden Unsicherheitsquellen

Im Folgenden wird die Berücksichtigung weiterer Unsicherheitsquellen dargestellt. Diese Unsicherheit ist wesentlich auf eine mögliche Inhomogenität von Populationen zurückzuführen. Populationen können statistisch inhomogen sein, da die verschiedenen Komponentengruppen einer Population ein verschiedenes Ausfallverhalten aufweisen können, z. B. aufgrund technischer Unterschiede, verschiedener Betriebsbedingungen oder verschiedener Instandhaltungsstrategien. Das Ausfallverhalten kann sich auch im Laufe der Zeit ändern. Deshalb ist die Übertragung von Betriebserfahrung einer Population von Komponentengruppen in verschiedenen Anlagen, die über einen längeren Zeitraum erfasst wurde, auf eine bestimmte in der PSA modellierte Komponentengruppe zu einem bestimmten Zeitpunkt mit einer Unsicherheit behaftet.

Diese Tatsache, dass über die für die Zielkomponentengruppe zutreffende GVA-Wahrscheinlichkeit eine weitere Unsicherheit vorhanden ist, die zusätzlich zu den im Kopplungsmodell berücksichtigten oben beschriebenen Unsicherheitsquellen existiert, kann wie folgt einbezogen werden: Wenn $q_{k \setminus r}$ die Wahrscheinlichkeit eines (k von r)-GVA bezeichnet, die anhand der in der betrachteten Population aufgetretenen Ereignisse geschätzt wurde, so wird aufgrund der weiteren Unsicherheiten im Allgemeinen die für die Zielkomponentengruppe zutreffende Wahrscheinlichkeit eines (k von r)-GVA, die im Folgenden als $\hat{q}_{k \setminus r}$ bezeichnet wird, von $q_{k \setminus r}$ im allgemeinen abweichen. Der genaue Wert der Abweichung ist jedoch nicht bekannt. Somit lässt sich für $\hat{q}_{k \setminus r}$ nur eine Wahrscheinlichkeitsverteilung angeben. Diese bedingte (d. h. vom Wert $q_{k \setminus r}$ abhängige) Verteilung $p(\hat{q}_{k \setminus r} | q_{k \setminus r})$ quantifiziert die zusätzlichen, im Modell nicht explizit berücksichtigten Unsicherheitsquellen. Jedes Verfahren zur Berücksichtigung weiterer Unsicherheitsquellen muss diese allgemeine Form aufweisen.

Die Verteilung der für die Zielkomponentengruppe gültigen GVA-Wahrscheinlichkeiten $p(\hat{q}_{k \setminus r})$ kann als Integral des Produktes der aus der beobachteten Population bestimmten Verteilung $p(q_{k \setminus r})$ mit der die weiteren Unsicherheitsquellen quantifizierenden bedingten Verteilung $p(\hat{q}_{k \setminus r} | q_{k \setminus r})$ berechnet werden:

$$p(\hat{q}_{k \setminus r}) = \int_0^1 p(\hat{q}_{k \setminus r} | q_{k \setminus r}) p(q_{k \setminus r}) dq_{k \setminus r} \quad (2.13)$$

Im Allgemeinen liegen keine Informationen vor, dass sich die Unsicherheitsquellen in den verschiedenen Komponentengruppen verschieden stark auswirken. Deshalb wird für alle Komponentengruppen dieselbe Verteilung $p(\hat{q}_{k\setminus r}|q_{k\setminus r})$ verwendet.

Angeichts der sehr geringen Anzahl von beobachteten Ereignissen in einzelnen Komponentengruppen ist es nicht möglich, die genaue Form der Verteilung bzw. ihre Charakteristika aus der Betriebserfahrung zu bestimmen. Deshalb sind plausible Annahmen über $p(\hat{q}_{k\setminus r}|q_{k\setminus r})$ zu treffen. Im Folgenden wird angenommen, dass $p(\hat{q}_{k\setminus r}|q_{k\setminus r})$ eine Betaverteilung ist. Zur Bestimmung der Parameter wird angenommen,

- dass der Erwartungswert unter der Verbreiterung erhalten bleibt, d. h. dass die Nicht-Verfügbarkeiten $\hat{q}_{k\setminus r}$ in der Zielkomponentengruppe im Mittel genauso groß sind wie $q_{k\setminus r}$,
- dass die Standardabweichung von $p(\hat{q}_{k\setminus r}|q_{k\setminus r})$ proportional zu $q_{k\setminus r}$ ist. Dies wird gefordert, damit die relative Unsicherheit unabhängig vom absoluten Wert $q_{k\setminus r}$ ist. Dann sind relative Unterschiede zwischen $\hat{q}_{k\setminus r}$ und $q_{k\setminus r}$ gleich wahrscheinlich. Z. B. ist die Wahrscheinlichkeit, dass $\hat{q}_{k\setminus r}$ um 10 % höher ist als $q_{k\setminus r}$, unabhängig von der absoluten Größe von $q_{k\setminus r}$. Der Proportionalitätsfaktor wird im Folgenden als ϱ bezeichnet.
- dass – wie bei der bisherigen, in /FAK 05/ beschriebenen Vorgehensweise – die „verbreiterten Verteilungen“ einen K-Faktor von mindestens 4 aufweisen. Der K-Faktor (auch K95-Faktor) ist definiert als Verhältnis des 95 %-Quantils zum Median:

$$K95: = Q_{95\%}/Q_{50\%}.$$

In /STI 09/ (Abschnitt 6) sind die genannten Annahmen im Detail begründet und mit Ergebnissen bei alternativen Annahmen sowie mit der bisherigen Vorgehensweise anhand von repräsentativen Beispielen der Betriebserfahrung verglichen.

Damit ergibt sich die Verbreiterungsverteilung zu

$$p(\hat{q}_{k\setminus r}|q_{k\setminus r}) = \frac{(1 - \hat{q}_{k\setminus r})^{\beta-1} \hat{q}_{k\setminus r}^{\alpha-1}}{f_{\beta}(\alpha, \beta)} \quad (2.14)$$

wobei $f_\beta(\alpha, \beta)$ die Betafunktion bezeichnet. Die Parameter α und β der Betaverteilung lassen sich aus den oben dargestellten Annahmen bestimmen zu:

$$\alpha = \frac{1 - q_{k \setminus r} - \varrho^2 q_{k \setminus r}}{\varrho^2} \quad (2.15)$$

und

$$\beta = \frac{1 - 2 q_{k \setminus r} - \varrho^2 q_{k \setminus r} + (q_{k \setminus r})^2 + \varrho^2 (q_{k \setminus r})^2}{\varrho^2} \quad (2.16)$$

wobei der die relative Breite der Verteilung bestimmende Faktor ϱ gegeben ist durch

$$\varrho = \varrho^{all} = 0.9463 \quad (2.17)$$

2.3 Diskussion der Modelleigenschaften

Wie oben erwähnt, wurde das Kopplungsmodell entwickelt, um auch bei wenig vorliegender Betriebserfahrung zu Schätzungen von Nicht-Verfügbarkeiten durch GVA zu kommen. Deshalb wurden einige einschränkende Modellannahmen getroffen. Diese Modellannahmen führen neben den erwünschten Eigenschaften auch zu unerwünschten Eigenschaften.

Die zwei unter diesem Aspekt wesentlichen Annahmen sind:

1. Wenn ein GVA-Schock mit Phänomen j aufgetreten ist, werden die einzelnen Komponenten mit einer Wahrscheinlichkeit von η_j (Kopplungsparameter) unverfügbar.
Deshalb genügt die Anzahl bei Auftreten eines GVA-Phänomens ausgefallener Komponenten einer Binomialverteilung (Gleichung (2.2)).
2. Verschiedene GVA-Phänomene haben im Allgemeinen verschiedene Kopplungsparameter.
Deshalb werden die Kopplungsparameter für alle aufgetretenen GVA-Ereignisse unabhängig voneinander geschätzt.

Diese Annahmen und Vorgehensweisen führen dazu, dass das Kopplungsmodell bei einer großen Menge Betriebserfahrung (hohe Ereigniszahl) nicht die erwünschten Konvergenzeigenschaften hat. Wenn die Anzahl beobachteter Ereignisse über alle Grenzen wächst, sollte die Unsicherheit – vor Berücksichtigung der Inhomogenität der Populationen – beliebig klein werden und die geschätzten Werte gegen die wahren Werte konvergieren:

$$p(q_{k \setminus r}) \rightarrow \delta(q_{k \setminus r} - q_{k \setminus r}^*) \quad (2.18)$$

wobei $q_{k \setminus r}^*$ die wahre Nicht-Verfügbarkeit aufgrund GVA und δ die Dirac-Verteilung ist, die charakterisiert ist durch $\int_{-\infty}^{\infty} \delta(x - y) f(x) dx = f(y)$ für beliebige Funktionen $f(y)$.

Dies wird vom Kopplungsmodell im Allgemeinen nicht erfüllt. Wenn die wahre Verteilung der verschiedenen Ausfallkombinationen nicht mit der zentralen Modellannahme des Kopplungsmodells (obige Annahme 1) kompatibel ist, ist eine Konvergenz der Schätzungen gegen die wahren Werte (Gleichung (2.18)) ausgeschlossen.

Dies lässt sich an einem einfachen Beispiel illustrieren: Es treten nur bestimmte Ausfallkombinationen auf, z. B. 1 von 4, 2 von 4 und 3 von 4, nicht aber 4 von 4. Bei jedem Kopplungsparameter $\eta_j \neq 0$ besteht aber eine nichtverschwindende bedingte Wahrscheinlichkeit, dass vier Komponenten ausfallen:

$$p(4 \setminus 4 | \eta_j) = (\eta_j)^4 > 0 \quad (2.19)$$

Somit ist das Kopplungsmodell nicht in der Lage, eine solche Situation zu beschreiben. Die Einschränkung aufgrund der Modellannahme beschränkt sich nicht nur auf solche einfachen Situationen, sondern tritt immer auf, wenn die Unterschiede zwischen den bedingten Wahrscheinlichkeiten der verschiedenen Ausfallkombinationen zu groß sind.

In der Praxis von erheblicher Bedeutung ist eine Folge von Annahme 2. Diese führt dazu, dass die Verteilungen $p(q_{k \setminus r})$ auch dann nicht schmaler werden, wenn die Anzahl der Ereignisse wächst. Dies ist darin begründet, dass die Unsicherheiten des Kopplungsparameter η_j und der Rate λ_j für jedes Ereignis unabhängig von den anderen Ereignissen geschätzt werden. Demzufolge sind die Ergebnisverteilungen identisch, wenn z. B. in einer Beobachtungszeit T ein 2 von 4-Ereignis und ein (3 von 4)-Ereignis beobachtet wurden, wie wenn in einer Beobachtungszeit $10 T$ zehn (2 von 4)-

Ereignisse und zehn (3 von 4)-Ereignisse beobachtet wurden. Dies steht im Widerspruch zu der Tatsache, dass die Evidenz über die Wahrscheinlichkeiten der verschiedenen Ausfallkombinationen durch die zehnfach größere Beobachtungszeit und die Zehnfache Ereignisanzahl gewachsen ist, was zu einer kleineren Schätzunsicherheit und somit geringer Breite der Verteilungen führen müsste.

Diese Eigenschaft kann auch zu nicht-konservativen Schätzabweichungen in dem Sinne führen, dass in der Unsicherheitsverteilung auch zu kleine Werte mit zu hoher Wahrscheinlichkeit vorkommen, da die Breite der Verteilung überschätzt wird.

In Bezug auf den Erwartungswert tritt allerdings keine Nicht-Konservativität auf, wie man an folgendem „worst case“-Szenario sieht: Es werden nur $(r \text{ von } r)$ -Ereignisse mit einer Rate $\lambda_{r \setminus r}$ beobachtet. Alle Ereignisse sind voll übertragbar und die Fehlerentdeckungszeit ist t . Deshalb sollte für sehr lange Beobachtungszeiten T die Wahrscheinlichkeit eines $(r \text{ von } r)$ -GVA gegen $q_{r \setminus r}^* = t \lambda_{r \setminus r}$ streben. Dies ist beim Kopplungsmodell nicht der Fall. Es ergibt sich vielmehr, unabhängig von T , wegen Gleichungen (2.7) und (2.9):

$$\langle q_{r \setminus r} \rangle = \frac{3}{2} t \lambda_{r \setminus r} \int_0^1 \eta^r p(\eta) d\eta \quad (2.20)$$

$$= \frac{3}{2} t \lambda_{r \setminus r} \int_0^1 \eta^r \frac{\Gamma(r+1)}{\Gamma(r+1/2) \Gamma(1/2)} \eta_j^{r-1/2} (1-\eta_j)^{-1/2} d\eta$$

Für $r = 2$ folgt $\langle q_{2 \setminus 2} \rangle = \frac{35}{32} t \lambda_{2 \setminus 2} \approx 1.093 q_{2 \setminus 2}^* > q_{2 \setminus 2}^*$, für $r = 4$ folgt $\langle q_{4 \setminus 4} \rangle = \frac{3861}{3584} t \lambda_{4 \setminus 4} \approx 1.077 q_{4 \setminus 4}^*$. Auch für größere Redundanzgrade ist $\langle q_{r \setminus r} \rangle > q_{r \setminus r}^*$.

In Abb. 2.1 ist die relative Überschätzung $\langle q_{r \setminus r} \rangle / q_{r \setminus r}^*$ dargestellt.

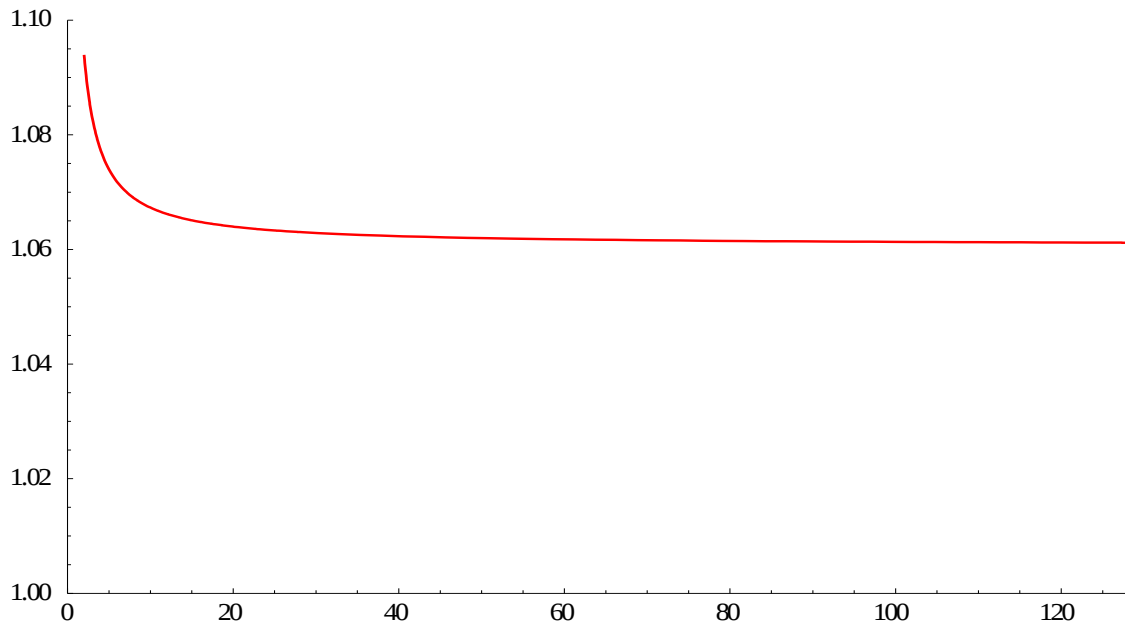


Abb. 2.1 Relative Überschätzung $\langle q_{r\setminus r} \rangle / q_{r\setminus r}^*$ für $r \in [2, 128]$

Die anderen Ausfallkombinationen $\langle q_{i\setminus r} \rangle, i \neq r$ werden ebenfalls überschätzt.

In der Vergangenheit wurden von der GRS bereits zwei Ansätze untersucht, um die oben genannten unerwünschten Eigenschaften zu vermeiden. Diese hatten den Ausgangspunkt, die Annahme 1 beizubehalten, aber die unabhängige Betrachtung aller Ereignisse aufzugeben. Jedoch zeigte sich, dass bei einem Ansatz eine Konvergenz nur bei sehr hohen Ereigniszahlen auftritt, die in der Praxis nicht erreicht wird. Der andere Ansatz war heuristischer Natur; bestimmte Annahmen und Festlegungen, die für die numerischen Ergebnisse wichtig sind, ließen sich nicht nachvollziehbar begründen. Beide Ansätze haben darüber hinaus die oben dargestellten mit der Annahmen einer Binomialverteilung verbundenen Nachteile. Deshalb wird im Rahmen dieses Vorhabens der Ansatz verfolgt, möglichst weitgehend auf Verteilungsannahmen zu verzichten, insbesondere auf die allgemeine Annahme einer Binomialverteilung.

Dies hat zur Folge, dass besondere Verfahren zum sogenannten Mapping erforderlich werden, die eine Übertragung von Ereignissen zwischen Komponentengruppen verschiedener Größen ermöglichen. Dieses Mapping wird im Kopplungsmodell durch die Annahme der Binomialverteilung automatisch geleistet. Separate Mappingverfahren beinhalten allerdings auch den Vorteil, bestimmte Charakteristika der Ereignisse zu berücksichtigen, die beim Kopplungsmodell nicht berücksichtigt werden können, wie bei

Phänomenen, die stets zur Nicht-Verfügbarkeit aller Komponenten führen (z. B. durch Auslegungsfehler). Darauf wird in Abschnitt 5.4.4 eingegangen.

3 Internationale Vorgehensweisen zur Modellierung gemeinsam verursachter Ausfälle

Im Folgenden werden die international gebräuchlichen etablierten Vorgehensweisen zur Modellierung von GVA dargestellt, die eine Quantifizierung mittels der Betriebserfahrung erlauben /MOS 88/, /MOS 89/, /MOS 98/ sowie /WIE 07/. Wesentliche Modelle, die im Folgenden diskutiert werden, sind das Basic-Parameter-Modell (BPM), das Beta-Faktor-Modell (BFM), das Multiple-Greek-Letter-Modell (MGLM), das Alpha-Faktor-Modell (AFM) und das Binomial-Failure-Rate-Model (BFRM).

Bei allen betrachteten Modellansätzen wird davon ausgegangen, dass alle Komponenten einer GVA-Gruppe statistisch äquivalent sind (Symmetrieannahme). Daraus folgt, dass die Nicht-Verfügbarkeiten von einzelnen Komponenten und von Komponentengruppen nicht von den individuellen Komponenten selbst, sondern nur von der Gruppengröße abhängig sind.

Es ist anzumerken, dass das Basic-Parameter-Modell, das Multiple-Greek-Letter-Modell und das Alpha-Faktor-Modell verschiedene Parametrisierungen desselben Modells sind, während Beta-Faktor-Modell, das Binomial-Failure-Rate-Modell und das oben beschriebene Kopplungsmodell jeweils abweichende Modellannahmen beinhalten.

Bevor die einzelnen Modelle erläutert werden, wird eine für alle Modelle verwendete Notation eingeführt.

3.1 Notation

Im Folgenden werden die in den Modellen betrachteten Größen definiert. Es wird eine Notation verwendet, die zu der für das Kopplungsmodell benutzten Notation konsistent ist. Diese weicht von der in der oben genannten Literatur verwendeten Notation ab, lehnt sich aber an sie an. Für die bessere Vergleichbarkeit werden die in der oben genannten Literatur verwendeten Bezeichnungen in Klammern oder Fußnoten angeführt.

$q_{1 \setminus r}$	Wahrscheinlichkeit, dass in einer Komponentengruppe der Größe r genau eine (beliebige) Komponente un verfügbar ist.
---------------------	---

- $q_{k \setminus r}$ Wahrscheinlichkeit, dass in einer Komponentengruppe der Größe r genau k (beliebige) Komponenten aufgrund gemeinsamer Ursache ausgefallen sind, wobei $k \geq 2$.
- $Q_{1 \setminus r}$ Wahrscheinlichkeit, dass in einer Komponentengruppe der Größe r eine bestimmte Komponente aufgrund eines unabhängigen Ausfalls un verfügbar ist (In der o. g. Literatur als $Q_1^{(r)}$ bezeichnet).
- $Q_{k \setminus r}$ Wahrscheinlichkeit, dass in einer Komponentengruppe der Größe r eine bestimmte Gruppe von Komponenten der Größe k aufgrund gemeinsamer Ursache ausgefallen sind, wobei $k \geq 2$ (In der o. g. Literatur als $Q_k^{(r)}$ bezeichnet).
- q_r^{BE} Wahrscheinlichkeit, dass in einer Komponentengruppe der Größe r ein beliebiges Ereignis (Einzelausfall einer der r Komponenten oder GVA) stattgefunden hat.
- Q^{EK} Wahrscheinlichkeit, dass eine bestimmte einzelne Komponente ausgefallen ist (In der o. g. Literatur als Q_t bezeichnet).

Somit werden Nicht-Verfügbarkeiten, die sich auf eine bestimmte Komponente oder Menge von bestimmten Komponenten beziehen, mit einem großen Q bezeichnet und Nicht-Verfügbarkeiten, die sich auf eine beliebige Komponente oder Ausfallkombinationen beziehen, mit kleinen q bezeichnet.

Wenn, wie oben dargestellt, alle Komponenten einer Gruppe statistisch äquivalent sind, lassen sich mittels kombinatorischer Überlegungen Beziehungen zwischen diesen Größen herleiten:

$$q_{k \setminus r} = \binom{r}{k} Q_{k \setminus r} \quad (3.1)$$

$$Q^{EK} = \sum_{k=1}^r \binom{r-1}{k-1} Q_{k \setminus r} \quad (3.2)$$

$$q_r^{BE} = \sum_{k=1}^r \binom{r}{k} Q_{k \setminus r} = \sum_{k=1}^r q_{k \setminus r} \quad (3.3)$$

3.2 Basic-Parameter-Modell

Beim Basic-Parameter-Modell sind die Wahrscheinlichkeiten eines (k von r)-Ausfalls $Q_{k \setminus r}$ unmittelbar die Modellparameter.

In der Literatur wird üblicherweise zusätzlich angenommen, dass die Ausfälle unabhängige Basisereignisse sind (siehe z. B. /MOS 98/). Dann können die Modellparameter unabhängig voneinander aus der Betriebserfahrung geschätzt werden und die Unsicherheitsverteilung faktorisiert. Es ist jedoch auch eine Schätzung der Modellparameter ohne diese Annahme möglich (siehe Abschnitt 4.3).

3.3 Beta-Faktor-Modell

Das Beta-Faktor-Modell beinhaltet folgende Annahmen:

- Ein konstanter Anteil β an der Komponentenausfallwahrscheinlichkeit wird durch GVA-Ereignisse verursacht.
- Bei einem GVA-Ereignis fallen alle Komponenten der Gruppe aus.

3.3.1 Parameter

Das Beta-Faktor-Modell hat zwei Parameter:

$Q^{EK} :=$ Wahrscheinlichkeit, dass eine bestimmte einzelne Komponente ausgefallen ist.

$\beta :=$ Anteil an der Komponentenausfallwahrscheinlichkeit aufgrund GVA

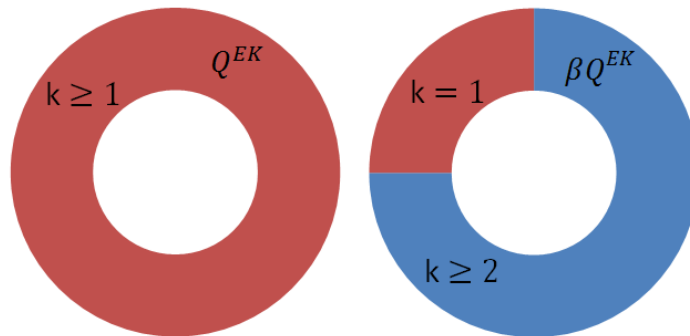


Abb. 3.1 Veranschaulichung der Parameter des Beta-Faktor Modells

links: alle Ausfallereignisse, rechts: Einteilung in unabhängigen Einzelausfall und Ausfall von mindestens zwei Komponenten aufgrund gemeinsamer Ursache

3.3.2 Berechnung von GVA-Wahrscheinlichkeiten

Mit Hilfe des Beta-Faktor-Modells können die Ausfallwahrscheinlichkeiten für eine Komponentengruppe mit der Größe r wie folgt berechnet werden:⁶

$$Q_{k \setminus r} = \begin{cases} (1 - \beta)Q^{EK} & k = 1 \\ 0 & 2 \leq k < r \\ \beta Q^{EK} & k = r \end{cases} \quad (3.4)$$

Vom Beta-Faktor-Modell werden also nur Einzelausfälle und komplette GVA beschrieben. Somit lassen sich auch keine mathematisch-statistisch wohlbegründeten

⁶ In der Literatur /MOS 88/, /MOS 89/, /MOS 98/: $Q_k^{(r)} = \begin{cases} (1 - \beta)Q_t & k = 1 \\ 0 & 2 \leq k < r \\ \beta Q_t & k = r \end{cases}$

Schätzalgorithmen für den allgemeinen Fall (Beobachtung auch partieller GVA) angeben⁷.

Die heutige Bedeutung des Beta-Faktor-Modells beschränkt sich deshalb im Wesentlichen auf Fälle, bei denen bezüglich GVA keine ausreichende Betriebserfahrung vorliegt, jedoch Schätzungen für die Ausfallwahrscheinlichkeit einzelner Komponenten Q^{EK} ermittelbar sind. Hier wird typischerweise der β -Faktor durch Expertenurteil konservativ grob abgeschätzt mit dem Ziel, eine insgesamt konservative Berücksichtigung von GVA in PSA-Rechnungen zu erreichen. Ein Beispiel für einen typischen Anwendungsfall sind software-basierte leittechnische Einrichtungen, bei denen eine Einzelausfallrate aus der Zuverlässigkeit von Bauteilen abgeschätzt werden kann, aber keine belastbare Erfahrung bzgl. GVA vorliegt. Hier kann die Einschätzung, dass Komponentenausfälle meist nicht durch GVA verursacht sind, zu einer groben Abschätzung von z. B. $\beta = 0.1$ führen.

3.4 Multiple-Greek-Letter-Modell

Das Multiple-Greek-Letter-Modell (MGL-Modell) ist eine Erweiterung des Beta-Faktor-Modells. Das MGL-Modell erlaubt auch den Ausfall von Untergruppen der Komponentengruppe, während bei dem Beta-Faktor Modell nur die Möglichkeiten eines unabhängigen Ausfalls oder eines kompletten Ausfalls aller Komponenten bestehen.

3.4.1 Parameter

Das Multiple-Greek-Letter-Modell hat r Parameter:

Q^{EK} := Wahrscheinlichkeit, dass eine bestimmte einzelne Komponente ausgefallen ist.

β := Bedingte Wahrscheinlichkeit für Ausfall von ≥ 2 Komponenten aufgrund von GVA, wenn ein Ausfall von einer Komponente bereits vorliegt.

⁷ Ein approximativer konservativer Ansatz kann darin bestehen, $Q_{1\setminus r}$ und $Q_{r\setminus r}$ wie in Abschnitt 3.2 zu schätzen, wobei aber für die Schätzung von $Q_{r\setminus r}$ nicht die Anzahl der kompletten GVA (d. h. der (r von r)-Ausfälle) $N_{r\setminus r}$ sondern die Anzahl **aller** Ereignisse mit Ausfall mehrerer Komponenten $N_{GVA} = \sum_{i=2}^r N_{i\setminus r}$ verwendet wird, und dann Q^{EK} und β mit Hilfe von Gleichung (3.4) aus diesen beiden Größen berechnet werden.

γ := Bedingte Wahrscheinlichkeit für Ausfall von ≥ 3 Komponenten aufgrund von GVA, wenn ein Ausfall von zwei Komponenten bereits vorliegt.

δ := Bedingte Wahrscheinlichkeit für Ausfall von ≥ 4 Komponenten aufgrund von GVA, wenn ein Ausfall von drei Komponenten bereits vorliegt.

Weitere Parameter werden in Abhängigkeit der Größe der Komponentengruppe entsprechend definiert.

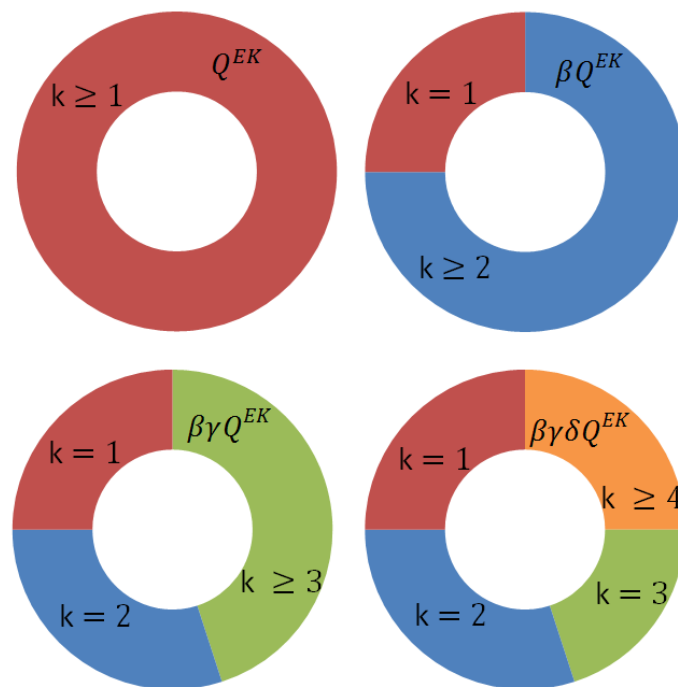


Abb. 3.2 Veranschaulichung des MGL-Parameter

Oben links: Alle Ausfallereignisse. Oben rechts: Einteilung in unabhängigen Einzelausfall und Ausfall von mindestens zwei Komponenten aufgrund gemeinsamer Ursache. Unten links: Einteilung in unabhängigen Einzelausfall, Ausfall von zwei Komponenten und Ausfall von mindestens drei Komponenten aufgrund gemeinsamer Ursache. Unten rechts: Einteilung in unabhängigen Einzelausfall, Ausfall von zwei und drei Komponenten sowie Ausfall von mindestens vier Komponenten aufgrund gemeinsamer Ursache.

3.4.2 Berechnung von GVA-Wahrscheinlichkeiten

Drückt man die $Q_{k \setminus r}$ mit Hilfe der Parameter des Modells aus⁸, so erhält man

$$Q_{k \setminus r} = \frac{1}{\binom{r-1}{k-1}} \prod_{i=1}^k \rho_i (1 - \rho_{k+1}) Q^{EK} \quad (3.5)$$

mit $k = 1, \dots, r = 0$ und $\rho_1 = 1, \rho_2 = \beta, \rho_3 = \gamma, \dots, \rho_{r+1} = 0$.

Es ist anzumerken, dass die Parameter ρ_i des Modells von der Komponentengruppengröße r abhängig sind.

3.5 Alpha-Faktor-Modell

Das Alpha-Faktor Modell stellt eine Reparametrisierung des Basic-Parameter-Modells (siehe Abschnitt 3.2) dar. Es ist das derzeit international am weitesten verbreitetste Modell zur Schätzung von Nicht-Verfügbarkeiten durch GVA aus der Betriebserfahrung.

3.5.1 Parameter

$Q^{EK} :=$ Wahrscheinlichkeit, dass eine bestimmte einzelne Komponente ausgefallen ist.

$\alpha_{1 \setminus r} :=$ Wahrscheinlichkeit, dass ein Ereignis ein (1 von r)-Ausfall ist

$\alpha_{k \setminus r} :=$ Wahrscheinlichkeit, dass ein Ereignis ein (k von r)-GVA ist ($k \geq 2$).

Für die Parameter $\alpha_{k \setminus r}$ gilt:

$$\sum_{k=1}^r \alpha_{k \setminus r} = 1 \quad (3.6)$$

⁸ In der Literatur /MOS 88/, /MOS 89/, /MOS 98/: $Q_k^{(r)} = \frac{1}{\binom{r-1}{k-1}} \prod_{i=1}^k \rho_i (1 - \rho_{k+1}) Q_t$

3.5.2 Berechnung von GVA-Wahrscheinlichkeiten

Die $q_{k \setminus r}$ können mit Hilfe von $\alpha_{k \setminus r}$ und Q^{EK} bzw. q_r^{BE} beschrieben werden.⁹

$$q_{k \setminus r} = \frac{r \alpha_{k \setminus r}}{\sum_{i=1}^r i \alpha_{i \setminus r}} Q^{EK} = \alpha_{k \setminus r} q_r^{BE} \quad (3.7)$$

Diese Gleichung ist für $k \in \{1, \dots, r\}$ gültig. Die Parameter $\alpha_{k \setminus r}$ des Alpha-Faktor-Modells haben also eine einfache Interpretation: Sie sind die Quotienten der Wahrscheinlichkeit eines (k von r)-Ausfalls mit der Wahrscheinlichkeit q_r^{BE} eines beliebigen Ereignisses (Einzelausfall oder Ausfall mehrerer Komponenten).

3.6 Binomial-Failure-Rate-Modell

Beim Binomial-Failure-Rate-Modell (BFR-Modell) werden die Ausfälle in Gruppen eingeteilt:

- Unabhängige Ausfälle,
- Gemeinsame Ausfälle von beliebig vielen Komponenten, die aufgrund von Schockereignissen verursacht werden. Hierbei wird in letale und nicht letale Schockereignisse unterschieden:
 - Bei letalen Schockereignissen fallen alle Komponenten des Systems infolge des Schocks aus.
 - Bei nicht letalen Schockereignissen fallen Komponenten unabhängig voneinander mit einer bestimmten, von der Komponente unabhängigen, Wahrscheinlichkeit aufgrund des Schockereignisses aus.

Aufgrund der Annahme, dass die Ausfallwahrscheinlichkeit bei nicht letalen Schockereignissen für jede Komponente konstant und gleich ist, ergibt sich eine Binomialverteilung der Anzahl ausgefallener Komponenten.

⁹ In der Literatur /MOS 88/, /MOS 89/, /MOS 98/: $Q_k^{(r)} = \frac{k}{\binom{r-1}{k-1}} \frac{\alpha_{k \setminus r}}{\alpha_t} Q_t$, mit $\alpha_t = \sum_{i=1}^r i \alpha_{i \setminus r}$

3.6.1 Parameter

Im Binomial-Failure-Rate-Modell werden folgende Parameter benutzt:

Q_I := Wahrscheinlichkeit für unabhängigen Ausfall,

ϕ := Wahrscheinlichkeit für nicht letalen Schock,

η := Wahrscheinlichkeit für Komponentenausfall bei nicht letalem Schock,

ω := Wahrscheinlichkeit für letalen Schock.

3.6.2 Berechnung von GVA-Wahrscheinlichkeiten

Die Ausfallwahrscheinlichkeit, die k Komponenten betrifft¹⁰, kann mit Hilfe der obigen Größen berechnet werden:

$$q_{k|r} = \binom{r}{k} \begin{cases} Q_I + \phi\eta(1-\rho)^{r-1}, & k = 1 \\ \phi\eta^k(1-\rho)^{r-k}, & 2 \leq k \leq r \\ \phi\eta^r + \omega, & k = r \end{cases} \quad (3.8)$$

Das BFR-Modell enthält die Annahme der Binomialverteilung bei nicht-letalen Schocks und beinhaltet damit noch stärker einschränkende Annahmen als das Kopplungsmodell. Da es gemäß Methodenband zum PSA-Leitfaden /FAK 05/ nur in einer weiterentwickelten Form (wie es das Kopplungsmodell darstellt) angewendet werden soll, wird es im Folgenden nicht weiter betrachtet.

3.7 Übersicht über die verschiedenen Modelle und ihre Parameter

In Tab. 3.1 werden die verschiedenen Modelle und ihre Parameter zusammenfassend aufgeführt.

¹⁰ In der Literatur /MOS 88/, /MOS 89/, /MOS 98/: $Q_k^{(r)} = \begin{cases} Q_I + \phi\eta(1-\rho)^{r-1}, & k = 1 \\ \phi\eta^k(1-\rho)^{r-k}, & 2 \leq k \leq r \\ \phi\eta^r + \omega, & k = r \end{cases}$,

Tab. 3.1 Übersicht der Modelle und ihrer Parameter

Modell	Modellparameter	Berechnung von GVA Wahrscheinlichkeiten
BPM	$Q_{1\setminus r}, Q_{2\setminus r}, \dots, Q_{r\setminus r}$	$Q_{k\setminus r}, k = 1, 2, \dots, r$
Beta-Faktor	Q^{EK}, β	$q_{r\setminus k} = \binom{r}{k} \begin{cases} (1-\beta)Q^{EK} & k = 1 \\ 0 & 2 \leq k < r \\ \beta Q^{EK} & k = r \end{cases}$
MGL	$Q^{EK}, \beta, \gamma, \delta, \dots$	$q_{k\setminus r} = \frac{r}{k} \prod_{i=1}^k \rho_i (1 - \rho_{k+1}) Q^{EK}$ mit $k = 1, \dots, r = 0$ und $\rho_1 = 1, \rho_2 = \beta, \rho_3 = \gamma, \dots, \rho_{r+1} = 0$
Alpha-Faktor	$Q^{EK}, \alpha_{1\setminus r}, \alpha_{2\setminus r}, \dots, \alpha_{r\setminus r}$	$q_{k\setminus r} = r \frac{\alpha_{k\setminus r}}{\alpha_t} Q^{EK} = \alpha_{k\setminus r} q_r^{BE}$
BFR	Q_I, ϕ, η, ω	$q_{k\setminus r} = \binom{r}{k} \begin{cases} Q_I + \phi\eta(1-\rho)^{r-1}, & k = 1 \\ \phi\eta^k(1-\rho)^{r-k}, & 2 \leq k \leq r \\ \phi\eta^r + \omega, & k = r \end{cases}$

4 Schätzung der Modellparameter und GVA-Wahrscheinlichkeiten

Im Folgenden wird die Schätzung der Modellparameter und der GVA-Wahrscheinlichkeiten der Modelle aus der Betriebserfahrung dargestellt. Hierzu wird grundsätzlich einem systematischen Ansatz unter Verwendung der Likelihood-Funktion in Verbindung mit Bayes'schen statistischen Verfahren gefolgt. Weiterhin werden die in /MOS 88/, /MOS 89/, /MOS 98/ angegebenen Vorgehensweisen erwähnt.

Für die Bestimmung der Likelihood-Funktion wird hier folgendes Szenario betrachtet:

- Die r Komponenten einer Komponentengruppe werden gleichzeitig getestet.
- Es wurden z Anforderungen beobachtet, bei denen $n_{1|r}$ (1 von r)-Ausfälle, $n_{2|r}$ (2 von r)-Ausfälle, ... und $n_{r|r}$ (r von r)-Ausfälle (komplette GVA) beobachtet wurden.

Die Schätzung der Modellparameter wird mit Bayes'schen statistischen Verfahren ausgeführt. Dies ermöglicht es insbesondere, die Schätzunsicherheit in Form einer a posteriori-Verteilung der Modellparameter, gegeben die Betriebserfahrung, abzubilden. Aus dieser Verteilung lässt sich dann die Unsicherheitsverteilung der GVA-Wahrscheinlichkeiten bestimmen.

Grundlage der Schätzung ist der Satz von Bayes. Angewendet auf das Problem, Modellparameter θ mit Beobachtungen $\mathfrak{N}$ zu verknüpfen, lautet dieser:

$$p(\theta|\mathfrak{N}) = \frac{p(\mathfrak{N}|\theta)\pi(\theta)}{p(\mathfrak{N})} \quad (4.1)$$

Hierbei ist $\pi(\theta)$ die a priori-Verteilung der Modellparameter. $p(\mathfrak{N}|\theta)$ charakterisiert das stochastische Modell, indem es für alle möglichen Beobachtungen $\mathfrak{N}$ die Wahrscheinlichkeiten angibt, dass $\mathfrak{N}$ beobachtet wird, wenn θ die wahren Modellparameter sind. $p(\mathfrak{N}|\theta)$ wird, wenn es als Funktion von θ betrachtet wird (d. h. bei festem $\mathfrak{N}$) auch als Likelihood-Funktion bezeichnet. Der Normierungsfaktor $p(\mathfrak{N})$ ergibt sich aus den anderen Größen durch Summation bzw. Integration über den gesamten Wertebereich von θ . Für $\theta \in \mathbb{R}$ gilt z. B. $p(\mathfrak{N}) = \int_{-\infty}^{\infty} p(\mathfrak{N}|\hat{\theta})\pi(\hat{\theta}) d\hat{\theta}$.

4.1 Likelihood-Funktion

Da die Ereigniszahlen unter den oben beschriebenen Annahmen einer Multinomialverteilung (auch Polynomialverteilung genannt) mit Parametern $\{1 - \sum_{k=1}^r q_{k \setminus r}, q_{1 \setminus r}, \dots, q_{r \setminus r}\}$ genügen, ist die Likelihood-Funktion¹¹, wenn man diese mit den $q_{k \setminus r}$ ausdrückt:

$$p(\mathfrak{N}|\theta) \propto \left(1 - \sum_{k=1}^r q_{k \setminus r}\right)^{z - \sum_{k=1}^r n_{k \setminus r}} \prod_{k=1}^r (q_{k \setminus r})^{n_{k \setminus r}} \quad (4.2)$$

Gleichung (4.2) hat die Form eines Produktes der Wahrscheinlichkeiten einer bestimmten Beobachtung (gegeben die Modellparameter) hoch die Anzahl der entsprechenden tatsächlichen Beobachtungen. Der erste Term ist z. B. die Wahrscheinlichkeit, dass kein Ausfall auftritt, hoch die Anzahl der Beobachtungen, bei denen kein Ausfall auftrat.

Mit den Definitionen

$$q_{0 \setminus r} = 1 - \sum_{k=1}^r q_{k \setminus r} \quad (4.3)$$

und

$$n_{0 \setminus r} = z - \sum_{k=1}^r n_{k \setminus r} \quad (4.4)$$

lässt sich Gleichung (4.2) einfach schreiben als:

$$p(\mathfrak{N}|\theta) \propto \prod_{k=0}^r (q_{k \setminus r})^{n_{k \setminus r}} \quad (4.5)$$

¹¹ Der Normierungsfaktor wird hier nicht betrachtet, da er für die Bestimmung der a posteriori-Verteilung nicht relevant ist.

Die Unterschiede der Modelle ergeben sich einerseits daraus, wie die $q_{k \setminus r}$ durch Modellparameter dargestellt werden und welche a priori-Verteilungen benutzt werden. Andererseits kommen auch verschiedene Approximationen zur Anwendung. Dies wird im Folgenden dargestellt. Zunächst wird jedoch die unmittelbare Schätzung der GVA-Wahrscheinlichkeiten mittels von Gleichung (4.5) betrachtet.

4.2 Direkte Schätzung der $q_{k \setminus r}$

Gleichung (4.2) bzw. (4.5) legt nahe, die $q_{k \setminus r}$ unmittelbar als Modellparameter zu verwenden. Dies ist jedoch keine international übliche Vorgehensweise.

Trotzdem wird im Folgenden dieser Ansatz entwickelt, um diese naheliegende „naive“ Vorgehensweise mit den anderen Vorgehensweisen vergleichen zu können.

Die konjugierte Verteilung zur Multinomialverteilung ist die Dirichlet-Verteilung. Für die Dichte der Dirichlet-Verteilung mit Parametern $\{a_{0 \setminus r}, a_{2 \setminus r}, \dots, a_{r \setminus r}\}$ gilt:

$$p(q_{1 \setminus r}, q_{2 \setminus r}, \dots, q_{r \setminus r} | a_{0 \setminus r}, a_{2 \setminus r}, \dots, a_{r \setminus r}) = \quad (4.6)$$

$$\frac{1}{B(a_{0 \setminus r}, a_{2 \setminus r}, \dots, a_{r \setminus r})} \prod_{i=0}^r (q_{i \setminus r})^{a_{i \setminus r} - 1}$$

Dabei ist der Normierungsfaktor $B(a_{0 \setminus r}, a_{1 \setminus r}, \dots, a_{r \setminus r})$ gegeben durch

$$B(a_0, \dots, a_r) = \frac{\prod_{i=0}^r \Gamma(a_i)}{\Gamma(\sum_{i=0}^r a_i)} \quad (4.7)$$

wobei Γ die Gammafunktion bezeichnet.

In (4.6) sind als erstes Argument der Dichte p nur $q_{1 \setminus r}$ bis $q_{r \setminus r}$ aufgeführt, da $q_{0 \setminus r}$ nach Gleichung (4.3) durch diese Größen bestimmt wird.

Wählt man die a priori-Verteilung nach dem Verfahren von Jeffreys, so erhält man als a priori-Verteilung eine Dirichlet-Verteilung mit Parametern $a_i = 1/2$:

$$\pi(q_{1\setminus r}, q_{2\setminus r}, \dots, q_{r\setminus r}) \propto \prod_{i=0}^r (q_{i\setminus r})^{-\frac{1}{2}} \quad (4.8)$$

Wenn $n_{k\setminus r}$ (k von r)-Ausfälle beobachtet wurden, ist die a posteriori-Verteilung eine Dirichlet-Verteilung mit den Parametern $a_k = n_{k\setminus r} + 1/2$:

$$p(q_{1\setminus r}, q_{2\setminus r}, \dots, q_{r\setminus r} | N) = \frac{\prod_{k=0}^r (q_{k\setminus r})^{n_{k\setminus r} - \frac{1}{2}}}{B\left(n_{0\setminus r} + \frac{1}{2}, n_{1\setminus r} + \frac{1}{2}, \dots, n_{r\setminus r} + \frac{1}{2}\right)} =$$

$$\frac{(1 - \sum_{k=1}^r q_{k\setminus r})^{z - \sum_{k=1}^r n_{k\setminus r} - 1/2} \prod_{k=1}^r (q_{i\setminus r})^{n_{k\setminus r} - 1/2}}{B(z - \sum_{k=1}^r n_{k\setminus r} + 1/2, n_{1\setminus r} + 1/2, \dots, n_{r\setminus r} + 1/2)}$$

Die Marginalverteilungen sind Betaverteilungen mit Parametern $n_{i\setminus r} + 1/2$ und $z - n_{i\setminus r} + \frac{r}{2}$:

$$p(q_{i\setminus r} | N) = \frac{(q_{i\setminus r})^{n_{i\setminus r} - 1/2} (1 - q_{i\setminus r})^{z - n_{i\setminus r} + \frac{r}{2} - 1}}{B(n_{i\setminus r} + 1/2, z - n_{i\setminus r} + \frac{r}{2})} \quad (4.10)$$

mit Normierungsfaktor B aus (4.7)

Für die Erwartungswerte gilt

$$\langle q_{k\setminus r} \rangle = \frac{n_{k\setminus r} + 1/2}{\sum_{i=0}^r (n_{i\setminus r} + 1/2)} = \frac{2n_{k\setminus r} + 1}{2z + r + 1} \quad (4.11)$$

Daraus folgt insbesondere für die Wahrscheinlichkeit, dass keine Komponente ausfällt:

$$\langle q_{0\setminus r} \rangle = \frac{2(z - \sum_{i=1}^r n_{i\setminus r}) + 1}{2z + r + 1} \quad (4.12)$$

Da in der Praxis stets $r \ll z$ gilt, folgt aus (4.11) für $k > 0$:

$$\langle q_{k\setminus r} \rangle \approx \frac{n_{k\setminus r} + \frac{1}{2}}{z} \quad (4.13)$$

Wenn bei z Anforderungen keinerlei Ausfälle beobachtet wurden (Nullfehlerstatistik), so folgt

$$\langle q_{k \setminus r} \rangle \approx \frac{1}{2z} \quad \text{für} \quad k \in \{1, 2, \dots, r\} \quad (4.14)$$

4.3 Basic-Parameter-Modell

Im Folgenden werden zwei Ansätze zur Bestimmung der Modellparameter und GVA-Wahrscheinlichkeiten mittels der Likelihood-Funktion bzw. unter Verwendung vereinfachender Annahmen dargestellt

4.3.1 Schätzung unter Verwendung der Likelihood-Funktion

Wenn man die Likelihood-Funktion als Funktion der Parameter des Basic-Parameter-Modells ausdrückt, erhält man

$$p(N|\theta) \propto \prod_{k=0}^r \left(\binom{r}{k} Q_{k \setminus r} \right)^{n_{k \setminus r}} \quad (4.15)$$

wobei analog zu Gleichung (4.3) $Q_{0 \setminus r}$ definiert wird als

$$Q_{0 \setminus r} = 1 - \sum_{k=1}^r \binom{r}{k} Q_{k \setminus r} \quad (4.16)$$

Der Wertebereich ist $Q_{k \setminus r} \in [0, 1/\binom{r}{k}]$ mit $\sum_{k=0}^r \binom{r}{k} Q_{k \setminus r} = 1$.

Im Gegensatz zu Gleichung (4.5) ist die Gleichung (4.15) komplizierter, da die Binomialfaktoren $\binom{r}{k}$ in der Summe stehen. Dies legt nahe, eine Variablentransformation durchzuführen, um diesen Faktor zu entfernen. Die Transformation, die dies leistet, ist $\{Q_{k \setminus r}\} \rightarrow \{q_{k \setminus r}\}$. Damit folgt man wieder der im vorigen Abschnitt beschriebenen Vorgehensweise. Die $q_{k \setminus r}$ können also als natürliche Parameter des Modells aufgefasst werden.

Die Verteilungen der $Q_{k \setminus r}$ erhält man, indem man für die Ergebnisverteilungen eine Rücktransformation durchführt. So erhält man z. B. für die Marginalverteilungen aus Gleichung (4.10):

$$p(Q_{i \setminus r} | N) = \binom{r}{i} \begin{cases} \frac{((\binom{r}{i} Q_{i \setminus r})^{n_{i \setminus r} - 1/2} (1 - (\binom{r}{i} Q_{i \setminus r}))^{z - n_{i \setminus r} + \frac{r}{2} - 1})}{B(n_{i \setminus r} + 1/2, z - n_{i \setminus r} + \frac{r}{2})} & \text{für } 0 < Q_{i \setminus r} < \frac{1}{\binom{r}{i}} \\ 0 & \text{sonst} \end{cases} \quad (4.17)$$

Für die Erwartungswerte gilt:

$$\langle Q_{k \setminus r} \rangle = \frac{1}{\binom{r}{k}} \frac{n_{k \setminus r} + 1/2}{z + r/2} \quad (4.18)$$

Da in der Praxis stets $z \gg r$ gilt, folgt

$$\langle Q_{k \setminus r} \rangle \approx \frac{n_{k \setminus r} + 1/2}{\binom{r}{k} z} \quad (4.19)$$

Wenn bei z Anforderungen keinerlei GVA beobachtet wurden (Nullfehlerstatistik), folgt

$$\langle Q_{k \setminus r} \rangle \approx \frac{1}{2 \binom{r}{k} z} \quad \text{für } k \in \{1, 2, \dots, r\} \quad (4.20)$$

Diese Ergebnisse entsprechen den Gleichungen (4.11) bis (4.14).

4.3.2 Vereinfachte Schätzung

In /MOS 98/ werden die Ausfälle der verschiedenen Ausfallkombinationen des Basic-Parameter-Modell als unabhängige Elementarereignisse angesehen. Die Wahrscheinlichkeiten eines (k von r)-Ausfalls $Q_{k \setminus r}$ sind unmittelbar die Modellparameter. Deshalb können die Modellparameter unabhängig voneinander aus der Betriebserfahrung geschätzt werden und ihre Realisierungen bei Unsicherheitsanalysen unabhängig voneinander aus ihren jeweiligen Verteilungen gezogen werden.

Bei der Wahl einer a priori-Verteilung nach dem Ansatz von Jeffreys /BOX 73/ gilt:

$$\pi(Q_{k \setminus r}) \propto \frac{1}{\sqrt{Q_{k \setminus r}(1 - Q_{k \setminus r})}} \quad (4.21)$$

und bei Beobachtungen von $n_{k \setminus r}$ (k von r)-Ausfällen bei N gleichzeitigen Anforderungen der Komponenten einer Komponentengruppe der Größe k wird in /MOS 98/ eine a posteriori-Verteilung hergeleitet unter impliziter Verwendung der Annahme, dass die GVA-Wahrscheinlichkeiten nicht von den betroffenen Komponenten, sondern nur von der Anzahl abhängen und der Tatsache, dass jede der Anforderungen einer Gruppe mit $\binom{r}{k}$ Möglichkeiten verbunden ist, dass k von r Anforderungen ausfallen. Es ergibt sich als a posteriori-Verteilung eine Betaverteilung mit den Parametern $n_{k \setminus r} + 1/2$ und $\binom{r}{k} z - n_{k \setminus r} + 1/2$. Für die Dichte dieser Verteilung gilt /MOS 98/:

$$p(Q_{k \setminus r} | \mathfrak{N}) = \frac{\Gamma(\binom{r}{k} z + 1)}{\Gamma(n_{k \setminus r} + 1/2) \Gamma(\binom{r}{k} z - n_{k \setminus r} + 1/2)} Q_{k \setminus r}^{n_{k \setminus r} - 1/2} (1 - Q_{k \setminus r})^{\binom{r}{k} z - n_{k \setminus r} - 1/2} \quad (4.22)$$

Somit werden $n_{k \setminus r}$ Ausfälle auf $\binom{r}{k} N$ „effektive Anforderungen“ bezogen. Die zugrundeliegende Argumentation ist nicht vollständig schlüssig, da tatsächlich nur N Anforderungen der Komponentengruppe vorliegen, und die Ausfälle verschiedener Kombinationen individueller Komponenten nicht als unabhängig angesehen werden können. Durch Gleichung (4.22) ist nicht berücksichtigt, dass $Q_{k \setminus r} \leq 1/\binom{r}{k}$ gelten muss. Demzufolge hat die Verteilung eine geringfügig größere Standardabweichung als (4.17).

Für die Erwartungswerte gilt

$$\langle Q_{k \setminus r} \rangle = \frac{n_{k \setminus r} + 1/2}{\binom{r}{k} z + 1} \quad (4.23)$$

Da in der Praxis stets $z \gg 1$ gilt, folgt:

$$\langle Q_{k \setminus r} \rangle \approx \frac{n_{k \setminus r} + 1/2}{\binom{r}{k} z} \quad (4.24)$$

Umgerechnet auf die $q_{k \setminus r}$ folgt

$$\langle q_{k \setminus r} \rangle = \binom{r}{k} \langle Q_{k \setminus r} \rangle \approx \frac{n_{k \setminus r} + \frac{1}{2}}{z} \quad (4.25)$$

Wenn bei z Anforderungen keinerlei GVA beobachtet wurden (Nullfehlerstatistik), folgt

$$\langle Q_{k \setminus r} \rangle \approx \frac{1}{2 \binom{r}{k} z} \quad \text{für} \quad k \in \{1, 2, \dots, r\} \quad (4.26)$$

Damit sind die Erwartungswerte approximativ gleich wie in Abschnitt 4.3.1 (Gleichung (4.18) und (4.19)). Allerdings sind die a posteriori-Verteilungen der Parameter hier unabhängig, während die $Q_{i \setminus r}$ bei dem Verfahren aus Abschnitt 4.3.1 einer Verbundwahrscheinlichkeit genügen, die nicht faktorisiert. Die Breite der entsprechenden Marginalverteilung ist dort auch geringfügig geringer.

4.4 Alpha-Faktor-Modell

Schätzungen für die Parameter des Alpha-Faktor-Modells lassen sich aus der oben diskutierten Likelihood-Funktion

$$p(\mathfrak{N}|\theta) \propto \prod_{k=0}^r (q_{k \setminus r})^{n_{k \setminus r}} = \left(1 - \sum_{k=1}^r q_{k \setminus r}\right)^{z - \sum_{k=1}^r n_{k \setminus r}} \prod_{k=1}^r (q_{k \setminus r})^{n_{k \setminus r}} \quad (4.27)$$

ableiten. Dazu wird die Relation $q_{k \setminus r} = \alpha_{k \setminus r} q_r^{BE}$ (Gleichung 3.6) für $k \in \{1, \dots, r\}$ benutzt:

$$p(\mathfrak{N}|\theta) \propto (-q_r^{BE})^{z - n^{BE}} \prod_{k=1}^r (\alpha_{k \setminus r} q_r^{BE})^{n_{k \setminus r}} = \quad (4.28)$$

$$(1 - q_r^{BE})^{z - n^{BE}} (q_r^{BE})^{n^{BE}} \prod_{k=1}^r (\alpha_{k \setminus r})^{n_{k \setminus r}}$$

wobei $n^{BE} = \sum_{k=1}^r n_{k \setminus r}$ die Anzahl beliebiger Ereignisse (d. h. Einzelausfälle und GVA) ist.

Die Likelihood-Funktion zerfällt in zwei Faktoren:

$$p(\mathfrak{N}|\theta) = p(n^{BE}|q_r^{BE}) p(\mathfrak{N}|\{\alpha_{1\setminus r}, \dots, \alpha_{r-1\setminus r}\}) \quad (4.29)$$

mit

$$p(n^{BE}|q_r^{BE}) \propto (1 - q_r^{BE})^{z-n^{BE}} (q_r^{BE})^{n^{BE}} \quad (4.30)$$

und

$$p(\mathfrak{N}|\{\alpha_{1\setminus r}, \dots, \alpha_{r-1\setminus r}\}) \propto \prod_{k=1}^r (\alpha_{k\setminus r})^{n_{k\setminus r}} \quad (4.31)$$

wobei $\sum_{k=1}^r \alpha_{k\setminus r} = 1$ gilt, weshalb in (4.31) nur $r - 1$ Argumente angegeben sind.

Da die Likelihood-Funktion faktorisiert, ist es möglich, q_r^{BE} und $\{\alpha_{1\setminus r}, \dots, \alpha_{r-1\setminus r}\}$ unabhängig voneinander zu schätzen.

4.4.1 Schätzung der Wahrscheinlichkeit von Ereignissen q_r^{BE}

Die Schätzung der Wahrscheinlichkeit von Ereignissen q_r^{BE} wird durch die entsprechende Likelihood-Funktion

$$p(n^{BE}|q_r^{BE}) \propto (1 - q_r^{BE})^{z-n^{BE}} (q_r^{BE})^{n_{k\setminus r}} \quad (4.32)$$

bestimmt. Wenn man die a priori-Verteilung nach dem Verfahren von Jeffreys als

$$\pi(q_r^{BE}) \propto \frac{1}{\sqrt{q_r^{BE}(1 - q_r^{BE})}} \quad (4.33)$$

wählt, ergibt sich als a posteriori-Verteilung eine Betaverteilung mit den Parametern $n^{BE} + 1/2$ und $z - n^{BE} + 1/2$:

$$p(q_r^{BE}) = \frac{\Gamma(z + 1)}{\Gamma(n^{BE} + \frac{1}{2}) \Gamma(z - n^{BE} + \frac{1}{2})} \times (q_r^{BE})^{n^{BE} - \frac{1}{2}} (1 - q_r^{BE})^{z - n^{BE} - \frac{1}{2}} \quad (4.34)$$

Für die Erwartungswerte gilt:

$$\langle q_r^{BE} \rangle = \frac{n^{BE} + \frac{1}{2}}{z + 1} \approx \frac{n^{BE} + \frac{1}{2}}{z} \quad (4.35)$$

da in der Praxis $z \gg 1$ gilt. Wurden keine Ereignisse beobachtet, so folgt

$$\langle q_r^{BE} \rangle = \frac{1}{2(z + 1)} \approx \frac{1}{2z}. \quad (4.36)$$

4.4.2 Schätzung der Alpha-Faktoren

Die Alpha-Faktoren werden mit Hilfe der entsprechenden Likelihood-Funktion

$$p(\mathfrak{N} | \{\alpha_{1 \setminus r}, \dots, \alpha_{r-1 \setminus r}\}) \propto \prod_{k=1}^r (\alpha_{k \setminus r})^{n_{k \setminus r}} \quad (4.37)$$

bestimmt. Die Konjugierte Verteilung zur Multinomialverteilung ist die Dirichlet-Verteilung. Für die Dichte der Dirichlet-Verteilung gilt, wie oben bereits erwähnt:

$$p(\alpha_{1 \setminus r}, \alpha_{2 \setminus r}, \dots, \alpha_{r-1 \setminus r} | a_{1 \setminus r}, a_{2 \setminus r}, \dots, a_{r \setminus r}) = \frac{1}{B(a_{1 \setminus r}, a_{2 \setminus r}, \dots, a_{r \setminus r})} \prod_{i=1}^r (\alpha_{i \setminus r})^{a_{i \setminus r}-1} \quad (4.38)$$

Der Normierungsfaktor $B(a_{1 \setminus r}, a_{2 \setminus r}, \dots, a_{r \setminus r})$ ist in Gleichung (4.7) angegeben.

In (4.38) sind wiederum als erstes Argument der Dichte p nur $\alpha_{1 \setminus r}$ bis $\alpha_{r-1 \setminus r}$ aufgeführt, da $\sum_{k=1}^r \alpha_{k \setminus r} = 1$ gilt.

Wählt man die a priori-Verteilung nach dem Verfahren von Jeffreys als:

$$\begin{aligned} \pi(\alpha_{1 \setminus r}, \alpha_{2 \setminus r}, \dots, \alpha_{r-1 \setminus r}) &\propto \prod_{i=1}^r (\alpha_{i \setminus r})^{-\frac{1}{2}} \\ &= \prod_{i=1}^{r-1} (\alpha_{i \setminus r})^{-\frac{1}{2}} \left(1 - \sum_{k=1}^{r-1} \alpha_{k \setminus r}\right)^{-\frac{1}{2}} \end{aligned} \quad (4.39)$$

ist die a posteriori-Verteilung eine Dirichlet-Verteilung mit den Parametern $a_k = n_{k \setminus r} + 1/2$:

$$p(\alpha_{1 \setminus r}, \alpha_{2 \setminus r}, \dots, \alpha_{r-1 \setminus r} | \mathfrak{N}) = \frac{\prod_{i=1}^r (\alpha_{i \setminus r})^{n_{k \setminus r} - \frac{1}{2}}}{B(n_{1 \setminus r} + 1/2, n_{2 \setminus r} + 1/2, \dots, n_{r \setminus r} + 1/2)} \quad (4.40)$$

Für die Erwartungswerte der Alpha-Faktoren gilt:

$$\langle \alpha_{k \setminus r} \rangle = \frac{n_{k \setminus r} + 1/2}{\sum_{i=1}^r (n_{k \setminus r} + 1/2)} = \frac{n_{k \setminus r} + 1/2}{n^{BE} + r/2} \quad (4.41)$$

Wurden keine Ereignisse beobachtet, so folgt $\langle \alpha_{k \setminus r} \rangle = \frac{1}{r}$.

4.4.3 Erwartungswerte der GVA-Wahrscheinlichkeiten

Da die Alpha-Faktoren $\alpha_{k \setminus r}$ und die Wahrscheinlichkeit von Ereignissen q_r^{BE} unabhängig sind, gegeben die Beobachtungen, ist der Erwartungswert der GVA-Wahrscheinlichkeiten das Produkt der Erwartungswerte der Alpha-Faktoren und der Wahrscheinlichkeit von Ereignissen:

$$\langle q_{k \setminus r} \rangle = \langle q_r^{BE} \alpha_{k \setminus r} \rangle = \langle q_r^{BE} \rangle \langle \alpha_{k \setminus r} \rangle = \frac{(n^{BE} + \frac{1}{2})(n_{k \setminus r} + \frac{1}{2})}{(z + 1)(n^{BE} + \frac{r}{2})} \quad (4.42)$$

Wenn keine Ausfälle beobachtet wurden, so folgt

$$\langle q_{k \setminus r} \rangle \approx \frac{1}{2rz} \quad (4.43)$$

Damit ist in diesem Fall der Erwartungswert um den Faktor r , d. h. den Redundanzgrad, kleiner als bei direkter Schätzung der $q_{k \setminus r}$ (Gleichung (4.14)). Die Ursache dafür liegt in den verschiedenen a priori-Annahmen (Gleichung (4.8) gegenüber Gleichungen (4.33) und (4.39)). Die Likelihood-Funktion (ausgedrückt in den $q_{k \setminus r}$) ist identisch.

Durch den a priori des Alpha-Faktor-Modells wird die Symmetrie zwischen den verschiedenen möglichen Ereignissen der Anforderung einer Komponentengruppe aufgehoben. Das Ergebnis „kein Ausfall“, das in Abschnitt 4.2 (siehe z. B. Gleichung (4.6)

völlig äquivalent zu den einzelnen möglichen Ausfallkombinationen behandelt wurde, wird speziell herausgehoben. Dies kann man veranschaulichen, indem man den Test als zweistufigen hierarchischen Prozess auffasst. In der ersten Stufe wird differenziert nach „Ausfall“ oder „kein Ausfall“. In der zweiten Stufe (falls ein Ausfall vorliegt), wird die Ausfallkombination bestimmt. Dies ist in Abb. 4.1 dargestellt.

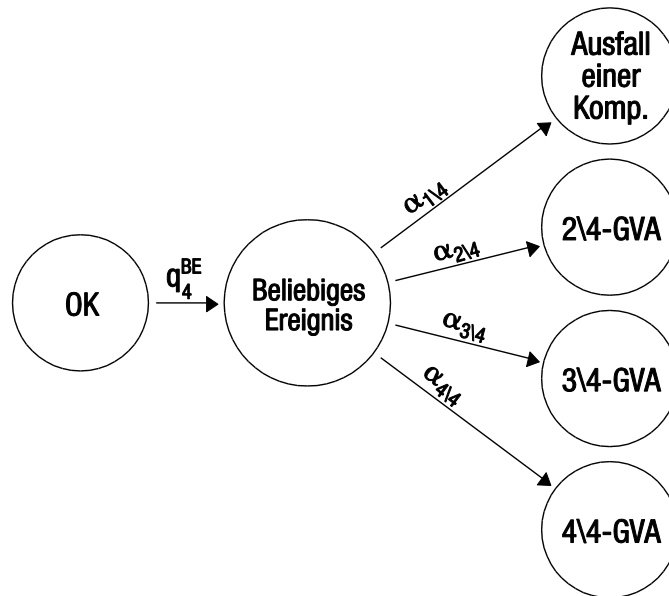


Abb. 4.1 Struktur des Alpha-Faktor-Modells

Es ist zu betonen, dass die Unterschiede im Schätzergebnis nicht von der verschiedenen Modelldarstellung bzw. Parametrierung, sondern nur durch die dadurch nahegelegte Wahl der a priori-Verteilungen herrühren. Deshalb gleichen sich auch die Ergebnisse im Grenzfall beliebig vieler Betriebserfahrung an.

4.5 Beta-Faktor-Modell

Beim Beta-Faktor-Modell ist die Likelihood-Funktion

$$\begin{aligned}
 p(\mathfrak{N}|\theta) &\propto \left(1 - \sum_{k=1}^r q_{k \setminus r}\right)^{z - \sum_{k=1}^r n_{k \setminus r}} \prod_{k=1}^r (q_{k \setminus r})^{n_{k \setminus r}} \\
 &= (1 - k(1 - \beta)Q^{EK} - \beta Q^{EK})^{z - \sum_{k=1}^r n_{k \setminus r}} \\
 &\quad (k(1 - \beta)Q^{EK})^{n_{1 \setminus r}} \left(\prod_{k=2}^{r-1} \delta_{n_{k \setminus r}, 0} \right) (\beta Q^{EK})^{n_{r \setminus r}}
 \end{aligned} \tag{4.44}$$

wobei $\delta_{a,b}$ das Kronecker-Symbol¹² ist.

Die Likelihood-Funktion ist identisch 0, falls $n_{k \setminus r} > 0$ für ein k mit $1 > k > r$ gilt, d. h. falls nicht vollständige GVA aufgetreten sind, da das Modell diesen allgemeinen Fall nicht beschreiben kann. Somit lassen sich auch keine Schätzalgorithmen für diesen Fall aus der Likelihood-Funktion ableiten.

Für den Fall, dass die Bedingung $\forall_{1 > k > r: n_{k \setminus r} = 0$ gilt (z. B. für $r = 2$), könnte man analog wie in Abschnitt 4.3.1 vorgehen und zunächst die Verteilung der $q_{k \setminus r}$ schätzen und daraus die Verteilung der Modellparameter ableiten.

Da das Beta-Faktor-Modell heute für die Schätzung von GVA-Wahrscheinlichkeiten aus der Betriebserfahrung nicht mehr gebräuchlich ist, wird es nicht vertieft diskutiert.

4.6 Multiple-Greek-Letter-Modell

Für das Multiple-Greek-Letter-Modell (MGL-Modell) hat die Likelihood-Funktion eine sehr komplizierte Gestalt. Die Parameter genügen keinen wohlbeschriebenen Wahr-

¹² Das Kronecker-Symbol ist definiert durch $\delta_{a,b} = \begin{cases} 1 & \text{für } a = b \\ 0 & \text{sonst} \end{cases}$

scheinlichkeitsverteilungen. Die Schätzung der Parameter wird nach /MOS 89/ deshalb indirekt vorgenommen, indem für die Parameter des Alpha-Faktor-Modells a posteriori-Verteilungen bestimmt werden und die Verteilungen der Parameter des Multiple-Greek-Letter-Modells aus diesen berechnet werden. Alternativ kann man analog wie in Abschnitt 4.3.1 vorgehen und zunächst die Verteilung der $q_{k \setminus r}$ schätzen und daraus die Verteilung der Modellparameter ableiten. Deshalb erübrigt sich hier eine weitere Diskussion.

5 Entwicklung von Modellen zur Quantifizierung von GVA aus der deutschen Betriebserfahrung

In den folgenden Abschnitten werden Modellansätze zur Weiterentwicklung der Schätzung von GVA-Wahrscheinlichkeiten aus der deutschen Betriebserfahrung entwickelt und diskutiert. Hierzu werden zunächst die zu erfüllenden Randbedingungen dargestellt.

5.1 Randbedingungen der Modelle

Für die Quantifizierung von GVA aus der deutschen Betriebserfahrung sind bestimmte Randbedingungen zu berücksichtigen, die im Folgenden dargestellt werden.

Entsprechend dem deutschen PSA-Leitfaden /BMU 05/ und seinen nachgeordneten Fachbänden zu PSA-Methoden /FAK 05/ und -Daten /FAK 05a/ sollen Basisereignisse wenn immer möglich aus der anlagenspezifischen Betriebserfahrung quantifiziert werden. Dies ist für unabhängige Ausfälle in den meisten Fällen möglich, nicht jedoch für GVA. Anlagenspezifische Schätzungen der Unverfügbarkeit durch unabhängige Ausfälle können anhand der beim jeweiligen Betreiber vorliegenden Betriebserfahrung bestimmt werden. Für GVA werden demgegenüber generische Daten verwendet, die auf der Betriebserfahrung der Gesamtheit der deutschen Kernkraftwerke basieren, da GVA seltene Ereignisse sind und deshalb die Anzahl der in einzelnen Kraftwerken aufgetretenen Ereignisse sehr gering ist. Deshalb wurden im Rahmen der systematischen Erfassung der entsprechenden Betriebserfahrung nur Ereignisse mit GVA und potentiellen GVA (Ereignisse, bei denen nur eine Komponente ausgefallen ist, jedoch weitere Komponenten Schädigungen aufgrund systematischer Ursache aufweisen) erfasst, bewertet und in das GVA-Datenbanksystem der GRS aufgenommen. Andere Einzelausfälle wurden nicht einbezogen.

Deshalb sind Modelle bzw. Schätzverfahren, bei denen die Kenntnis der Anzahl von Einzelausfällen notwendig für die Schätzung von GVA-Wahrscheinlichkeiten ist, für die deutschen Rahmenbedingungen nicht geeignet. Dies betrifft insbesondere das Alpha-Faktor-Modell (Abschnitt 3.5). Bei einer Verwendung von anlagenspezifischer und generischer Betriebserfahrung, wie sie sich aus dem deutschen PSA-Regelwerk ergibt, wäre eine Verwendung von Einzelfehlern bei der Schätzung von GVA-Wahrscheinlichkeiten auch insofern problematisch, als sie zu einer intransparenten Vermischung

von anlagenspezifischen und generischen Informationen beiträgt, weil dann die Schätzungen von GVA-Wahrscheinlichkeiten von den Einzelfehlern in anderen Kraftwerken abhängen, während die geschätzten Wahrscheinlichkeiten von Einzelfehlern dies nicht tun. Insbesondere die Verwendung von Schätzungen der Wahrscheinlichkeit eines beliebigen Ereignisses (q_r^{BE} in Gleichung (3.6) aus anlagenspezifischer Betriebserfahrung und Schätzung der Alpha-Faktoren ($\alpha_{k \setminus r}$ in Gleichung (3.6) anhand generischer Betriebserfahrungen ist nicht sachgerecht, da dann z. B. das Auftreten vieler Einzelfehler in anderen Anlagen ein Unterschätzen der GVA-Wahrscheinlichkeiten verursachen kann.

Deshalb werden im Folgenden nur Ansätze betrachtet, die ein Schätzen der GVA-Wahrscheinlichkeiten ohne Bezugnahme auf unabhängige Ausfälle (so genannte absolute Modelle) ermöglichen. Hierdurch ist die erforderliche klare Trennung anlagenspezifischer und anlagenübergreifender (generischer) Betriebserfahrung gegeben.

Diese Randbedingungen wurden auch in der bisherigen, in Abschnitt 2 beschriebenen Vorgehensweise berücksichtigt.

Aus den genannten Randbedingungen können in Verbindung mit den bereits zuvor diskutierten Eigenschaften der bisherigen Vorgehensweise zur GVA-Quantifizierung folgende anzustrebende Eigenschaften abgeleitet werden:

1. Die mit den einschränkenden Annahmen des Kopplungsmodells verbundenen Nachteile (siehe Abschnitt 2.3) sollen vermieden werden.
2. Die Berücksichtigung von Schätzunsicherheiten soll im bisherigen Umfang erhalten bleiben. Dies umfasst
 - a) statistische Unsicherheiten,
 - b) Unsicherheiten der Expertenbewertungen,
 - c) Interpretationsunsicherheiten der Komponentenschädigungen,
 - d) mit der möglichen Inhomogenität der Populationen verknüpfte Unsicherheiten und
 - e) sonstige Unsicherheiten.

3. Das Modell muss mit der deutschen Betriebserfahrung in Bezug auf GVA, wie sie in der GVA-Datenbank erfasst ist, quantifizierbar sein. Relevante Eigenschaften dieser Daten sind, wie oben bereits dargestellt,
 - a) Bei den Ereignissen liegen jeweils quantitative Abschätzungen der Komponentenschädigungen durch mehrere Experten vor.
 - b) Auch potentielle GVA sind enthalten, d. h. Ereignisse mit systematischer Ursache, bei denen höchstens ein Ausfall und sonst nur Komponentenschädigungen aufgetreten sind.
 - c) Einzelausfälle sind in der Datenbank nicht enthalten, da für die Quantifizierung von Einzelfehlern im Allgemeinen anlagenspezifische Betriebserfahrung verwendet wird.
 - d) Keine Zeitpunkte oder Anzahlen der Anforderungen der einzelnen Komponenten sind verfügbar.
4. Da nicht für alle in der PSA zu modellierenden Komponentengruppengrößen hinreichend Betriebserfahrung (Komponentenart und Gruppengröße) vorliegt, muss ggf. Mapping möglich sein.

Der Punkt d) impliziert, dass GVA-Raten geschätzt werden müssen.

Im Folgenden werden verschiedene entwickelte Modellansätze diskutiert.

5.2 Beschreibung der Modelle

5.2.1 Modell A

Beim Modell A werden die GVA-Ereignisse mit verschiedenen Ausfallkombinationen als unabhängige Elementarereignisse angesehen. Diese Annahme ist berechtigt, wenn GVA-Ereignisse sehr selten sind, d. h. das Auftreten zweier GVA an einer Komponentengruppe innerhalb einer Fehlerentdeckungszeit vernachlässigt werden kann. (k von r)-GVA treten mit einer Rate $\lambda_{k \setminus r}$ auf. Die Nicht-Verfügbarkeit aufgrund eines (k von r)-GVA ist, wobei t die Fehlerentdeckungszeit bezeichnet:

$$q_{k \setminus r} = t \lambda_{k \setminus r} \quad (5.1)$$

In Abb. 5.1 ist die Modellstruktur des Modells A am Beispiel einer Komponentengruppe mit vier Komponenten dargestellt. Die Komponentengruppe befindet sich zunächst in einem nicht durch GVA-Phänomene beeinflussten Zustand; alle Komponenten sind verfügbar („ok“). Mit einer Rate $\lambda_{k \setminus r}$ geht es in einen Zustand „ $k \setminus r$ -GVA“ über, bei dem k der vorhandenen r Komponenten durch GVA un verfügbar sind ($k = 2, \dots, r$).

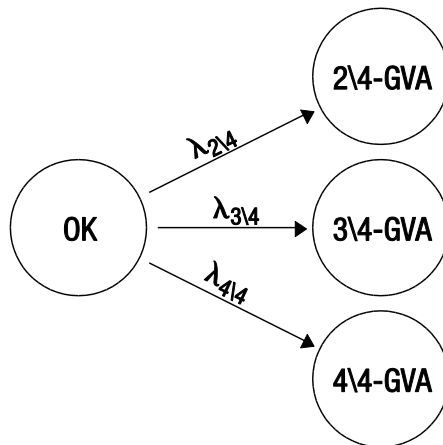


Abb. 5.1 Modellstruktur Modell A

Wie zuvor diskutiert, ist $q_{k \setminus r} \ll 1$ Voraussetzung für die Gültigkeit von Gleichung (5.1), was im Allgemeinen erfüllt wird.

5.2.1.1 Phänomenabhängige Fehlerentdeckungszeiten

Für den Spezialfall, dass mehrere phänomenabhängige Fehlerentdeckungszeiten vorliegen, kann Gleichung (5.1) abgewandelt werden zu

$$q_{k \setminus r} = \sum_{i=1}^I t_{(i)} \lambda_{k \setminus r; (i)} \quad (5.2)$$

wobei I die Anzahl verschiedener Fehlerentdeckungszeiten bezeichnet. Entsprechend bezeichnet $\lambda_{k \setminus r; (i)}$ die Rate, mit der GVA-Phänomene mit Fehlerentdeckungszeit $t_{(i)}$ zu einem (k von r)-Ausfall führen.

In der Praxis kommt $I = 2$ vor beim Frischdampfabblasseregelventil: $t_{(1)} = 336 \text{ h}$ für die monatlichen, versetzten Fahrprüfungen der Abblaseregelventile bei geschlossenen Absperrventilen und $t_{(2)} = 8736 \text{ h}$ für die nicht versetzten jährlichen „scharfen“ Prüfungen der Abblaseregelventile vor dem Anfahren mit geöffneten Absperrventilen.

Da die Verallgemeinerung offensichtlich ist, wird im Folgenden auf die Darstellung der Abhängigkeit von der Fehlerentdeckungszeit verzichtet.

5.2.2 Modell B

Beim Modell B treten GVA mit einer Rate λ auf. Die Anteile der verschiedenen Ausfallkombinationen wird durch Parameter $\omega_{k \setminus r}$ ($k = 2 \dots r$) beschrieben. Die Nicht-Verfügbarkeit aufgrund eines (k von r)-GVA ist, wobei t die Fehlerentdeckungszeit bezeichnet:

$$q_{k \setminus r} = t \lambda \omega_{k \setminus r} \quad (5.3)$$

In Abb. 5.2 ist die Modellstruktur des Modells B am Beispiel einer Komponentengruppe mit vier Komponenten dargestellt. Die Komponentengruppe befindet sich zunächst in einem nicht durch GVA-Phänomene beeinflussten Zustand; alle Komponenten sind verfügbar („ok“). Mit einer Rate λ geht es in einen GVA-Zustand über. Instantan geht es mit Wahrscheinlichkeit $\omega_{k \setminus r}$ in einen Zustand über, bei dem k der vorhandenen r Komponenten unverfügbar sind ($k = 2, \dots, r$).

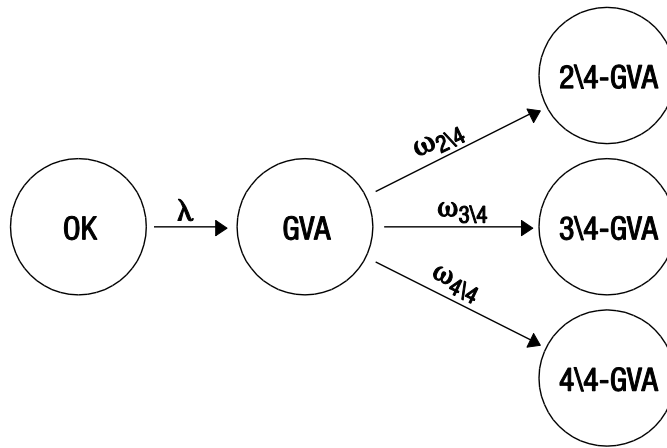


Abb. 5.2 Modellstruktur Modell B

Analog zu der vorhergehenden Diskussion wird $t \lambda \omega_{k \setminus r} \ll 1$ angenommen. Da die Komponentengruppe sicher in einen der möglichen GVA-Zustände übergeht, gilt:

$$\sum_{k=2}^r \omega_{k \setminus r} = 1 \quad (5.4)$$

Es ist anzumerken, dass die Modelle A und B äquivalent sind, da

$$\lambda_{k \setminus r} = \lambda \omega_{k \setminus r} \quad (5.5)$$

gilt. Allerdings legt die verschiedene Modellstruktur eine unterschiedliche Wahl der a-priori-Verteilungen der Modellparameter nahe, die im Allgemeinen zu abweichenden Schätzergebnissen führt. Dadurch kann bewirkt werden, dass die erwartete Gesamt-GVA-Rate im Fall der Nullfehlerstatistik nicht mit dem Redundanzgrad r anwächst. Dies wird unten diskutiert.

Dieses Modell ist in seiner Struktur ähnlich dem Alpha-Faktor-Modell. Es unterscheidet sich von ihm allerdings dadurch, dass nur GVA und keine Einzelausfälle beschrieben werden. Deshalb ist im Gegensatz zum Alpha-Faktor-Modell eine Schätzung aus der in der GVA-Datenbank enthaltenen Information möglich (Bedingung 3).

5.2.3 Modell C

In Hinblick auf günstige Eigenschaften bei der Übertragung von Betriebserfahrung aus Komponentengruppen abweichender Größe ('mapping') wird Modell C eingeführt. Dieses Modell ist ähnlich Modell B, aber es werden nicht nur GVA, sondern beliebige Ereignisse mit systematischer Ursache beschrieben.

Die Nicht-Verfügbarkeit aufgrund eines (k von r)-GVA ist, wobei t die Fehlerentdeckungszeit bezeichnet:

$$q_{k \setminus r} = t \zeta x_{k \setminus r} \quad (5.6)$$

In Abb. 5.3 ist die Modellstruktur des Modells C dargestellt. Die Komponentengruppe befindet sich zunächst in einem nicht durch GVA-Phänomene beeinflussten Zustand; alle Komponenten sind verfügbar („ok“). Mit einer Rate ζ geht es in einen Zustand über, bei dem ein GVA-Phänomen auf die Komponentengruppe eingewirkt hat. Instantan geht es mit Wahrscheinlichkeit $x_{k \setminus r}$ in einen Zustand über, bei dem k der vorhandenen r Komponenten unverfügbar sind ($k = 0, \dots, r$). Das heißt, es werden auch Fälle beschrieben, in denen keine Komponenten ausgefallen sind. Dies steht damit in Beziehung, dass die Interpretationsvektoren (siehe Abschnitt 2.2.1) für Ereignisse, in denen nur Komponentenschädigungen, keine Ausfälle beobachtet wurden, Elemente $w_{0 \setminus r} > 0$ aufweisen. Allerdings zeigt die Betriebserfahrung, dass bei Beobachtung nur einer Komponentenschädigung nicht immer eine systematische Ursache erkannt wird. Deshalb können für die Quantifizierung dieses Modells benötigten Anzahlen der Ereignisse mit nur einem oder keinem Komponentenausfall nicht genau bestimmt werden; der Schätzfehler ist im Allgemeinen von unbekannter Größe.

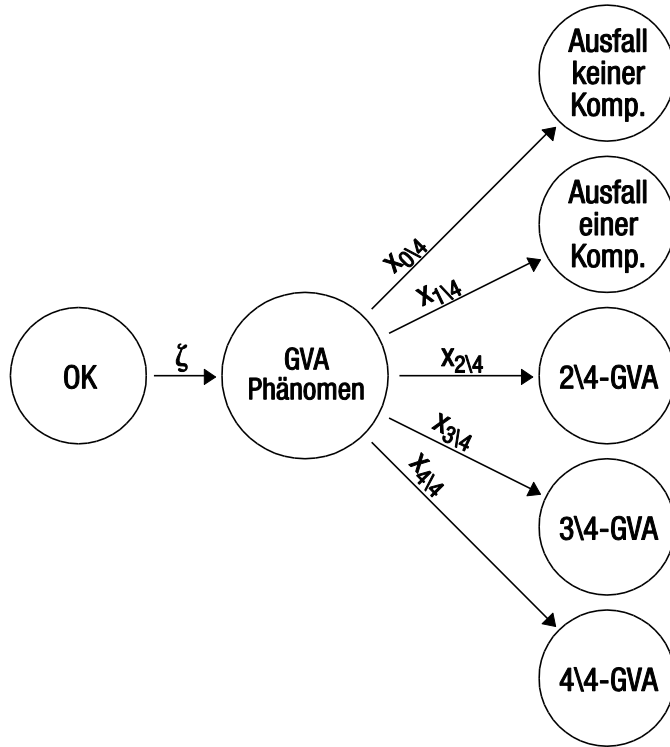


Abb. 5.3 Modellstruktur Modell C

Die Einführung von $x_{0\setminus r}$ ermöglicht es, die Rate ζ unabhängig von den Interpretationsvektoren zu schätzen. Dies führt im Ergebnis dazu, dass die Unsicherheitsverteilungen von ζ und $X = \{x_{0\setminus r}, x_{1\setminus r} \dots x_{r\setminus r}\}$ faktorisieren.

Analog zu der vorhergehenden Diskussion wird $t \zeta \omega_{k\setminus r} \ll 1$ angenommen. Da die Komponentengruppe sicher in einen der möglichen Endzustände übergeht, gilt:

$$\sum_{k=0}^r x_{k\setminus r} = 1 \quad (5.7)$$

Modell C wird hier insbesondere für weitere theoretische Betrachtungen eingeführt, da es in Verbindung mit dem Mapping eine einfachere Analyse erlaubt (siehe Abschnitt 6.2.3). Der Grund ist folgender: Wenn man z. B. eine Komponentengruppe der Größe 2 als eine zufällige Teilmenge der Komponenten einer Gruppe 4 ansieht, und ein (2 von 4)-GVA aufgetreten ist, so sind in der Zweiergruppe mit Wahrscheinlichkeit 1/4 zwei Komponenten, mit Wahrscheinlichkeit 1/2 eine Komponente und mit Wahrscheinlichkeit 1/4 keine Komponenten ausgefallen. Nur bei Modell C werden (0 von 2)-Ereignisse und (1 von 2)-Ereignisse vom Modell erfasst. Über diese kombinatorischen Überlegun-

gen können Beziehungen zwischen den $x_{k \setminus r}$ für verschiedene Redundanzgrade hergeleitet werden, während λ für alle r identisch ist (siehe Abschnitt 6.2.2). Nur Modell C weist diese Eigenschaft auf.

Im Folgenden wird die Schätzung der Parameter und GVA-Wahrscheinlichkeiten aus der Betriebserfahrung dargestellt.

5.3 Schätzung der Modellparameter aus Ereignisanzahlen

Zur Schätzung der Modellparameter und GVA-Wahrscheinlichkeiten wird zuerst davon ausgegangen, dass die Anzahl der verschiedenen Ausfallkombinationen bekannt ist. Diese wird mit $\mathfrak{N} = \{n_{0 \setminus r}, n_{1 \setminus r} \dots n_{r \setminus r}\}$ bezeichnet, während die Gesamtbeobachtungszeit als T bezeichnet wird. Die Gesamtheit der geschätzten GVA-Wahrscheinlichkeiten wird mit $\mathfrak{Q}$ bezeichnet.

5.3.1 Modell A

Da beim Modell A die GVA-Ereignisse mit verschiedenen Ausfallkombinationen als unabhängige Elementarereignisse angesehen werden, können die Raten $\lambda_{k \setminus r}$ der einzelnen GVA-Ausfallkombinationen $k = 2, \dots, r$ unabhängig voneinander geschätzt werden. Die Schätzung der Rate $\lambda_{k \setminus r}$ hängt nur von der Anzahl $n_{k \setminus r}$ der (k von r)-GVA und der Gesamtbeobachtungszeit als T ab.

Wählt man eine nichtinformative a priori-Verteilung über die Jeffreys'sche Regel

$$\pi(\lambda_{k \setminus r}) \propto \frac{1}{\sqrt{\lambda_{k \setminus r}}} \quad (5.8)$$

so folgt für die a posteriori-Verteilung:

$$p(\lambda_{k \setminus r} | \mathfrak{N}) = p(\lambda_{k \setminus r} | n_{k \setminus r}) = \frac{T^{n_{k \setminus r} + 1/2}}{\Gamma(n_{k \setminus r} + 1/2)} (\lambda_{k \setminus r})^{n_{k \setminus r} - 1/2} e^{-\lambda_{k \setminus r} T} \quad (5.9)$$

$\lambda_{k \setminus r}$ genügt einer Gammaverteilung mit Parametern $n_{k \setminus r} + 1/2$ und T . Für die Verbundwahrscheinlichkeit der Modellparameter gilt somit:

$$p(\mathcal{M}|\mathfrak{N}) = \prod_{k=2}^r \frac{T^{n_{k\setminus r}+1/2}}{\Gamma(n_{k\setminus r} + 1/2)} (\lambda_{k\setminus r})^{n_{k\setminus r}-1/2} e^{-\lambda_{k\setminus r}T} \quad (5.10)$$

Da nach Gleichung (6.1) $q_{k\setminus r} = t \lambda_{k\setminus r}$ gilt, folgt:

$$p(q_{k\setminus r}|n_{k\setminus r}) = \frac{\left(\frac{T}{t}\right)^{n_{k\setminus r}+1/2}}{\Gamma(n_{k\setminus r} + 1/2)} (q_{k\setminus r})^{n_{k\setminus r}-1/2} e^{-q_{k\setminus r}T/t} \quad (5.11)$$

$q_{k\setminus r}$ genügt einer Gammaverteilung mit Parametern $n_{k\setminus r} + 1/2$ und T/t . Die Verteilung der Gesamtheit der geschätzten GVA-Wahrscheinlichkeiten $\mathfrak{Q} = \{q_{2\setminus r}, q_{3\setminus r} \dots q_{r\setminus r}\}$ faktorisiert:

$$p(\mathfrak{Q}|\mathfrak{N}) = \prod_{k=2}^r \frac{\left(\frac{T}{t}\right)^{n_{k\setminus r}+1/2}}{\Gamma(n_{k\setminus r} + 1/2)} (q_{k\setminus r})^{n_{k\setminus r}-1/2} e^{-q_{k\setminus r}T/t} \quad (5.12)$$

Für die Erwartungswerte gilt:

$$\langle q_{k\setminus r} \rangle = \left(n_{k\setminus r} + \frac{1}{2}\right) \frac{t}{T} \quad (5.13)$$

Im Fall einer Nullfehlerstatistik gilt

$$\langle q_{k\setminus r} \rangle = \frac{t}{2T} \quad (5.14)$$

Daraus folgt, dass für die erwartete Gesamtwahrscheinlichkeit für GVA $\langle \sum_{k=2}^r q_{k\setminus r} \rangle$ im Fall einer Nullfehlerstatistik

$$\left\langle \sum_{k=2}^r q_{k\setminus r} \right\rangle = (r-1) \frac{t}{2T} \quad (5.15)$$

mit dem Redundanzgrad stark ansteigt (linear für große r). Dieser Erwartungswert entspricht dem Ergebnis bei Anwendung des Maximum-Likelihood-Punktschätzers $n_{k\setminus r}/T$ auf einen Datensatz mit $(r-1)$ GVA-Ereignissen. Für den Redundanzgrad 5 entspricht dies zwei GVA-Ereignissen, für $r=128$ über 63 GVA-Ereignissen.

Modell B ist so konstruiert, dass diese Eigenschaft vermieden wird, wie im Folgenden gezeigt wird.

5.3.2 Modell B

Beim Modell B beschreibt die GVA-Rate λ den Übergang vom Zustand „ok“ zum Zustand „GVA“. Deshalb kann λ aus der Gesamtzahl der beobachteten GVA geschätzt werden, die definiert werden kann als

$$n_{GVA} = \sum_{k=2}^r n_{k \setminus r} \quad (5.16)$$

Wählt man eine nichtinformative a priori-Verteilung über die Jeffreys'sche Regel

$$\pi(\lambda) \propto \frac{1}{\sqrt{\lambda}} \quad (5.17)$$

so folgt für die a posteriori-Verteilung:

$$p(\lambda | \mathfrak{N}) = p(\lambda | n_{GVA}) = \frac{T^{n_{GVA} + 1/2}}{\Gamma(n_{GVA} + 1/2)} (\lambda)^{n_{GVA} - 1/2} e^{-\lambda T} \quad (5.18)$$

Für den Erwartungswert gilt:

$$\langle \lambda \rangle = \frac{n_{GVA} + \frac{1}{2}}{T} \quad (5.19)$$

Die Schätzung der Verteilung der Ausfallkombinationen erfolgt analog Abschnitt 4.2. $\omega_{k \setminus r}$ stellen die Parameter einer Multinomialverteilung dar. Die konjugierte Verteilung einer Multinomialverteilung ist wie oben erwähnt die Dirichletverteilung. Für die Dichte der Dirichletverteilung gilt:

$$p(\omega_{2 \setminus r}, \omega_{3 \setminus r}, \dots, \omega_{r-1 \setminus r} | a_{2 \setminus r}, a_{3 \setminus r}, \dots, a_{r \setminus r}) = \frac{1}{B(a_{2 \setminus r}, a_{3 \setminus r}, \dots, a_{r \setminus r})} \prod_{i=2}^r (\omega_{i \setminus r})^{a_{i \setminus r} - 1} \quad (5.20)$$

Hierbei sind als erstes Argument der Dichte p nur $\omega_{2 \setminus r}$ bis $\omega_{r-1 \setminus r}$ aufgeführt, da $\omega_{r \setminus r}$ nach Gleichung (6.4) durch diese Größen bestimmt wird. Der Normierungsfaktor $B(a_{2 \setminus r}, a_{3 \setminus r}, \dots, a_{r \setminus r})$ ist gegeben durch

$$B(a_{0 \setminus r}, a_{1 \setminus r}, \dots, a_{r \setminus r}) = \frac{\prod_{i=2}^r \Gamma(a_{i \setminus r})}{\Gamma(\sum_{i=2}^r a_{i \setminus r})} \quad (5.21)$$

wobei Γ die Gammafunktion bezeichnet.

Wählt man die a priori-Verteilung nach dem Verfahren von Jeffreys, so erhält man als a priori-Verteilung eine Dirichletverteilung mit Parametern $\alpha_i = 1/2$:

$$\pi(\omega_{2\setminus r}, \omega_{3\setminus r}, \dots, \omega_{r\setminus r}) \propto \prod_{i=2}^r (\omega_{i\setminus r})^{-\frac{1}{2}} \quad (5.22)$$

Wenn $n_{k\setminus r}$ (k von r)-Ausfälle beobachtet wurden, ist die a posteriori-Verteilung eine Dirichletverteilung mit den Parametern $a_k = n_{k\setminus r} + 1/2$:

$$p(\omega_{2\setminus r}, \omega_{3\setminus r}, \dots, \omega_{r\setminus r} | \mathfrak{N}) = \frac{\prod_{i=2}^r (\omega_{i\setminus r})^{n_{i\setminus r} - 1/2}}{B(n_{2\setminus r} + 1/2, n_{3\setminus r} + 1/2, \dots, n_{r\setminus r} + 1/2)} \quad (5.23)$$

Für die Erwartungswerte gilt:

$$\langle \omega_{k\setminus r} \rangle = \frac{n_{k\setminus r} + 1/2}{\sum_{i=2}^r (n_{i\setminus r} + 1/2)} = \frac{2n_{k\setminus r} + 1}{2n_{GVA} + r - 1} \quad (5.24)$$

Für die Verbundwahrscheinlichkeit der Modellparameter gilt:

$$p(\mathcal{M} | \mathfrak{N}) = \frac{T^{n_{GVA} + 1/2}}{\Gamma(n_{GVA} + 1/2)} (\lambda)^{n_{GVA} - 1/2} e^{-\lambda T} \times \frac{\prod_{i=2}^r (\omega_{i\setminus r})^{n_{i\setminus r} - 1/2}}{B(n_{2\setminus r} + 1/2, n_{3\setminus r} + 1/2, \dots, n_{r\setminus r} + 1/2)} \quad (5.25)$$

Für Modell B lässt sich kein einfacher geschlossener Ausdruck für die Verteilung der GVA-Wahrscheinlichkeiten angeben, da $q_{k\setminus r}$ ein Produkt aus $\omega_{k\setminus r}$ und λ ist (Gleichung (6.3)). Daraus folgt auch, dass die Verteilung $p(\mathfrak{Q} | \mathfrak{N})$ für Modell B im Gegensatz zu Modell A nicht faktorisiert. Auch die Marginalverteilungen $p(q_{k\setminus r} | \mathfrak{N})$ lassen sich nicht analytisch angeben. Die Verteilung $p(\mathfrak{Q} | \mathfrak{N})$ und ihre Marginalverteilungen können aber im Rahmen eines Monte-Carlo-Verfahrens einfach und effizient realisiert werden, indem Samples (Stichproben) aus (6.12) und (6.17) gezogen werden und Samples von $q_{k\setminus r}$ nach Gleichung (6.3) bestimmt werden.

Da λ und $\omega_{k \setminus r}$ unabhängig sind, gegeben $\mathfrak{N}$, gilt für den Erwartungswert der GVA-Unverfügbarkeiten

$$\langle q_{k \setminus r} \rangle = t \langle \lambda \rangle \langle \omega_{k \setminus r} \rangle = t \frac{n_{GVA} + \frac{1}{2}}{T} \frac{2n_{k \setminus r} + 1}{2n_{GVA} + r - 1} \quad (5.26)$$

Sind keine Ereignisse beobachtet worden, folgt:

$$\langle q_{k \setminus r} \rangle = t \langle \lambda \rangle \langle \omega_{k \setminus r} \rangle = \frac{t}{2T} \frac{1}{r - 1} \quad (5.27)$$

Daraus folgt, dass für die erwartete Gesamtwahrscheinlichkeit für GVA $\langle \sum_{k=2}^r q_{k \setminus r} \rangle$ im Fall einer Nullfehlerstatistik

$$\langle \sum_{k=2}^r q_{k \setminus r} \rangle = \frac{t}{2T} \quad (5.28)$$

im Gegensatz zu Modell A nicht vom Redundanzgrad abhängt.

5.3.3 Modell C

Beim Modell C beschreibt die GVA-Rate ζ den Übergang vom Zustand „ok“ zum Zustand „GVA-Phänomen hat eingewirkt“. Deshalb kann ζ aus der Gesamtzahl der Ereignisse mit Einwirkung eines GVA-Phänomens geschätzt werden. Diese wird definiert als

$$n_{ph} = \sum_{k=0}^r n_{k \setminus r} = N \quad (5.29)$$

Wählt man eine nichtinformative a priori-Verteilung über die Jeffreys'sche Regel

$$\pi(\zeta) \propto \frac{1}{\sqrt{\zeta}} \quad (5.30)$$

so folgt für die a posteriori-Verteilung:

$$p(\zeta|\mathfrak{N}) = p(\zeta|n_{ph}) = \frac{T^{n_{ph}+1/2}}{\Gamma(n_{ph}+1/2)} (\zeta)^{n_{ph}-1/2} e^{-\zeta T} \quad (5.31)$$

Im Gegensatz zu (6.12) ist die Schätzung auch von $n_{0\setminus r}$ und $n_{1\setminus r}$ abhängig.

Für den Erwartungswert gilt:

$$\langle \zeta \rangle = \frac{n_{ph} + \frac{1}{2}}{T} \quad (5.32)$$

Die Schätzung der Verteilung der Ausfallkombinationen erfolgt ebenfalls analog dem vorigen Abschnitt. Wählt man die a priori-Verteilung wieder nach dem Verfahren von Jeffreys, so erhält man eine Dirichlet-Verteilung mit den Parametern $\alpha_i = 1/2$:

$$(x_{0\setminus r}, x_{1\setminus r}, \dots, x_{r\setminus r}) \propto \prod_{i=0}^r (x_{i\setminus r})^{-\frac{1}{2}} \quad (5.33)$$

Wenn $n_{k\setminus r}$ (k von r)-Ausfälle beobachtet wurden, ist die a posteriori-Verteilung eine Dirichlet-Verteilung mit den Parametern $a_k = n_{k\setminus r} + 1/2$:

$$p(x_{0\setminus r}, x_{1\setminus r}, \dots, x_{r\setminus r}|\mathfrak{N}) = \frac{\prod_{i=0}^r (x_{i\setminus r})^{n_{i\setminus r}-1/2}}{B(n_{0\setminus r} + 1/2, n_{1\setminus r} + 1/2, \dots, n_{r\setminus r} + 1/2)} \quad (5.34)$$

In Übereinstimmung mit (6.22) und abweichend von (6.17) sind die Schätzungen auch von $n_{0\setminus r}$ und $n_{1\setminus r}$ abhängig.

Für die Verbundwahrscheinlichkeit der Modellparameter gilt:

$$p(\mathcal{M}|\mathfrak{N}) = \frac{T^{n_{ph}+1/2}}{\Gamma(n_{ph}+1/2)} (\zeta)^{n_{ph}-1/2} e^{-\zeta T} \times \frac{\prod_{i=0}^r (x_{i\setminus r})^{n_{i\setminus r}-1/2}}{B(n_{0\setminus r} + 1/2, n_{1\setminus r} + 1/2, \dots, n_{r\setminus r} + 1/2)} \quad (5.35)$$

Für Modell C lässt sich wie bei Modell B kein einfacher geschlossener Ausdruck für die Verteilung der GVA-Wahrscheinlichkeiten angeben, da $q_{k\setminus r}$ ein Produkt aus $x_{k\setminus r}$ und ζ ist (Gleichung (6.6)). Daraus folgt auch, dass die Verteilung $p(\mathfrak{Q}|\mathfrak{N})$ für Modell C nicht faktorisiert. Auch die Marginalverteilungen $p(q_{k\setminus r}|\mathfrak{N})$ lassen sich nicht analytisch angeben. Die Verteilung $p(\mathfrak{Q}|\mathfrak{N})$ und ihre Marginalverteilungen können aber wie bei Modell B im Rahmen eines Monte-Carlo-Verfahrens einfach und effizient realisiert werden, indem Samples aus (6.22) und (6.25) gezogen werden und Samples von $q_{k\setminus r}$ nach Gleichung (6.6) bestimmt werden.

Für die Erwartungswerte gilt:

$$\langle x_{k\setminus r} \rangle = \frac{n_{k\setminus r} + 1/2}{\sum_{i=0}^r (n_{i\setminus r} + 1/2)} = \frac{2n_{k\setminus r} + 1}{2n_{ph} + r + 1} \quad (5.36)$$

Da ζ und $x_{k\setminus r}$ bei gegebenem $\mathfrak{N}$ unabhängig sind, gilt für den Erwartungswert der GVA-Unverfügbarkeiten

$$\langle q_{k\setminus r} \rangle = t \langle \zeta \rangle \langle x_{k\setminus r} \rangle = t \frac{n_{ph} + \frac{1}{2}}{T} \frac{2n_{k\setminus r} + 1}{2n_{ph} + r + 1} \quad (5.37)$$

Sind keine Ereignisse beobachtet worden, so folgt:

$$\langle q_{k\setminus r} \rangle = t \langle \lambda \rangle \langle \omega_{k\setminus r} \rangle = \frac{t}{2T} \frac{1}{r + 1} \quad (5.38)$$

Daraus folgt, dass für die erwartete Gesamtwahrscheinlichkeit für GVA $\langle \sum_{k=2}^r q_{k\setminus r} \rangle$ im Fall einer Nullfehlerstatistik

$$\langle \sum_{k=2}^r q_{k\setminus r} \rangle = \frac{t}{2T} \frac{r - 1}{r + 1} < \frac{t}{2T} \quad (5.39)$$

mit dem Redundanzgrad ansteigt aber asymptotisch für große r von ihm unabhängig wird.

5.3.4 Konvergenzeigenschaften der Modelle

Aufgrund ihrer Konstruktion ist gesichert, dass die Parameter der Modelle gegen ihre wahren Werte konvergieren. Wie diese Konvergenz mit zunehmender Betriebserfahrung erfolgt, wird im Folgenden diskutiert.

Zunächst wird der Fall betrachtet, dass in einer endlichen Beobachtungszeit T keine Ereignisse auftreten (so genannte Nullfehlerstatistik). Dies ist in Tab. 5.1 dargestellt. t bezeichnet die Fehlerentdeckungszeit.

Tab. 5.1 Erwartete GVA-Wahrscheinlichkeit bei Nullfehlerstatistik

Modell	Einzelne Ausfallkombination	alle GVA-Ausfallkombinationen
A	$\langle q_{k \setminus r} \rangle = \frac{t}{2T}$	$\langle \sum_{k=2}^r q_{k \setminus r} \rangle = (r-1) \frac{t}{2T}$
B	$\langle q_{k \setminus r} \rangle = \frac{t}{2(r-1)T}$	$\langle \sum_{k=2}^r q_{k \setminus r} \rangle = \frac{t}{2T}$
C	$\langle q_{k \setminus r} \rangle = \frac{t}{2(r+1)T}$	$\langle \sum_{k=2}^r q_{k \setminus r} \rangle = \frac{r-1}{r+1} \frac{t}{2T}$

Somit erfolgt die Konvergenz proportional zum Kehrwert der Beobachtungszeit. Die Konvergenzgeschwindigkeit wird durch den Vorfaktor bestimmt, der bei den Modellen verschieden ist.

Die erwarteten GVA-Wahrscheinlichkeiten sind bei Modell B und Modell C um einen Faktor $(r-1)$ bzw. $(r+1)$ kleiner als bei Modell A. Deshalb kann – insbesondere bei großen Komponentengruppen – die Wahl des Modells einen erheblichen Einfluss auf das Ergebnis haben. Bei Modell B ist die erwartete Gesamtwahrscheinlichkeit von GVA unabhängig vom Redundanzgrad, Bei Modell A steigt sie mit r stark an (proportional r für große r), bei Modell C weist sie nur eine schwache Abhängigkeit von r auf. Diese unterschiedliche Abhängigkeit vom Redundanzgrad gilt nicht nur bei Nullfehlerstatistik, sondern auch bei einer geringen Anzahl von beobachteten Ereignissen, wie im Folgenden illustriert wird.

Hierbei wird folgendes Szenario betrachtet: In einer Komponentengruppe der Größe r treten (i von r)-Ereignisse mit einer Rate ι auf ($i \geq 2$); andere Ereignisse werden nicht beobachtet. Z. B. kann $i = r$ gewählt werden; dann treten nur komplette GVA auf.

Für die erwarteten Werte $\langle q_{k \setminus r} \rangle$ gilt bei Modell A:

$$\langle q_{i \setminus r} \rangle = \left(\langle n_{i \setminus r} \rangle + \frac{1}{2} \right) \frac{t}{T} = t \frac{\iota T + \frac{1}{2}}{T} = \iota t + \frac{t}{2T} \quad (5.40)$$

und für $k \neq i$

$$\langle q_{k \setminus r} \rangle = \frac{t}{2T} \quad (5.41)$$

Für Modell B gilt, da die Anzahl der Ereignisse poissonverteilt ist:

$$\langle q_{i \setminus r} \rangle = t \sum_{n_{i \setminus r}=0}^{\infty} \frac{n_{i \setminus r} + \frac{1}{2}}{T} \frac{2n_{i \setminus r} + 1}{2n_{i \setminus r} + r - 1} p_{POISSON}(n_{i \setminus r} | \iota T) \quad (5.42)$$

mit $p_{POISSON}(n|l) = l^n e^{-l} / n!$ der Wahrscheinlichkeitsfunktion der Poissonverteilung. (6.32) folgt aus (6.13) und (6.18), da die Parameter λ und $\omega_{i \setminus r}$ unabhängig sind, gegeben $n_{i \setminus r}$.

Für $k \neq i$ gilt

$$\langle q_{k \setminus r} \rangle = t \sum_{n_{i \setminus r}=0}^{\infty} \frac{n_{i \setminus r} + \frac{1}{2}}{T} \frac{1}{2n_{i \setminus r} + r - 1} p_{POISSON}(n_{i \setminus r} | \iota T) \quad (5.43)$$

Für Modell C gilt analog:

$$\langle q_{i \setminus r} \rangle = t \sum_{n_{i \setminus r}=0}^{\infty} \frac{n_{i \setminus r} + \frac{1}{2}}{T} \frac{2n_{i \setminus r} + 1}{2n_{i \setminus r} + r + 1} p_{POISSON}(n_{i \setminus r} | \iota T) \quad (5.44)$$

und für $k \neq i$

$$\langle q_{k \setminus r} \rangle = t \sum_{n_{i \setminus r}=0}^{\infty} \frac{n_{i \setminus r} + \frac{1}{2}}{T} \frac{1}{2n_{i \setminus r} + r + 1} p_{POISSON}(n_{i \setminus r} | \iota T) \quad (5.45)$$

Es ist anzumerken, dass hier der Erwartungswert bezüglich der Ereignisanzahl und der GVA-Wahrscheinlichkeiten gebildet wird. Die Summationen in den Gleichungen (5.31) bis (5.34) lassen sich analytisch ausführen, wobei relativ komplizierte Ausdrücke mit Summen verallgemeinerter unvollständiger Gammafunktionen erhalten werden. Auf eine Darstellung wurde hier verzichtet, da diese Ausdrücke einer unmittelbaren Interpretation nicht zugänglich sind.

Zuerst wird das Verhalten der Schätzer für eine Komponentengruppe der Größe $r = 4$ dargestellt. Als Beispiel wird hier $i = 4$ gewählt, d. h. es treten nur komplette GVA auf. In Abb. 5.4 wird das Verhalten von $\langle q_{4 \setminus 4} \rangle$ und in Abb. 5.5 das Verhalten von $\langle q_{k \setminus 4} \rangle$ für $k \neq 4$ dargestellt. Es gilt $\langle q_{3 \setminus 4} \rangle = \langle q_{2 \setminus 4} \rangle$, für Modell C auch $\langle q_{3 \setminus 4} \rangle = \langle q_{2 \setminus 4} \rangle = \langle q_{1 \setminus 4} \rangle = \langle q_{0 \setminus 4} \rangle$. Dabei ist die Skalierung der Zeitachse mit $\iota = 1$ so gewählt, dass sie der erwarteten Anzahl von Ereignissen entspricht, d. h. bei z. B. $T = 10$ sind im Mittel 10 Ereignisse aufgetreten. Es gilt $t = 0.0001$, d. h. die wahre Wahrscheinlichkeit eines (2 von 4)-Ereignisses in jedem Testintervall ist 0.0001.

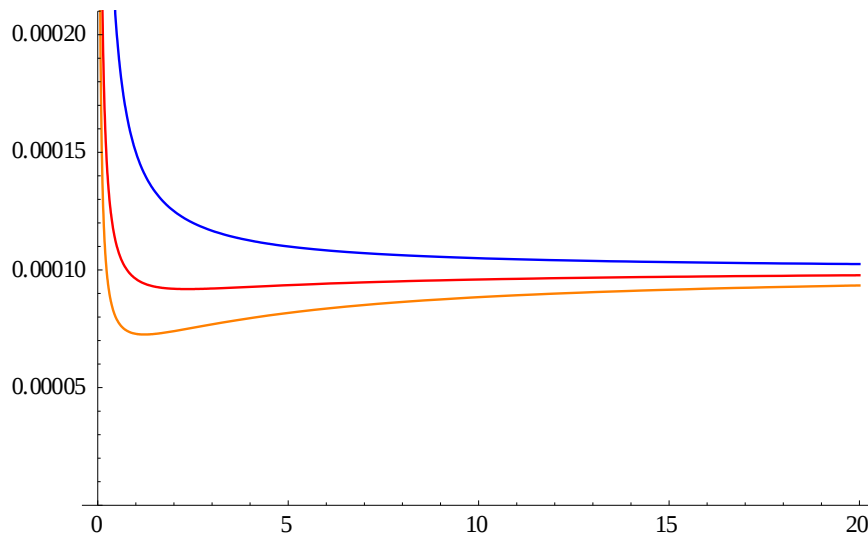


Abb. 5.4 Erwartungswerte $\langle q_{4 \setminus 4} \rangle$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange)

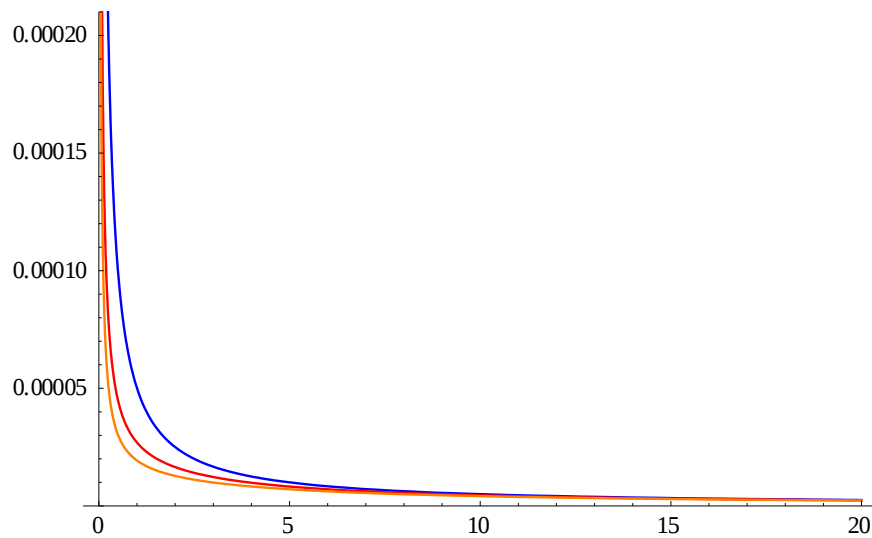


Abb. 5.5 Erwartungswerte $\langle q_{k \setminus 4} \rangle$ für $k \neq 4$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange)

Es ist erkennbar, dass für Modell A die Erwartungswerte des Schätzers stets über dem wahren Wert liegen, während Modelle B und C bzgl. $\langle q_{4 \setminus 4} \rangle$ einen „Unterschwinger“ zeigen. Für $T \geq 1$ liegen die Erwartungswerte der Schätzer stets unter dem wahren Wert. Dies ist bei Modell C stärker ausgeprägt als bei Modell B. Der minimale Wert bei Modell B ist über 8 %, der von Modell C über 27 % kleiner als der wahre Wert.

Für $k \neq 4$ liegen auch bei Modell B und C die Erwartungswerte des Schätzers $\langle q_{k \setminus 4} \rangle$ stets über dem wahren Wert.

Der „Unterschwinger“ ist umso stärker ausgeprägt, je größer r ist. Als Beispiel wird folgendes dem oben betrachteten Beispiel analoges Szenario in einer Komponentengruppe der Größe 8 betrachtet: Es treten (8 von 8)-Ereignisse mit einer Rate ι auf; andere Ereignisse werden nicht beobachtet.

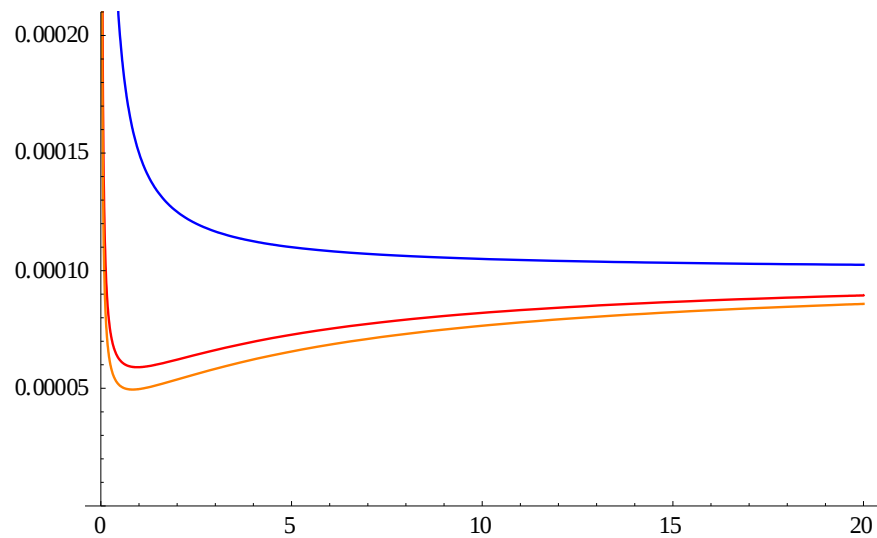


Abb. 5.6 Erwartungswerte $\langle q_{8 \setminus 8} \rangle$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange)

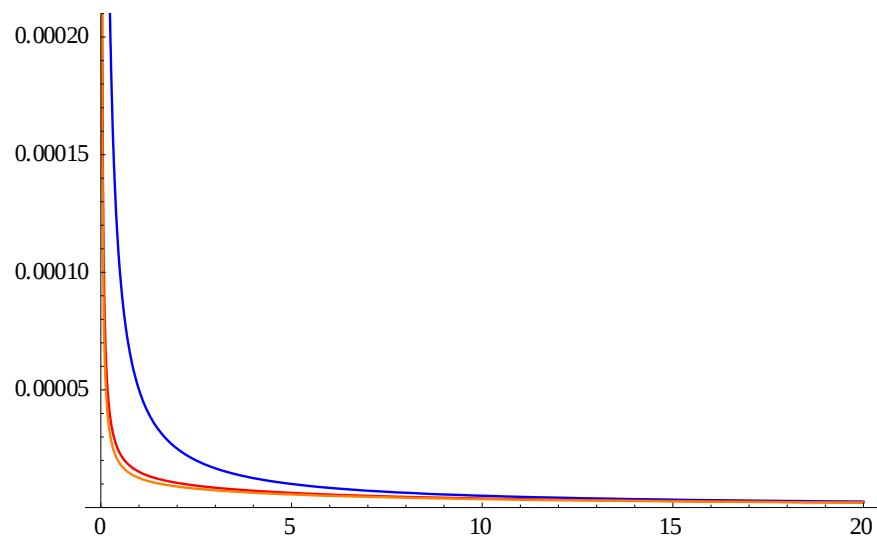


Abb. 5.7 Erwartungswerte $\langle q_{k \setminus 8} \rangle$ für $k \neq 8$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange)

Modell A zeigt, da die Raten der GVA-Ausfallkombinationen unabhängig geschätzt werden, ein identisches Verhalten zu oben. Bei Modell B und C konvergiert $\langle q_{8 \setminus 8} \rangle$ nur langsam. Der minimale Wert bei Modell B ist über 40 %, der von Modell C über 50 % kleiner als der wahre Wert.

Abschließend wird eine sehr große Komponentengruppe mit $r = 128$ betrachtet: Es treten wieder komplette GVA mit einer Rate ι auf; andere Ereignisse werden nicht beobachtet.

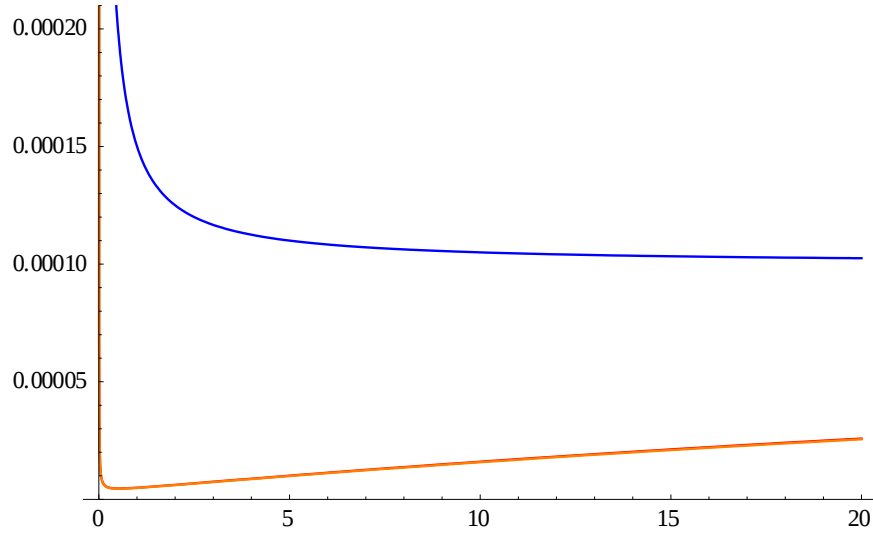


Abb. 5.8 Erwartungswerte $\langle q_{128 \setminus 128} \rangle$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange); die Kurven für Modell B und Modell C liegen fast übereinander

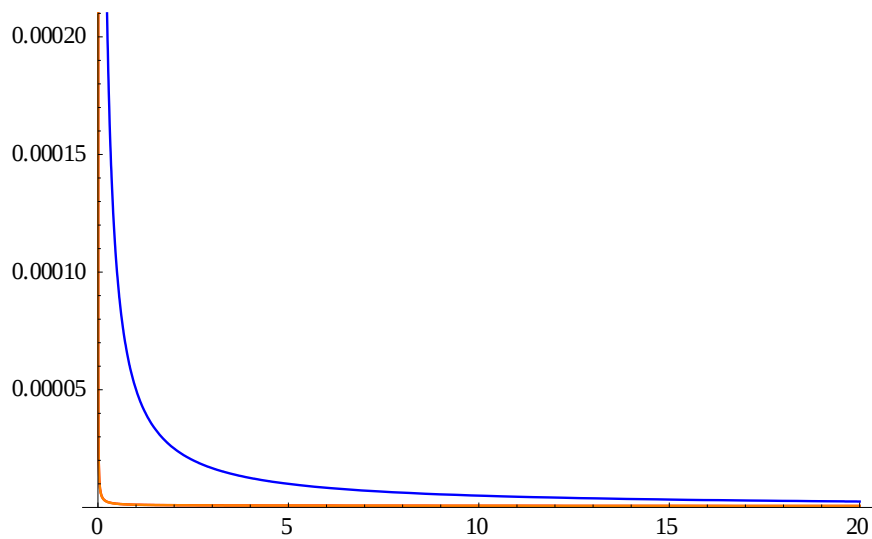


Abb. 5.9 Erwartungswerte $\langle q_{k \setminus 128} \rangle$ für $k \neq 128$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange); die Kurven für Modell B und Modell C liegen fast übereinander

Bei Modell B und C konvergiert $\langle q_{128 \setminus 128} \rangle$ extrem langsam. Für in der Betriebserfahrung typischerweise auftretende Ereigniszahlen ist keine auch nur annähernde Konvergenz zu beobachten. Vielmehr fällt der erwartete Schätzwert auf weniger als 5 % des wahren Wertes ab.

Dieses Verhalten kann verstanden werden, wenn man die Konvergenz der einzelnen Modellparameter betrachtet. Hierzu werden die Erwartungswerte der Schätzungen der Modellparameter λ und $\omega_{k \setminus r}$ bzw. ζ und $x_{k \setminus r}$ untersucht. Für das hier betrachtete Beispiel gilt:

$$\langle \lambda \rangle = \langle \zeta \rangle = \langle q_{i \setminus r} \rangle = \iota t + \frac{t}{2T} \quad (5.46)$$

Analog zu Gleichung (6.32) gilt:

$$\langle \omega_{i \setminus r} \rangle = \sum_{n_{i \setminus r}=0}^{\infty} \frac{2n_{i \setminus r} + 1}{2n_{i \setminus r} + r - 1} p_{POISSON}(n_{i \setminus r} | \iota T) \quad (5.47)$$

bzw.

$$\langle x_{i \setminus r} \rangle = \sum_{n_{i \setminus r}=0}^{\infty} \frac{2n_{i \setminus r} + 1}{2n_{i \setminus r} + r + 1} p_{POISSON}(n_{i \setminus r} | \iota T) \quad (5.48)$$

Für $k \neq i$ gilt

$$\langle \omega_{k \setminus r} \rangle = \sum_{n_{i \setminus r}=0}^{\infty} \frac{1}{2n_{i \setminus r} + r - 1} p_{POISSON}(n_{i \setminus r} | \iota T) \quad (5.49)$$

bzw.

$$\langle x_{k \setminus r} \rangle = \sum_{n_{i \setminus r}=0}^{\infty} \frac{1}{2n_{i \setminus r} + r + 1} p_{POISSON}(n_{i \setminus r} | \iota T) \quad (5.50)$$

Für $r = 128$ sind diese Abhängigkeiten dargestellt. Man beachte den zu den vorherigen Abbildungen abweichenden Maßstab der Abszisse.

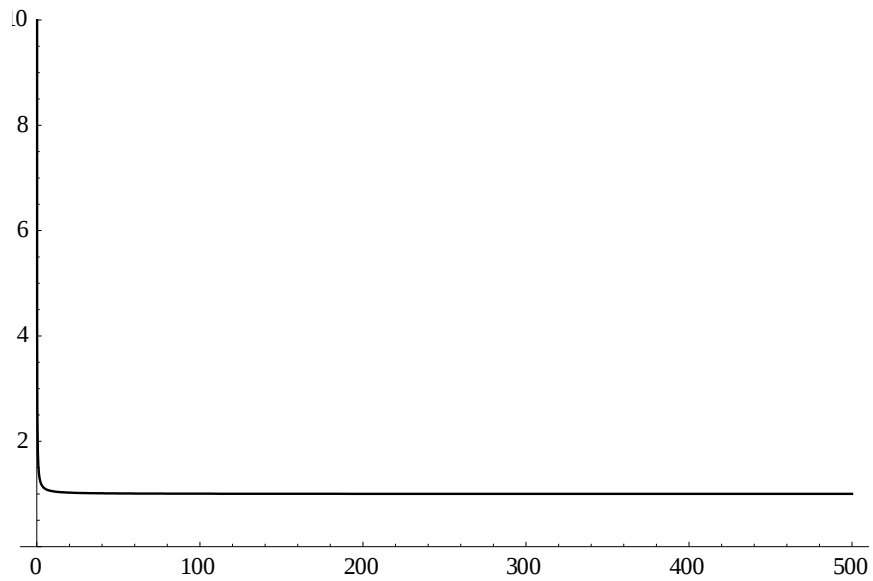


Abb. 5.10 Erwartungswerte $\langle \lambda \rangle = \langle \zeta \rangle$ in Abhängigkeit von T für $\iota = 1$

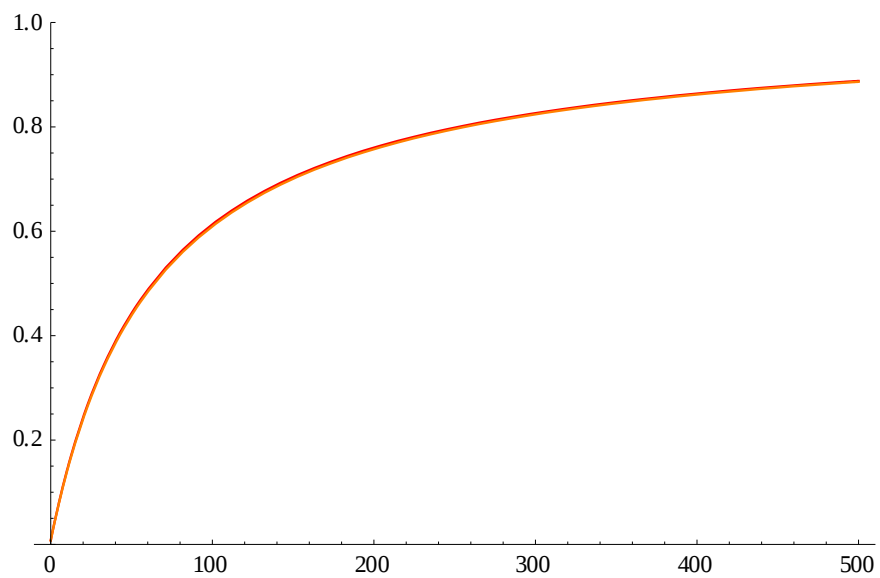


Abb. 5.11 Erwartungswerte $\langle \omega_{128 \setminus 128} \rangle$ bzw. $\langle x_{128 \setminus 128} \rangle$ in Abhängigkeit von T für $\iota = 1$ für Modell B (rot) und Modell C (orange); die Kurven liegen fast übereinander

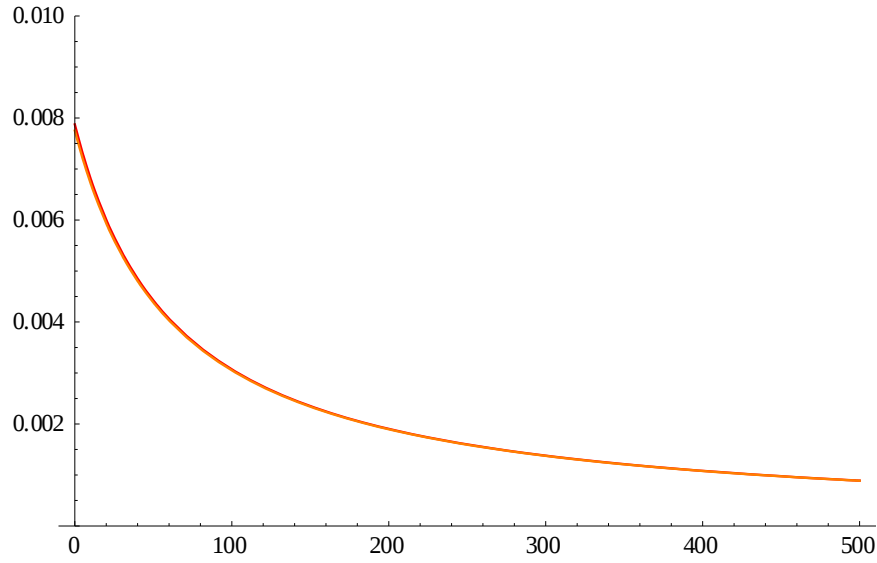


Abb. 5.12 Erwartungswerte $\langle \omega_{k \setminus 128} \rangle$ bzw. $\langle x_{k \setminus 128} \rangle$ für $k \neq 128$ in Abhängigkeit von T für $\iota = 1$ für Modell B (rot) und Modell C (orange); die Kurven liegen fast übereinander

Man erkennt, dass die Konvergenz der Erwartungswerte der Parameter, die die Gesamtrate der Modellereignisse beschreiben ($\langle \lambda \rangle$ bzw. $\langle \zeta \rangle$) sehr schnell, aber der Parameter, die die Verteilung auf die verschiedenen Ausfallkombinationen beschreiben, langsam ist. Da die „Startwerte“ (Nullfehlerstatistik) $\langle \omega_{k \setminus r} \rangle = 1/(r-1)$ bzw. $\langle x_{k \setminus r} \rangle = 1/(r+1)$ für große Komponentengruppen sehr klein sind und $q_{k \setminus r} = t \lambda \omega_{k \setminus r}$ bzw. $q_{k \setminus r} = t \zeta x_{k \setminus r}$ gilt, ergibt sich ein Unterschwingen.

Die langsame Konvergenz kann man auch an Gleichungen (5.26) bzw. (5.36) erkennen:

$$\langle \omega_{k \setminus r} \rangle = \frac{2n_{k \setminus r} + 1}{2n_{GVA} + r - 1} \approx \frac{n_{k \setminus r}}{n_{GVA} + r/2} \quad (5.51)$$

bzw.

$$\langle x_{k \setminus r} \rangle = \frac{2n_{k \setminus r} + 1}{2n_{GVA} + r + 1} \approx \frac{n_{k \setminus r}}{n_{Ph} + r/2} \quad (5.52)$$

Eine Konvergenz kann somit erfolgen, wenn die Anzahl der Ereignisse deutlich größer als die Hälfte des Redundanzgrades ist.

Daraus lässt sich schlussfolgern, dass Modelle B und C für größere Komponenten-
gruppen im Allgemeinen nicht geeignet sind, da dann eine sehr starke Unterschätzung
von GVA-Wahrscheinlichkeiten nicht auszuschließen ist. Somit wird das Ziel von Mo-
dell B, auch bei großen Redundanzgraden eine realistische Modellierung zu erreichen,
nicht erreicht.

Auch beim Alpha-Faktor-Modell existiert ein entsprechender Effekt. Beim Alpha-Faktor-
Modell sind jedoch nicht nur systematische Ausfälle, sondern auch Einzelfehler einbe-
zogen. In der Praxis ist die Anzahl der Ereignisse mit einem Ausfall deutlich größer als
die Anzahl GVA. Eine Unterschätzung des Anteils von Einzelausfällen entspricht einer
Überschätzung der GVA-Wahrscheinlichkeiten. Somit sind auch für das Alpha-Faktor-
Modell erhebliche Schätzabweichungen zu erwarten, allerdings unter den genannten
Bedingungen keine Unterschätzungen, sondern nur Überschätzungen von GVA-
Wahrscheinlichkeiten.

5.4 Schätzung der Modellparameter aus der Betriebserfahrung

Nun wird dargestellt, wie die Modellparameter und GVA-Wahrscheinlichkeiten aus der
Betriebserfahrung geschätzt werden können. Wesentliche Elemente dabei sind die
Einbeziehung der Interpretationsunsicherheit der Komponentenschädigungen, des
Übertragbarkeitsfaktors sowie jeweils mehrerer Expertenbewertungen. Dies führt dazu,
dass im Allgemeinen die Gesamtheit der Anzahlen $\mathfrak{N} = \{n_{0 \setminus r}, n_{1 \setminus r} \dots n_{r \setminus r}\}$ nur unsi-
cher bekannt ist. Wenn man die Gesamtheit der relevanten Betriebserfahrung als $\mathfrak{E}$
bezeichnet, kann man diese Beziehung ausdrücken durch die bedingte Wahrschein-
lichkeit $p(\mathfrak{N}|\mathfrak{E})$. Die gesuchte Wahrscheinlichkeitsverteilung der Modellparameter $\mathcal{M}$,
gegeben die Betriebserfahrung $p(\mathcal{M}|\mathfrak{E})$, lässt sich dann schreiben als

$$p(\mathcal{M}|\mathfrak{E}) = \sum_{\text{alle } \mathfrak{N}} p(\mathcal{M}|\mathfrak{N})p(\mathfrak{N}|\mathfrak{E}) \quad (5.53)$$

Die konkrete Bedeutung der zu schätzenden Modellparameter für die verschiedenen
Modelle ist in Tab. 5.2 aufgeführt (siehe Abschnitt 5.5).

5.4.1 Bestimmung der bedingten Wahrscheinlichkeit $p(\mathfrak{N}|\mathfrak{E})$

Im Folgenden wird dargestellt, wie die Wahrscheinlichkeit $p(\mathfrak{N}|\mathfrak{E})$ bestimmt werden kann. Dazu werden zunächst die Wahrscheinlichkeiten für ein Einzelereignis aus den Interpretationsvektoren und den Übertragbarkeitsfaktoren der Experten bestimmt. Wie aus den abgeschätzten Komponentenschädigungen ein Interpretationsvektor berechnet wird, ist in Abschnitt 2.2.1 dargestellt. $W_{i,e}$ bezeichnet im Folgenden den so gebildeten Interpretationsvektor des Experten i zum Ereignis e und $f_{i,e}$ bezeichnet den Übertragbarkeitsfaktor des Experten i zum Ereignis e . Zunächst wird ein modifizierter Interpretationsvektor gebildet, der die Übertragbarkeit beinhaltet. Eine eingeschränkte Übertragbarkeit bedeutet, dass nur mit einer gewissen Wahrscheinlichkeit $f_i < 1$ das dem Ereignis i zugrundeliegende GVA-Phänomen in der Zielgruppe wirksam wird. Statistisch äquivalent ist die Interpretation, dass es zwar mit derselben Wahrscheinlichkeit auftritt, aber mit Wahrscheinlichkeit $1 - f_i > 0$ grundsätzlich keine Schädigungen verursacht. Das kann durch einen modifizierten Interpretationsvektor $V_{i,e}$ ausgedrückt werden, der definiert wird als

$$V_{i,e} = f_{i,e}W_{i,e} + (1 - f_{i,e})(1,0, \dots, 0) \quad (5.54)$$

Für die Elemente gilt also (hierbei wurden die die Ereignisse und Experten bezeichnende Indizes nicht geschrieben):

$$\begin{aligned} v_{0 \setminus r} &= f w_{0 \setminus r} + (1 - f) \\ v_{k \setminus r} &= f w_{k \setminus r} \quad \text{für } 1 \leq k \leq r \end{aligned} \quad (5.55)$$

Der für ein Ereignis insgesamt maßgebliche Interpretationsvektor V_i ergibt sich dann, indem man die Verteilungen der verschiedenen Experten mit gleichem Gewicht mischt, da alle Experten als gleich kompetent in Bezug auf die Ereignisbewertung angesehen werden. Dies entspricht der Mittelung über die expertenspezifischen Interpretationsvektoren $V_{i,e}$:

$$V_i = \frac{1}{NE} \sum_{e=1}^{NE} V_{i,e} \quad (5.56)$$

wobei NE die Anzahl der Experten ist, die Ereignis i bewertet haben.

Nun sind aus den N Interpretationsvektoren V_i die Wahrscheinlichkeiten, dass $z_{k \setminus r}$ mal ein (k von r)-Ereignis beobachtet wurde, für alle relevanten $z_{k \setminus r}$ zu berechnen. Die $z_{k \setminus r}$ sind offensichtlich nicht unabhängig, so dass die Verbundwahrscheinlichkeit von $Z = (z_{0 \setminus r}, \dots, z_{r \setminus r})$ zu betrachten ist. Es gilt

$$0 \leq z_{k \setminus r} \leq N \quad (5.57)$$

und

$$\sum_{k=0}^r z_{k \setminus r} = N \quad (5.58)$$

Die Wahrscheinlichkeit einer bestimmten Kombination von Anzahlen $p(z_{0 \setminus r}, \dots, z_{r \setminus r})$ kann geschrieben werden als Summe über alle denkbaren Kombinationen der Wahrscheinlichkeit der Kombination mal einem Term, der genau dann 1 ist, wenn die Gesamtzahl aller (k von r)-Ausfälle gleich $z_{k \setminus r}$ ist.

Die Wahrscheinlichkeit einer Kombination $(k_1, k_2, \dots, k_N)$ ist die Wahrscheinlichkeit, dass im ersten Ereignis k_1 Komponenten ausgefallen sind und im zweiten k_2 usw., und somit $\prod_{i=1}^N v_{i; k_i \setminus r}$.

Es wird nach oben gesagtem nur zu $p(z_{0 \setminus r}, \dots, z_{r \setminus r})$ beigetragen, wenn $z_{0 \setminus r}$ die Anzahl der k_i ist, die gleich 0 sind und $z_{1 \setminus r}$ die Anzahl der k_i ist, die gleich 1 sind, usw., d. h. wenn $z_{i \setminus r} = \sum_{j=0}^N \delta_{k_j, i}$ für alle $i = 0 \dots r$ gilt.

Es gilt somit

$$p(Z) = \sum_{k_1=0}^r \sum_{k_2=0}^r \dots \sum_{k_N=0}^r \left(\prod_{i=1}^N v_{i; k_i \setminus r} \right) \left(\prod_{i=0}^r \delta_{z_{i \setminus r}, \sum_{j=0}^N \delta_{k_j, i}} \right) \quad (5.59)$$

Zur Illustration werden aus (MD.14) ausgewählte $p(Z)$ berechnet.

1) Wahrscheinlichkeit von N (r von r)-Ereignissen:

Hier ist $Z = (0, \dots, 0, N)$. Das letzte Produkt in (MD.14) ist genau dann 1, wenn alle $k_j = r$ sind. Somit ist $p(Z) = \prod_{i=1}^N v_{i; r \setminus r}$.

2) Wahrscheinlichkeit von 1 (r von r)-Ereignis und $N - 1$ (r-1 von r)-Ereignissen:

Hier ist $Z = (0, \dots, N - 1, 1)$. Das letzte Produkt in (MD.14) ist genau dann 1, wenn genau ein $k_j = r$ ist und alle anderen $k_j = r - 1$ sind. Somit ist

$$p(Z) = v_{1;r \setminus r} \prod_{i=2}^N v_{i;r-1 \setminus r} + v_{1;r-1 \setminus r} v_{2;r \setminus r} \prod_{i=2}^N v_{i;r-1 \setminus r} + \dots + \prod_{i=1}^{N-1} v_{i;r-1 \setminus r} v_{N;r \setminus r}$$

Da die Interpretationsvektoren V_i die gesamten Informationen aus den Expertenbewertungen $\mathfrak{E}$ beinhalten und sich $\mathfrak{N}$ mit Z identifizieren lässt, ist $p(Z)$ aus (MD.14) die gesuchte bedingte Wahrscheinlichkeit.

5.4.2 Schätzung der Modellparameter aus der Betriebserfahrung

Im Folgenden sind die Ausdrücke für die Schätzung der Modellparameter aus der Betriebserfahrung $\mathfrak{E}$ unter Verwendung der oben eingeführten modifizierten Interpretationsvektoren angegeben.

5.4.2.1 Modell A

Für die Verbundwahrscheinlichkeit der Modellparameter gilt:

$$\begin{aligned} p(\mathcal{M}|\mathfrak{E}) &= \sum_{z_{0 \setminus r}=0}^N \sum_{z_{1 \setminus r}=0}^N \dots \sum_{z_{r \setminus r}=0}^N \delta_{\sum_{i=0}^r z_{i \setminus r}, N} \\ &\times \sum_{k_1=0}^r \sum_{k_2=0}^r \dots \sum_{k_N=0}^r \left(\prod_{i=1}^N v_{i; k_i} \right) \left(\prod_{i=2}^r \delta_{z_{i \setminus r}, \sum_{j=0}^N \delta_{k_j, i}} \right) \\ &\times \prod_{k=2}^r \frac{T^{z_{k \setminus r} + 1/2}}{\Gamma(z_{k \setminus r} + 1/2)} (\lambda_{k \setminus r})^{z_{k \setminus r} - 1/2} e^{-\lambda_{k \setminus r} T} \end{aligned} \quad (5.60)$$

Man erkennt, dass die Verteilung nicht faktorisiert. Grund ist, dass die Anzahlen der verschiedenen Ausfallkombinationen nicht unabhängig sind (siehe Gleichung MD.14).

5.4.2.2 Modell B

Für die Verbundwahrscheinlichkeit der Modellparameter gilt:

$$\begin{aligned}
p(\mathcal{M}|\mathfrak{E}) &= \sum_{z_{0 \setminus r}=0}^N \sum_{z_{1 \setminus r}=0}^N \dots \sum_{z_{r \setminus r}=0}^N \delta_{\sum_{i=0}^r z_{i \setminus r}, N} \\
&\times \sum_{k_1=0}^r \sum_{k_2=0}^r \dots \sum_{k_N=0}^r \left(\prod_{i=1}^N v_{i; k_i} \right) \left(\prod_{i=2}^r \delta_{z_{i \setminus r}, \sum_{j=0}^N \delta_{k_j, i}} \right) \\
&\times \frac{T(\sum_{k=2}^r z_{k \setminus r}) + 1/2}{\Gamma((\sum_{k=2}^r z_{k \setminus r}) + 1/2)} (\lambda)^{(\sum_{k=2}^r z_{k \setminus r}) - 1/2} e^{-\lambda T} \\
&\times \frac{\prod_{i=2}^r (\omega_{i \setminus r})^{n_{k \setminus r} - 1/2}}{B(n_{2 \setminus r} + 1/2, n_{3 \setminus r} + 1/2 \dots n_{r \setminus r} + 1/2)}
\end{aligned} \tag{5.61}$$

Diese Verteilung faktorisiert ebenfalls nicht.

5.4.2.3 Modell C

Für die Verbundwahrscheinlichkeit der Modellparameter gilt:

$$\begin{aligned}
p(\mathcal{M}|\mathfrak{N}) &= \sum_{z_{0 \setminus r}=0}^N \sum_{z_{1 \setminus r}=0}^N \dots \sum_{z_{r \setminus r}=0}^N \delta_{\sum_{i=0}^r z_{i \setminus r}, N} \\
&\times \sum_{k_1=0}^r \sum_{k_2=0}^r \dots \sum_{k_N=0}^r \left(\prod_{i=1}^N v_{i; k_i} \right) \left(\prod_{i=2}^r \delta_{z_{i \setminus r}, \sum_{j=0}^N \delta_{k_j, i}} \right) \\
&\times \frac{T(\sum_{k=0}^r z_{k \setminus r}) + 1/2}{\Gamma((\sum_{k=0}^r z_{k \setminus r}) + 1/2)} (\lambda)^{(\sum_{k=0}^r z_{k \setminus r}) - 1/2} e^{-\lambda T} \\
&\times \frac{\prod_{i=2}^r (\omega_{i \setminus r})^{n_{k \setminus r} - 1/2}}{B(n_{2 \setminus r} + 1/2, n_{3 \setminus r} + 1/2 \dots n_{r \setminus r} + 1/2)}
\end{aligned} \tag{5.62}$$

Da aufgrund der Deltafunktion in der ersten Zeile von Gleichung (6.62) für den Term in der dritten Zeile $\sum_{i=0}^r z_{i \setminus r} = N$ gilt, lässt sich dieser Term nach vorne ziehen und es gilt

$$p(\mathcal{M}|\mathfrak{N}) = p(\zeta|\mathfrak{E}) p(x_{0 \setminus r}, x_{1 \setminus r}, \dots, x_{r-1 \setminus r}|\mathfrak{E}) \quad (5.63)$$

mit

$$p(\zeta|\mathfrak{E}) = p(\zeta|\mathfrak{N}) = \frac{T^{N+1/2}}{\Gamma(N+1/2)} (\zeta)^{N-1/2} e^{-\zeta T} \quad (5.64)$$

und

$$\begin{aligned} p(x_{0 \setminus r}, x_{1 \setminus r}, \dots, x_{r-1 \setminus r}|\mathfrak{E}) &= \sum_{z_{0 \setminus r}=0}^N \sum_{z_{1 \setminus r}=0}^N \dots \sum_{z_{r-1 \setminus r}=0}^N \delta_{\sum_{i=0}^r z_{i \setminus r}, N} \\ &\times \frac{\prod_{i=0}^r (x_{i \setminus r})^{z_{i \setminus r}-1/2}}{B(z_{0 \setminus r}+1/2, z_{1 \setminus r}+1/2, z_{2 \setminus r}+1/2, \dots, z_{r-1 \setminus r}+1/2)} \\ &\times \sum_{k_1=0}^r \sum_{k_2=0}^r \dots \sum_{k_N=0}^r \left(\prod_{i=1}^N v_{i; k_i \setminus r} \right) \left(\prod_{i=0}^r \delta_{z_{i \setminus r}, \sum_{j=0}^N \delta_{k_j, i}} \right) \end{aligned} \quad (5.65)$$

Dieser Ausdruck ist recht kompliziert, er beinhaltet $r \times N$ Summationen.

5.4.3 Implementation der Schätzalgorithmen mit Monte-Carlo-Verfahren

Während die analytischen Formeln kompliziert erscheinen, ist die Vorgehensweise zur Implementation mittels Monte-Carlo einfacher nachvollziehbar. Welche Schritte für eine Monte-Carlo-Rechnung durchzuführen sind, wird zunächst für Modell A erläutert:

1. Berechnung der expertenspezifischen modifizierten Interpretationsvektoren aus den von den Experten abgeschätzten Komponentenschädigungen und dem Übertragbarkeitsfaktor nach Gleichung (5.54) für alle Ereignisse und Experten.
2. Berechnung der nicht-expertenspezifischen modifizierten Interpretationsvektoren nach Gleichung (5.56) für alle Ereignisse.

3. Führe für jedes Sample aus:

- a) Setze den $r + 1$ dimensionalen Vektor von Anzahlen $Z = (0, \dots, 0)$
- b) Führe für alle Ereignisse aus:
 - i. Ziehe eine Ausfallkombination aus der Kategorienverteilung mit Parametervektor v_i . d. h. mit Wahrscheinlichkeit $v_{i;k \setminus r}$ wird ein $(k \text{ von } r)$ -Ereignis gezogen.
 - ii. erhöhe das $k + 1$ te Element von Z um 1.
Die Elemente von Z enthalten nun Samples der Anzahlen der $(k \text{ von } r)$ -Ereignissen, d. h. $Z = (n_{0 \setminus r}, n_{1 \setminus r}, \dots, n_{r \setminus r})$
- c) Führe für alle Ausfallkombinationen $(2 \text{ von } r) \dots (r \text{ von } r)$ aus:
 - iii. Ziehe ein Sample von $q_{k \setminus r}$ aus $p(q_{k \setminus r} | n_{k \setminus r})$ (Gleichung (5.11))
Dann stellt $Q = (q_{2 \setminus r}, q_{3 \setminus r}, \dots, q_{r \setminus r})$ ein Sample der gesuchten Verteilung $p(Q | \mathcal{E})$ dar.

Die Schätzung nach Modell B lässt sich implementieren, indem statt Schritt 3 c des oben beschriebenen Algorithmus ausgeführt wird:

- c) Ziehe ein Sample von $(\omega_{2 \setminus r}, \omega_{3 \setminus r}, \dots, \omega_{r \setminus r})$ aus $p(\omega_{2 \setminus r}, \omega_{3 \setminus r}, \dots, \omega_{r \setminus r} | Z)$ (Gleichung (6.17)).
- d) Berechne die Anzahl GVA $n_{GVA} = \sum_{k=2}^r n_{k \setminus r}$
- e) Ziehe ein Sample von $p(\lambda | n_{GVA})$ (Gleichung (5.18))
- f) Führe für alle Ausfallkombinationen $(2 \text{ von } r) \dots (r \text{ von } r)$ aus:
 - i. Berechne ein Sample von $q_{k \setminus r}$ als $q_{k \setminus r} = t \lambda \omega_{k \setminus r}$ (Gleichung (5.3))
Dann stellt $Q = (q_{2 \setminus r}, q_{3 \setminus r}, \dots, q_{r \setminus r})$ ein Sample der gesuchten Verteilung $p(Q | \mathcal{E})$ dar.

Die Schätzung nach Modell C lässt sich implementieren, indem statt Schritt 3 c des oben beschriebenen Algorithmus ausgeführt wird:

- c) Ziehe ein Sample von $(x_{0 \setminus r}, x_{1 \setminus r}, \dots, x_{r \setminus r})$ aus $p(x_{0 \setminus r}, x_{1 \setminus r}, \dots, x_{r \setminus r} | Z)$ (Gleichung (6.25))
- d) Ziehe ein Sample von $p(\zeta | N)$ (Gleichung ((5.31))

- e) Führe für alle Ausfallkombinationen (0 von r) ... (r von r) aus:
- i. Berechne ein Sample von $q_{k \setminus r}$ als $q_{k \setminus r} = t \zeta x_{k \setminus r}$ (Gleichung (5.6))
- Dann stellt $Q = (q_{0 \setminus r}, q_{1 \setminus r}, \dots, q_{r \setminus r})$ ein Sample der gesuchten Verteilung $p(Q|\mathcal{E})$ dar.

5.4.4 Berücksichtigung der verbleibenden Unsicherheitsquellen

Die verbleibenden Unsicherheitsquellen können analog wie bei der aktuellen Version des Kopplungsmodells mit dem in Abschnitt 2.2.2 beschriebenen Verfahren einbezogen werden. Das Verfahren lässt sich problemlos in die oben beschriebene Monte-Carlo-Rechnung integrieren:

1. Führe für jedes Sample $q_{k \setminus r}$ aus:

- a) Ziehe ein Sample aus $p(\hat{q}_{k \setminus r} | q_{k \setminus r})$

Dann stellt $\hat{Q} = (\hat{q}_{2 \setminus r}, \hat{q}_{3 \setminus r}, \dots, \hat{q}_{r \setminus r})$ bzw. für Modell C $\hat{Q} = (\hat{q}_{0 \setminus r}, \hat{q}_{1 \setminus r}, \dots, \hat{q}_{r \setminus r})$ ein Sample der Verteilung der GVA-Wahrscheinlichkeiten bzw. Wahrscheinlichkeiten systematischer (k von r)-Ausfälle unter Berücksichtigung der verbleibenden Unsicherheiten $p(\hat{Q}|\mathcal{E})$ dar.

5.5 Vergleich der Modelle

Die zu schätzenden Modellparameter für die verschiedenen Modelle ist in Tab. 5.2 aufgeführt. Unabhängig von der Betriebserfahrung festzulegende Modellparameter (hier die Fehlerentdeckungszeit t) sind nicht enthalten.

Tab. 5.2 Zu schätzende Modellparameter

Modell	Zu schätzender Modellparameter	Anzahl der unabhängigen Parameter	Für die Schätzung benötigte Ereignisanzahlen
A	$\mathcal{M} = \{\lambda_{2 \setminus r}, \lambda_{3 \setminus r}, \dots, \lambda_{r \setminus r}\}$	$r - 1$	$n_{2 \setminus r}, n_{3 \setminus r}, \dots, n_{r \setminus r}$
B	$\mathcal{M} = \{\lambda, \omega_{2 \setminus r}, \omega_{3 \setminus r}, \dots, \omega_{r \setminus r}\}$	$r - 1$	$n_{2 \setminus r}, n_{3 \setminus r}, \dots, n_{r \setminus r}$
C	$\mathcal{M} = \{\zeta, x_{0 \setminus r}, x_{1 \setminus r}, \dots, x_{r \setminus r}\}$	$r + 1$	$n_{0 \setminus r}, n_{1 \setminus r}, \dots, n_{r \setminus r}$

Wie oben bereits beschrieben, ist $p(\Omega|\mathfrak{N})$ nur für Modell A analytisch darstellbar. Nur für dieses Modell faktorisiert diese Wahrscheinlichkeitsdichte, so dass $q_{k \setminus r}$ unabhängig sind, gegeben $\mathfrak{N}$.

Allerdings ist $\mathfrak{N}$ im Allgemeinen nicht genau bekannt, da in der Praxis in allen GVA-Datensätzen (im Folgenden als $\mathfrak{E}$ bezeichnet) Ereignisse enthalten sind, in denen nicht nur ausgefallene und nicht betroffene, sondern auch geschädigte Komponenten beobachtet wurden. Dies führt dann dazu, dass $p(\Omega|\mathfrak{E})$ auch für Modell A im Allgemeinen nicht faktorisiert.

Modell A basiert nur auf der Annahme, dass die Wahrscheinlichkeit eines GVA sehr klein ist, so dass das Wirksamwerden zweier GVA-Phänomene vernachlässigt werden kann, bevor der GVA entdeckt wird.

Modell B und C enthalten zusätzlich einen hypothetischen generischen Zustand ‘GVA’ bzw. „GVA-Phänomen hat eingewirkt“. Diese Annahme beeinflusst die Wahl der a priori-Verteilungen und somit die Schätzergebnisse. Es kann eine starke Unterschätzung der Wahrscheinlichkeiten einzelner GVA-Ausfallkombinationen (z. B. kompletter GVA) auftreten.

Demgegenüber treten bei Modell A nur Überschätzungen auf. Deshalb ist nur mit Modell A eine konservative Schätzung der GVA-Wahrscheinlichkeiten möglich.

Für alle Modelle erfolgt eine Berücksichtigung der Unsicherheiten in dem Umfang, wie sie beim Kopplungsmodell etabliert ist. Die Modelle lassen sich effizient in Form von Monte-Carlo-Verfahren implementieren.

Für alle Modelle sind separate Mappingverfahren erforderlich, wenn Ereignisse an Komponentengruppen abweichender Größe Teil der zur Quantifizierung verwendeten Betriebserfahrung sind.

Zusammenfassend scheint Modell A für die Schätzung von GVA-Wahrscheinlichkeiten am besten geeignet, während Modell C Vorteile für theoretische Betrachtungen des Mappings aufweist (siehe Abschnitt 6.2.2).

6 Mapping

Wie oben bereits erwähnt, liegt im Allgemeinen nicht für jede Komponentengruppe, die in der PSA modelliert ist, ausreichend Betriebserfahrung für Gruppen derselben Größe vor, um eine hinreichend genaue Schätzung von GVA-Wahrscheinlichkeiten zu ermöglichen. In diesen Fällen muss auch die Betriebserfahrung aus Komponentengruppen abweichender Größe verwendet werden. Dann ist die Übertragung von Ereignissen zwischen Komponentengruppen verschiedener Größe, das so genannte Mapping, erforderlich.

Beim Mapping wird häufig unterschieden nach einer Übertragung von einer größeren in eine kleinere Komponentengruppe ('Mapping Down') und einer Übertragung von einer kleineren in eine größere Komponentengruppe ('Mapping Up').

Wie oben dargestellt, ist im Kopplungsmodell aufgrund der Modellannahme des unabhängigen Ausfalls der Komponenten bei Auftritt des GVA-Phänomens mit Wahrscheinlichkeit η (Kopplungsparameter) ein „automatisches“ Mapping im Modell vorgesehen. Für andere Modelle, insbesondere die in Abschnitt 5.2 beschriebenen, sind gesonderte Modellansätze und Verfahren für das Mapping erforderlich. Diese werden im Folgenden diskutiert. Hierbei werden zunächst grundlegende verschiedenen Möglichkeiten diskutiert, wie GVA-Phänomene wirken können und welche Auswirkungen sich daraus auf das Mapping ergeben.

Die verschiedenen Verfahren werden zunächst anhand der Bewertungen eines Experten dargestellt. Mehrere Expertenbewertungen können wie in Abschnitt 2 (Gleichung (2.12)) einbezogen werden.

6.1 Phänomene mit Ausfall aller Komponenten

Manche GVA-Phänomene können stets alle Komponenten betreffen, unabhängig von ihrer Anzahl. Denkbare GVA-Phänomene umfassen die folgenden Komponentengruppen:

- Komponenten können nicht vollständig anforderungsgerecht geprüft werden. Deshalb bleibt ein systematischer Auslegungs- oder Herstellungsfehler unentdeckt, der im Anforderungsfall zum Ausfall führt.

- Bei gleichzeitiger Modifikation aller redundanten Komponenten wird ein Fehler eingebracht.
- Die redundanten Komponenten werden einer Einwirkung (z. B. Wasserschlag) ausgesetzt, die zum Ausfall führt.

Bei dieser Art von GVA-Phänomenen fallen stets alle Komponenten aus, unabhängig von der Größe der Komponentengruppe.

Bei der Modellierung von GVA mit dem Kopplungsmodell in der jetzigen Form wird eine solche Information nicht berücksichtigt. Dies führt dazu, dass auch bei einem beobachteten vollständigen GVA in der Unsicherheitsverteilung der Kopplungsparameter auch Werte $\eta < 1$ vorkommen und deshalb auch mit einer bestimmten Wahrscheinlichkeit in der Zielkomponentengruppe unvollständige GVA auftreten. Dies ist eine unmittelbare Folge des Satzes von Bayes: Weil auch bei $\eta < 1$ mit einer gewissen Wahrscheinlichkeit alle Komponenten ausfallen, ist auch bei Ausfall aller Komponenten der wahre Wert von η mit einer gewissen Wahrscheinlichkeit kleiner 1. Somit fallen bei einem GVA mit diesem Phänomen mit endlicher Wahrscheinlichkeit nicht alle Komponenten aus. Dies hat insbesondere bei kleinen Komponentengruppen eine hohe Bedeutung, da in diesem Fall die Unsicherheit der Schätzung des Kopplungsparameters besonders groß ist.

6.1.1 Deutsche Betriebserfahrung

In der deutschen Betriebserfahrung sind etwa 30 Ereignisse mit einem Ausfall aller Komponenten einer GVA-Komponentengruppe aufgetreten. Darunter befinden sich auch Ereignisse, bei denen davon auszugehen ist, dass – unabhängig von der Größe der Komponentengruppe – alle Komponenten ausgefallen wären. Im Folgenden sind einige Beispiele kurz angeführt:

- Aufgrund einer missverständlichen Prüfanweisung wurde bei einer wiederkehrenden Prüfung (WKP) am Reaktorschutzsystem die automatische Dieselanregung für alle Diesel unterbrochen.
- Ein Stabeinwurf erfolgte nicht aufgrund fehlerhafter Simulationen in der digitalen Leittechnik. Als Ursache wurden die zugrunde gelegten Planungsunterlagen identifiziert, die wichtige Zusammenhänge nicht in ausreichender Tiefe und Klarheit zeigten.

- Bypassklappen konnten aufgrund eines Auslegungsfehlers nicht gegen auftretenden Differenzdruck öffnen.
- Bei einer Modifikation der Notstromdieselversorgung wurde ein neuer Spannungsregler installiert, der schneller als vorher die erforderliche Nennspannung erreicht. Damit wird das Verbraucherzuschaltprogramm schon freigegeben, wenn sich der Diesel noch in der Hochlaufphase befindet. Das Zuschalten des ersten Verbrauchers kann damit noch nicht vom Drehzahlregler ausgeregelt werden, so dass eine Unterspannung auftritt, die zu einem Ausfall der Notstromschiene führt.
- Kondensatgefäße bei Niveaumesseinrichtungen wurden bei Errichten der Anlage falsch montiert. Dies kann bei bestimmten Anlagenzuständen zu einem erheblichen systematischen Fehler der Messwerte führen.

Diese Beispiele illustrieren, dass Ausfälle aller Komponenten häufig dann auftreten, wenn die Fehler nicht durch wiederkehrende Prüfungen (WKP) erfasst werden, weil diese nicht für alle Anforderungsfälle vollständig abdeckend sind bzw. wenn alle Komponenten modifiziert worden sind, bevor eine WKP durchgeführt wurde. Die uneingeschränkte Übertragbarkeit (Übertragbarkeitsfaktor $f = 1$) ist deshalb meist nur für bestimmte Anlagenzustände gegeben (z. B. Anlagenzustände vor Anfahrprüfungen). Während für den Leistungsbetrieb daher der Beitrag dieser Ereignisse im Allgemeinen nicht dominant ist, kann für PSA-Rechnungen, die sich auf für die Ereignisse jeweils zutreffende Anlagenzustände beziehen, ein erheblicher Einfluss auf das quantitative Ergebnis nicht ausgeschlossen werden.

Somit erscheint es sinnvoll, bei der Weiterentwicklung der Verfahren zum Mapping von Ereignissen sicherzustellen, dass Phänomene, die zum Ausfall aller Komponenten führen, angemessen berücksichtigt werden.

Hierzu können entweder Mappingverfahren verwendet werden, die die Eigenschaft aufweisen, dass GVA mit Ausfall aller Komponenten immer auf GVA mit Ausfall aller Komponenten abgebildet werden, oder es werden die GVA in zwei Kategorien aufgeteilt:

- GVA-Phänomene, die notwendig zum Ausfall aller Komponenten führen,
- GVA-Phänomene, die nicht notwendig zum Ausfall aller Komponenten führen.

Für erstere ist das Mapping trivial: Alle Komponenten sind ausgefallen, unabhängig von der Komponentengruppengröße. Für die verbleibenden wird ein anderes (z. B. auf dem Kopplungsmodell basierendes) Mappingverfahren angewandt. Um diesen Ansatz anwenden zu können, müssten jedoch alle GVA-Ereignisse mit Ausfall aller Komponenten dahingehend bewertet werden, ob ein Phänomen vorliegt, das zum Ausfall aller Komponenten führt. Eine solche Bewertung wurde im Rahmen der GVA-Ereignisbewertungen bisher nicht durchgeführt. Deshalb sind die Eingabe, Speicherung und Verarbeitung entsprechender Daten in dem Datenbanksystem POOL/PEAK der GRS nicht vorgesehen. Eine Verwendung dieser Informationen würde also eine Bewertung aller GVA-relevanten Ereignisse und eine Erweiterung der Datenbanken, der Oberflächen und Werkzeuge sowie ihrer Schnittstellen erfordern.

6.2 Phänomene ohne notwendige Ausfall aller Komponenten

Bei Phänomenen, die nicht notwendig zum Ausfall aller Komponenten führen, sind umfangreichere Betrachtungen, ggf. unter Verwendung von einschränkenden Annahmen, erforderlich. Im Folgenden werden solche Ansätze im Detail untersucht.

6.2.1 Modellbasierter Ansatz

Bei einem modellbasierten Ansatz wird ein GVA-Modell verwendet, um abzuschätzen, wie das aufgetretene GVA-Phänomen in der Komponentengruppe abweichender Größe gewirkt hätte. Als Modell könnte z. B. das Kopplungsmodell eingesetzt werden. Dazu ist folgende Vorgehensweise möglich:

Zunächst wird – wie beim Kopplungsmodell üblich – der Kopplungsparameter η_j aus den Komponentenschädigungen bestimmt. Wie in Abschnitt 2 dargestellt, ergibt sich aus den subjektiven Wahrscheinlichkeiten w_i der verschiedenen Ausfallkombinationen die Wahrscheinlichkeitsdichte $p(\eta_j)$ für den Kopplungsparameter η_j des GVA-Ereignisses j :

$$p(\eta_j) = \sum_{i=0}^r w_{i \setminus r} \frac{\Gamma(r+1)}{\Gamma(i+1/2) \Gamma(r-i+1/2)} \eta_j^{i-1/2} (1-\eta_j)^{r-i-1/2} \quad (6.1)$$

Mittels dieser Verteilung des Kopplungsparameters werden die Wahrscheinlichkeiten der verschiedenen Ausfallkombinationen berechnet. Hat die Zielkomponentengruppe

die Größe $\hat{r}$, so folgt für die Wahrscheinlichkeit der verschiedenen Ausfallkombinationen, da sie einer Binomialverteilung genügen,

$$p(\hat{k} \setminus \hat{r}) = \int_0^1 p(\hat{k} \setminus \hat{r} | \eta_j) p(\eta_j) d\eta_j = \binom{\hat{r}}{\hat{k}} \int_0^1 \eta_j^{\hat{k}} (1 - \eta_j)^{\hat{r} - \hat{k}} p(\eta_j) d\eta_j \quad (6.2)$$

mit $p(\eta_j)$ aus Gleichung (2.7). Daraus folgt

$$p(\hat{k} \setminus \hat{r}) = \sum_{i=0}^r w_{i \setminus r} \frac{\Gamma(r+1)}{\Gamma(i+1/2) \Gamma(r-i+1/2)} \binom{\hat{r}}{\hat{k}} \times \int_0^1 \eta_j^{\hat{k}} (1 - \eta_j)^{\hat{r} - \hat{k}} \eta_j^{i-1/2} (1 - \eta_j)^{r-i-1/2} d\eta_j \quad (6.3)$$

Durch Auswertung des Integrals ergibt sich

$$p(\hat{k} \setminus \hat{r}) = \binom{\hat{r}}{\hat{k}} \sum_{i=0}^r w_{i \setminus r} \frac{\Gamma(r+1)}{\Gamma(i+1/2) \Gamma(r-i+1/2)} \times \frac{\Gamma(i + \hat{k} + 1/2) \Gamma(r + \hat{r} - i - \hat{k} + 1/2)}{\Gamma(r + \hat{r} + 1)} \quad (6.4)$$

Dies stellt die einfachste Möglichkeit des Mappings dar. Allerdings basiert es wie das Kopplungsmodell auf der Annahme der Binomialverteilung, die in der Praxis nicht bzw. nicht streng erfüllt ist.

Es stellt sich also die Frage, wie ein Verfahren entwickelt werden kann, das auf möglichst wenig einschränkenden Annahmen beruht.

6.2.2 Probabilistisch-kombinatorischer Ansatz

Die oben diskutierte Annahme der Binomialverteilung mit einem Kopplungsparameter η_j beinhaltet die Annahme, dass eine Komponentengruppe der Größe r statistisch

äquivalent einer beliebigen Untergruppe der Größe r einer größeren Gruppe mit Größe $\hat{r} > r$ ist¹³. Mit dieser Annahme alleine können auch Mappingverfahren begründet werden. Dies wird im Folgenden diskutiert.

Diskussion der Grundannahme des probabilistisch-kombinatorischen Ansatzes

Der oben dargestellte probabilistisch-kombinatorischen Ansatz erscheint zunächst sehr naheliegend, allerdings gibt es auch Argumente, die seine Gültigkeit in Frage stellen. Dies sind insbesondere:

- Einfluss der Entdeckungsmechanismen der GVA
- Auswirkungen der Annahme auf die Unsicherheitsanalyse

Für die Nicht-Verfügbarkeiten durch GVA sind neben den Mechanismen für die Entstehung von gemeinsam verursachten Ausfällen auch die Mechanismen für ihre Entdeckung ausschlaggebend.

Für GVA-Entstehungsmechanismen erscheint es zunächst plausibel, dass sich eine kleine Gruppe verhält wie eine Untergruppe einer großen Gruppe, da bei Wirksamwerden des GVA-Phänomens die Wirkung auf die einzelnen Komponenten nicht voneinander abhängig ist. Beispielsweise hängt die Schädigung einer Komponente durch unzulässige Belastung nicht davon ab, wie die Belastung auf andere Komponenten wirkt, sondern nur von der Belastung und der Komponente selbst.

Bei der Entdeckung von GVA gilt kein analoges Argument. Vielmehr hat die Entdeckung von Schädigungen von Komponenten Einfluss auf die Untersuchung von weiteren Komponenten: Wird eine systematische Ursache vermutet, so werden typischerweise die weiteren Komponenten zeitnah geprüft und damit ein GVA entdeckt. Aus diesem Grund wird beim Kopplungsmodell die Fehlerentdeckungszeit bei versetztem Testen als Doppeltes des zeitlichen Abstands aufeinanderfolgender Prüfungen einzelner Komponenten angesetzt.

In Bezug auf die Unsicherheitsanalyse ist die Kompatibilität der Annahme mit den gewählten a priori-Verteilungen sicherzustellen, da diese Annahme eine Information über

¹³ Daraus folgt unter anderem, dass alle Komponenten statistisch äquivalent sind.

den a priori darstellt und damit eine nichtinformative Wahl der a priori-Verteilung nicht möglich ist. Dies wird in Abschnitt 6.3 detaillierter diskutiert.

Probabilistisch-kombinatorisches Mapping Down

Der einfachste Fall stellt das ‘Mapping Down’ dar. Dabei wurde das GVA-Ereignis in einer Komponentengruppe der Größe r beobachtet. Die Zielkomponentengruppe hat die Größe $\check{r}$, wobei $\check{r} < r$ ist. Für das ‘Mapping Down’ wird wie oben dargestellt angenommen, dass sich die Zielkomponentengruppe im statistischen Sinne verhält wie eine Untergruppe der Größe $\check{r}$. Anders ausgedrückt ergibt sich das Mapping durch Weglassen von jeweils $r - \check{r}$ zufälligen Komponenten und Berücksichtigung aller möglichen Fälle, welche Komponenten weggelassen werden, mit gleicher Wahrscheinlichkeit. In der gleichen Wahrscheinlichkeit drückt sich die statistische Äquivalenz aller Komponenten aus; die Wahrscheinlichkeiten der verschiedenen Komponentenauswahlen berücksichtigt die mit dem Mapping verknüpfte Unsicherheit.

Beispiel: In einer Komponentengruppe der Größe 3 wurden

- ein Ausfall,
- eine starke Schädigung und
- eine sehr schwache Schädigung

beobachtet. In Tab. 6.1 sind die möglichen Kombinationen in einer Komponentengruppe der Größe 2 mit ihren Gewichten aufgeführt:

Tab. 6.1 Schädigungsvektor und jeweilige Wahrscheinlichkeit

Schädigungsvektor in der Komponentengruppe der Größe 2	Wahrscheinlichkeit
(Ausfall, starke Schädigung)	1/3
(Ausfall, sehr schwache Schädigung)	1/3
(starke Schädigung, sehr schwache Schädigung)	1/3

Wie in Abschnitt 2 beschrieben, werden aus den Schädigungsvektoren die Wahrscheinlichkeiten w_i der verschiedenen Ausfallkombinationen berechnet; jetzt sind aber zusätzlich noch die Wahrscheinlichkeiten der verschiedenen Schädigungsvektoren mit

zu berücksichtigen, die die mit dem Mapping verknüpften statistischen Unsicherheiten quantifizieren.

Als Beispiel sind für eine Komponentengruppe mit zwei Komponenten die Berechnungsvorschriften in Tab. 6.2 dargestellt, wobei ein Ereignis an einer Komponentengruppe der Größe 3 beobachtet wurde, bei dem die Komponentenschädigungswerte d_1, d_2 und d_3 waren (zur Definition der numerischen Schädigungswerte siehe (Tab. 2.1).

Tab. 6.2 Berechnungsvorschriften für das Mapping von einer Komponentengruppe mit drei Komponenten auf eine Komponentengruppe mit zwei Komponenten

Ausfallkombination	Berechnung des zugehörigen Wertes der Komponente des Interpretationsvektors
0 von 2	$w_{0\setminus 2} = \frac{1}{3}((1 - d_1)(1 - d_2) + (1 - d_1)(1 - d_3) + (1 - d_2)(1 - d_3))$
1 von 2	$w_{1\setminus 2} = \frac{1}{3}(d_1(1 - d_2) + (1 - d_1)d_2 + d_1(1 - d_3) + (1 - d_1)d_3 + d_2(1 - d_3) + (1 - d_2)d_3)$
2 von 2	$w_{2\setminus 2} = \frac{1}{3}(d_1 d_2 + d_1 d_3 + d_2 d_3)$

Die rechte Seite der Gleichungen in Tab. 6.2 lässt sich statt als Funktion der Komponentenschädigungen auch als Funktionen der Elemente der Interpretationsvektoren in der Komponentengruppe der Größe 3 (d. h. der Wahrscheinlichkeiten, dass k von 3 Komponenten ausgefallen sind, $w_{k\setminus 3}$, $k = 0, \dots, 3$) schreiben. Entsprechende Ausdrücke sind z. B. in /MOS 89/ für kleine Komponentengruppengrößen tabelliert (Tabelle C-3 auf Seite C-7).

Tab. 6.3 Ergebnisse für das Mapping eines Ereignisses in einer Komponentengruppe mit drei Komponenten, bei der ein Ausfall ($d_1 = 1$), eine starke Schädigung ($d_2 = 1/2$) und eine sehr schwache Schädigung ($d_3 = 1/100$) beobachtet wurde, auf eine Komponentengruppe mit zwei Komponenten

Ausfallkombination	Zugehöriger Wert der Komponente des Interpretationsvektors
0 von 2	$w_0 = \frac{99}{600}$
1 von 2	$w_1 = \frac{199}{300}$
2 von 2	$w_2 = \frac{103}{600}$

Mapping Up

Beim Mapping Up ist die Zielkomponentengruppengröße größer als die Komponenten-gruppengröße, wo das GVA-Ereignis aufgetreten ist. Hier können keine Komponenten weggelassen werden, sondern aus den Informationen über die beobachteten Komponenten muss geschlossen werden, wie sich die verbleibenden verhalten hätten. Dies ist folgendem Szenario äquivalent: Eine Komponentengruppe der Größe $\hat{r}$ wurde unvollständig beobachtet, indem nur die Schädigungen von $r < \hat{r}$ Komponenten erfasst wurden. Die Schädigungen der $\hat{r} - r$ nicht beobachteten Komponenten sollen aus denjenigen der beobachteten Komponenten abgeleitet werden.

Für das Schätzen der Schädigungen der $\hat{r} - r$ nicht beobachteten Komponenten sind verschiedene Ansätze denkbar, die im Folgenden dargestellt werden.

Probabilistisch-kombinatorisches Mapping Up

Ähnlich den probabilistisch-kombinatorischem Mapping Down lässt sich ein probabilistisch-kombinatorisches Mapping Up konstruieren. Grundlage des Mappings ist alleine die Annahme, dass eine kleine Komponentengruppe als Teil einer großen angesehen werden kann, die nicht vollständig beobachtet wird. Aus den Beobachtungen der kleinen Gruppe wird auf die große Gruppe geschlossen. Allerdings unterscheidet sich das statistische Mapping Up vom statistischen Mapping Down, dass sich die Ereignisse

nicht einzeln, sondern nur in ihrer Gesamtheit übertragen lassen. Dies lässt sich an einem einfachen hypothetischen Beispiel erläutern:

In einer Komponentengruppe der Größe 3 wird eine große Anzahl von (3 von 3)-GVA beobachtet, jedoch weder (2 von 3)-Ereignisse noch Ereignisse mit systematischer Ursache, bei der nur eine Komponente ausgefallen ist (1 von 3)-Ereignisse). Wird die Komponentengruppe als Teil einer Gruppe der Größe 4 angesehen und angenommen, dass die Komponenten unabhängig sind, so lässt sich schließen, dass in der Gruppe der Größe 4 mit sehr hoher Wahrscheinlichkeit (4 von 4)-GVA auftreten. Dies folgt aus dem Satz von Bayes: Würden in der Gruppe der Größe 4 auch (2 von 4)-GVA oder (3 von 4)-GVA auftreten, so würden in der beobachteten Gruppe der Größe 3 auch mit hoher Wahrscheinlichkeit (2 von 3)-Ereignisse oder Ereignisse mit nur einem Ausfall beobachtet werden. Da dies nicht der Fall ist, folgt, dass mit sehr hoher Wahrscheinlichkeit (4 von 4)-GVA auftreten.

Es ist offensichtlich, dass ein solcher Schluss nur aus der Gesamtheit der Beobachtungen in der Komponentengruppe der Größe 3 folgt. Das probabilistisch-kombinatorische Mapping-Up-Verfahren ist – im Gegensatz zum Mapping Down – nicht als Mapping einzelner Ereignisse darstellbar.

Im Folgenden wird das Verfahren mathematisch ausformuliert.

Mathematische Formulierung

Es soll aus Beobachtungen in Komponentengruppen der Größe r auf GVA-Wahrscheinlichkeiten in Komponentengruppen der Größe $\hat{r} > r$ geschlossen werden.

Um eine Bayes'sche Vorgehensweise zu entwickeln, ist die bedingte Wahrscheinlichkeit der Modellparameter für die Komponentengruppen der Größe $\hat{r}$, im Folgenden als X bezeichnet, gegeben die Beobachtungen (dies sind die Beobachtungen in der Komponentengruppe der Größe r , im Folgenden als M bezeichnet), zu bestimmen.

Im Folgenden wird nur das Modell C betrachtet, da, wie in Abschnitt 5.2.3 diskutiert, bei diesem Modell Ereignisse in Komponentengruppen, von denen nur eine Untermenge beobachtet wird, wieder in jedem Fall Ereignisse erhalten werden, die durch das Modell beschrieben werden. Eignet sich zum Beispiel in einer Komponentengruppe der Größe $\hat{r} = 3$ ein Ereignis mit zwei Ausfällen, so kann dies in einer Untergruppe der

Größe zwei als (2 von 2)-Ereignis oder (1 von 2)-Ereignis erscheinen. Von Modell A und B werden nur erstere erfasst. Nur in Modell C stellen alle möglichen Ergebnisse modellierte Ereignisse dar. Somit ist die Ereignisanzahl im Modell C vom Mapping unabhängig; deshalb kann hier der die Ereignisrate beschreibende Parameter ζ unabhängig vom Mapping geschätzt werden (siehe Gleichung (6.22)). Wegen dieser Eigenschaft ist eine analytische Betrachtung, wie sie unten erfolgt, sinnvoll möglich.

Um eine Bayes'sche Vorgehensweise zu entwickeln, könnte man die Beziehung der tatsächlichen Beobachtungen M von (hypothetischen) Beobachtungen in der Komponentengruppe der Größe $\hat{r}$, im Folgenden als $\hat{M}$ bezeichnet, betrachten, die wiederum unmittelbar von den Modellparametern X abhängen. Die Abhängigkeit der M von den $\hat{M}$ ergibt sich daraus, dass die M die Beobachtungen von $r < \hat{r}$ Komponenten umfassen, während die $\hat{M}$ alle $\hat{r}$ Komponenten umfassen.

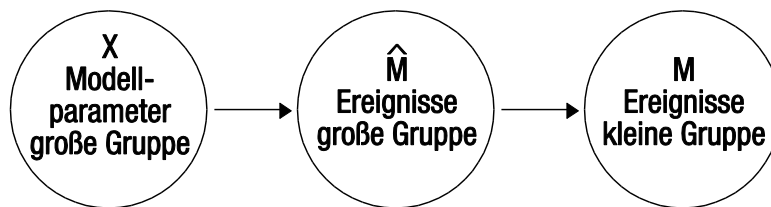


Abb. 6.1 Indirekte Abhängigkeit der Beobachtungen M in der „kleinen“ Komponentengruppe von den Modellparametern X in der „großen“ Komponentengruppe über nicht beobachtete Ereignisse $\hat{M}$ in der „großen“ Komponentengruppe

Die Abhängigkeit ist stochastischer Natur, da die beobachtete Untergruppe als zufällig angenommen wird. Besteht z. B. $\hat{M}$ aus einem (2 von 3)-Ausfall, und ist $r = 2$, so ist M mit Wahrscheinlichkeit $1/3$ ein (2 von 2)-Ausfall und mit Wahrscheinlichkeit $2/3$ ein (1 von 2)-Ausfall.

Äquivalent kann die indirekte Abhängigkeit über die Parameter der Verteilung der Ausfallereignisse in der beobachteten Komponentengruppe der Größe r (als Y bezeichnet) dargestellt werden. Dies führt zu einfacheren Ausdrücken, da die Abhängigkeit nicht stochastisch, sondern deterministisch ist: Y lässt sich eindeutig aus X berechnen, weil die Parameter gerade die bedingten Wahrscheinlichkeiten eines (k von r)-Ausfalls bzw. eines (k von $\hat{r}$)-Ausfalls sind. Auf diese Weise können die Ergebnisse einfacher analytisch dargestellt werden. Deshalb wird im Folgenden so vorgegangen.

Die betrachtete Abhängigkeit ist in Abb. 6.2 dargestellt: Die Beobachtungen M in der „kleinen“ Komponentengruppe sind von den bedingten Wahrscheinlichkeiten X in der „großen“ Komponentengruppe über die bedingten Wahrscheinlichkeiten Y in der „kleinen“ Komponentengruppe abhängig.

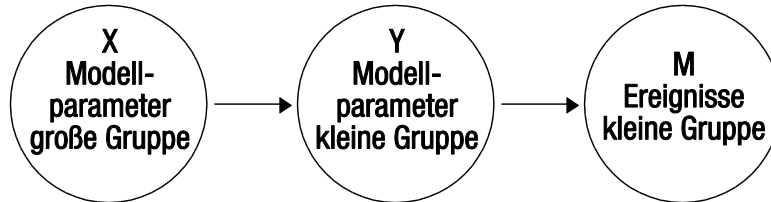


Abb. 6.2 Betrachtete indirekte Abhängigkeit der Beobachtungen M in der „kleinen“ Komponentengruppe von den Modellparametern X in der „großen“ Komponentengruppe über die Modellparameter Y in der „kleinen“ Komponentengruppe

Wie in Kapitel 3 bereits diskutiert, genügen die Beobachtungen M einer Multinomialverteilung

$$\begin{aligned}
 p(M|Y, z) = & \frac{z!}{\prod_{k=1}^r (m_{k\setminus r})! (z - \sum_{k=1}^r m_{k\setminus r})!} \\
 & \times \left(1 - \sum_{k=1}^r y_{k\setminus r}\right)^{z - \sum_{k=1}^r m_{k\setminus r}} \prod_{k=1}^r (y_{k\setminus r})^{m_{k\setminus r}}
 \end{aligned} \tag{6.5}$$

Hierbei bezeichnet z die Gesamtzahl der beobachteten Ereignisse.

Y lässt sich aus X mittels kombinatorischer Überlegungen bestimmen. Am Beispiel von $r = 2$, $\hat{r} = 3$ wird im Folgenden die Herleitung demonstriert:

Wenn 0 von 3 Komponenten ausfallen, so werden in einer Untergruppe der Größe 2 sicher 0 Ausfälle beobachtet. Wenn 1 von 3 Komponenten ausfällt, so werden in einer Untergruppe der Größe 2 keine Ausfälle beobachtet, wenn die ausgefallene Komponente nicht der Untergruppe der Größe 2 angehört. Dies ist mit Wahrscheinlichkeit $1/3$ der Fall. Andernfalls wird ein Ausfall beobachtet. Wenn 2 von 3 Komponenten ausfallen, so werden in einer Untergruppe der Größe 2 zwei Ausfälle beobachtet, wenn die Untergruppe der Größe 2 identisch mit der Menge der ausgefallenen Komponenten ist.

Dies ist mit Wahrscheinlichkeit $1/3$ der Fall. Andernfalls wird genau ein Ausfall beobachtet. Bei einem Ausfall von 3 von 3 Komponenten ist es sicher, dass zwei Ausfälle beobachtet werden. Dies ist in Abb. 6.3 dargestellt.

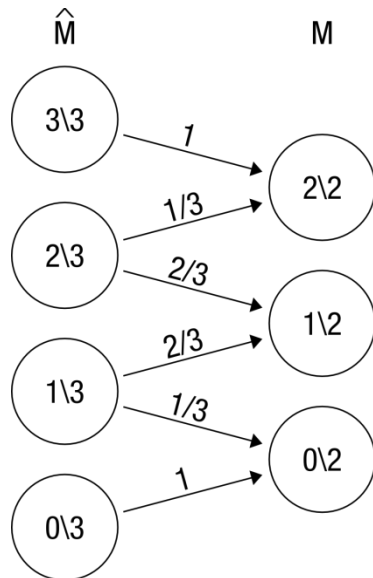


Abb. 6.3 Abhängigkeit der Beobachtungen M in der „kleinen“ Komponentengruppe von den Beobachtungen $\hat{M}$ in der „großen“ Komponentengruppe

Da die Wahrscheinlichkeit, dass k von 3 Komponenten ausfallen, $x_{k\backslash 3}$ ist, gilt:

$$y_{0\backslash 2} = x_{0\backslash 3} + \frac{1}{3} x_{1\backslash 3}$$

$$y_{1\backslash 2} = \frac{2}{3} x_{1\backslash 3} + \frac{2}{3} x_{2\backslash 3} \quad (6.6)$$

$$y_{2\backslash 2} = \frac{1}{3} x_{2\backslash 3} + x_{3\backslash 3}$$

Allgemein lässt sich schreiben

$$Y = SX \quad (6.7)$$

wobei S eine stochastische $r \times \hat{r}$ -Matrix ist, die wie oben beschrieben durch die Kombinatorik bestimmt ist und $Y = (y_{0\backslash 2}, y_{1\backslash 2}, y_{2\backslash 2})$ und $X = (x_{0\backslash 3}, x_{1\backslash 3}, x_{2\backslash 3}, x_{3\backslash 3})$. Hier gilt:

$$S = \begin{pmatrix} 1 & \frac{1}{3} & 0 & 0 \\ 0 & \frac{2}{3} & \frac{2}{3} & 0 \\ 0 & 0 & \frac{1}{3} & 1 \end{pmatrix} \quad (6.8)$$

In Anhang A sind die Elemente für weitere Kombinationen von r und $\hat{r}$ angegeben.

Damit lässt sich die bedingte Wahrscheinlichkeit von M gegeben X (und z) bestimmen

$$p(M|X, z) = \frac{z!}{(z - \sum_{k=1}^r m_{k \setminus r})! \prod_{k=1}^r (m_{k \setminus r})!} \times \left(1 - \sum_{k=1}^r (SX)_{k \setminus r}\right)^{z - \sum_{k=1}^r m_{k \setminus r}} \prod_{k=1}^r ((SX)_{k \setminus r})^{m_{k \setminus r}} \quad (6.9)$$

Im obigen Beispiel erhält man

$$p(M|X, z) = \frac{z!}{m_{0 \setminus 2}! m_{1 \setminus 2}! m_{2 \setminus 2}!} \left(x_{0 \setminus 3} + \frac{1}{3} x_{1 \setminus 3}\right)^{m_{0 \setminus 2}} \quad (6.10)$$

$$\left(\frac{2}{3} x_{1 \setminus 3} + \frac{2}{3} x_{2 \setminus 3}\right)^{m_{1 \setminus 2}} \left(\frac{1}{3} x_{2 \setminus 3} + x_{3 \setminus 3}\right)^{m_{2 \setminus 2}}$$

Nun kann der Satz von Bayes angewandt werden, um die gesuchte Größe $p(X|M, z)$ zu bestimmen. Er lautet:

$$p(X|M, z) = \frac{p(M|X, z)p(X)}{p(M)} \quad (6.11)$$

Dabei ist die Normierungskonstante $p(M)$ gegeben durch

$$p(M) = \int_{\text{Wertebereich von } X} p(M|\hat{X}, z) p(\hat{X}) d\hat{X} \quad (6.12)$$

Die a priori-Verteilung $p(X)$ (a priori in Bezug auf Beobachtung der M , d. h. der Ereignisse in den kleineren Komponentengruppen der Größe r) kann Informationen über etwaige Beobachtungen in Komponentengruppen der Größe $\hat{r}$ enthalten.

Wenn man wieder, wie oben diskutiert, das Verfahren von Jeffreys zur Bestimmung einer a priori-Verteilung wählt, so ist $p(X)$ wiederum eine Dirichletverteilung:

$$p(X) = p(x_{1 \setminus \hat{r}}, x_{2 \setminus \hat{r}}, \dots, x_{\hat{r} \setminus \hat{r}}) = \frac{1}{B(a_{0 \setminus \hat{r}}, a_{2 \setminus \hat{r}}, \dots, a_{\hat{r} \setminus \hat{r}})} \prod_{i=0}^{\hat{r}} (x_{i \setminus \hat{r}})^{a_{i \setminus \hat{r}} - 1} \quad (6.13)$$

Der Normierungsfaktor $B(a_{0 \setminus \hat{r}}, a_{2 \setminus \hat{r}}, \dots, a_{\hat{r} \setminus \hat{r}})$ ist in Gleichung (4.7) definiert.

Für die Parameter $a_{k \setminus \hat{r}}$ gilt:

$$a_{k \setminus \hat{r}} = \frac{1}{2} + n_{k \setminus \hat{r}} \quad (6.14)$$

wobei $n_{k \setminus \hat{r}}$ die Beobachtungen in den Komponentengruppen der Größe $\hat{r}$ sind. Liegen keine Beobachtungen vor, so gilt $a_{k \setminus \hat{r}} = \frac{1}{2}$. Damit ist $p(X)$ nichtinformativ.

Für den mathematischen Inhalt und damit für das Ergebnis ist es irrelevant, ob man aus den Beobachtungen in der Komponentengruppe der Größe $\hat{r}$ einen informativen a priori bestimmt, aus dem man mit den Beobachtungen in der Komponentengruppe der Größe r eine a posteriori-Verteilung berechnet, oder in umgekehrter Reihenfolge aus den Beobachtungen in der Komponentengruppe der Größe r eine informative a priori-Verteilung bestimmt und dann mit den Beobachtungen in der Komponentengruppe der Größe $\hat{r}$ eine a posteriori-Verteilung berechnet. Hier wird die erste Darstellung gewählt, da davon ausgegangen wird, dass in der Komponentengruppe der Größe $\hat{r}$ nur wenig Betriebserfahrung vorliegt, während die Hauptinformationsquelle die Beobachtungen in der Komponentengruppe der Größe r sind.

Somit ergibt sich:

$$\begin{aligned}
p(X|M, z) &\propto p(M|X, z)p(X) \\
&\propto \prod_{i=0}^{\hat{r}} (x_{i \setminus \hat{r}})^{n_{k \setminus \hat{r}} - \frac{1}{2}} \\
&\quad \left(1 - \sum_{k=1}^r (SX)_{k \setminus r}\right)^{z - \sum_{k=1}^r m_{k \setminus r}} \prod_{k=1}^r ((SX)_{k \setminus r})^{m_{k \setminus r}}
\end{aligned} \tag{6.15}$$

Für das oben erwähnte Beispiel $r = 2, \hat{r} = 3$ ergibt sich

$$\begin{aligned}
p(X|M, z) &\propto (x_{0 \setminus 3})^{n_{0 \setminus 3} - \frac{1}{2}} (x_{1 \setminus 3})^{n_{1 \setminus 3} - \frac{1}{2}} (x_{2 \setminus 3})^{n_{2 \setminus 3} - \frac{1}{2}} (x_{3 \setminus 3})^{n_{3 \setminus 3} - \frac{1}{2}} \left(x_{0 \setminus 3} \right. \\
&\quad \left. + \frac{1}{3} x_{1 \setminus 3}\right)^{m_{0 \setminus 2}} \left(\frac{2}{3} x_{1 \setminus 3} + \frac{2}{3} x_{2 \setminus 3}\right)^{m_{1 \setminus 2}} \left(\frac{1}{3} x_{2 \setminus 3} + x_{3 \setminus 3}\right)^{m_{2 \setminus 2}}
\end{aligned} \tag{6.16}$$

Betrachten wir nun den Fall, dass nur (2 von 2)-Ereignisse beobachtet wurden, d. h.:

$\forall_k n_{k \setminus 3} = 0, m_{0 \setminus 2} = m_{1 \setminus 2} = 0, m_{2 \setminus 2} = z \neq 0$. Damit vereinfacht sich (6.16) zu

$$\begin{aligned}
p(X|M, z) &\propto (x_{0 \setminus 3})^{-\frac{1}{2}} (x_{1 \setminus 3})^{-\frac{1}{2}} (x_{2 \setminus 3})^{-\frac{1}{2}} (x_{3 \setminus 3})^{-\frac{1}{2}} \left(\frac{1}{3} x_{2 \setminus 3} + x_{3 \setminus 3}\right)^z \\
&\propto (x_{0 \setminus 3})^{-\frac{1}{2}} (x_{1 \setminus 3})^{-\frac{1}{2}} (x_{2 \setminus 3})^{-\frac{1}{2}} (x_{3 \setminus 3})^{-\frac{1}{2}} \left(\frac{1}{4} x_{2 \setminus 3} + \frac{3}{4} x_{3 \setminus 3}\right)^z
\end{aligned} \tag{6.17}$$

Somit hat $p(X|M, z)$ die Form einer Mischverteilung aus $z + 1$ Verteilungen. Dieser Ausdruck weist eine Symmetrie bzgl. $x_{0 \setminus 3}$ und $x_{1 \setminus 3}$ auf. Daraus folgt, dass die Marginalverteilungen und die Erwartungswerte von $x_{0 \setminus 3}$ und $x_{1 \setminus 3}$ identisch sein müssen.

Für $z = 1$ folgt aus Gleichung (6.17):

$$\begin{aligned}
p(X|M, z) &\propto \frac{1}{4} (x_{0 \setminus 3})^{-\frac{1}{2}} (x_{1 \setminus 3})^{-\frac{1}{2}} (x_{2 \setminus 3})^{\frac{1}{2}} (x_{3 \setminus 3})^{-\frac{1}{2}} \\
&\quad + \frac{3}{4} (x_{0 \setminus 3})^{-\frac{1}{2}} (x_{1 \setminus 3})^{-\frac{1}{2}} (x_{2 \setminus 3})^{-\frac{1}{2}} (x_{3 \setminus 3})^{\frac{1}{2}}
\end{aligned} \tag{6.18}$$

Damit hat $p(X|M, z)$ die Form einer gewichtete Summe zweier a posteriori-Verteilungen, die auftreten, wenn

- ein (2 von 3)-Ereignis beobachtet wird (Gewicht $\frac{1}{4}$) bzw.
- ein (3 von 3)-Ereignis beobachtet wird (Gewicht $\frac{3}{4}$).

Gleichung (6.18) lässt sich entnehmen, dass die Gewichte $\frac{1}{4}$ bzw. $\frac{3}{4}$ sind, da die einzelnen Summanden dieselbe funktionale Form haben und somit der Normierungskonstanten der Verteilungen, wenn ein (2 von 3)-Ereignis beobachtet wird bzw. wenn ein (3 von 3)-Ereignis beobachtet wird, identisch sind. Dies ist bei allen weiteren Fällen $z > 1$ nicht mehr der Fall.

Für $z = 2$ folgt aus Gleichung (6.17)

$$\begin{aligned}
 p(X|M, z) &\propto (x_{0\setminus 3})^{-\frac{1}{2}}(x_{1\setminus 3})^{-\frac{1}{2}}(x_{2\setminus 3})^{-\frac{1}{2}}(x_{3\setminus 3})^{-\frac{1}{2}} \left(\frac{1}{3}x_{2\setminus 3} + x_{3\setminus 3} \right)^2 \\
 &\propto (x_{0\setminus 3})^{-\frac{1}{2}}(x_{1\setminus 3})^{-\frac{1}{2}}(x_{2\setminus 3})^{-\frac{1}{2}}(x_{3\setminus 3})^{-\frac{1}{2}} \left(\frac{1}{9}x_{2\setminus 3}^2 \right. \\
 &\quad \left. + \frac{1}{3}x_{2\setminus 3}x_{3\setminus 3} + x_{3\setminus 3}^2 \right) \\
 &\propto \frac{1}{16} (x_{0\setminus 3})^{-\frac{1}{2}}(x_{1\setminus 3})^{-\frac{1}{2}}(x_{2\setminus 3})^{\frac{3}{2}}(x_{3\setminus 3})^{-\frac{1}{2}} \\
 &\quad + \frac{6}{16}(x_{0\setminus 3})^{-\frac{1}{2}}(x_{1\setminus 3})^{-\frac{1}{2}}(x_{2\setminus 3})^{\frac{1}{2}}(x_{3\setminus 3})^{\frac{1}{2}} \\
 &\quad + \frac{9}{16}(x_{0\setminus 3})^{-\frac{1}{2}}(x_{1\setminus 3})^{-\frac{1}{2}}(x_{2\setminus 3})^{-\frac{1}{2}}(x_{3\setminus 3})^{\frac{3}{2}}
 \end{aligned} \tag{6.19}$$

Damit hat $p(X|M, z)$ die Form einer gewichteten Summe dreier a posteriori-Verteilungen (Mischverteilung), die auftreten, wenn

- zwei (2 von 3)-Ereignisse beobachtet werden bzw.
- ein (2 von 3)-Ereignis und ein (3 von 3)-Ereignis beobachtet werden bzw.
- zwei (3 von 3)-Ereignisse beobachtet werden.

Wie oben erwähnt folgt hier nicht, dass die Beiträge 1/16 bzw. 6/16 seien, da die Normierungsfaktoren berücksichtigt werden müssen. Diese sind im Gegensatz zu oben nicht identisch, da die einzelnen Summanden nicht dieselbe funktionale Form haben.

Für die a posteriori-Verteilung gilt, wenn zwei (2 von 3)-Ereignisse beobachtet werden:

$$\begin{aligned}
 p(X|\text{zwei (2 von 3) – Ereignisse}) &= \\
 &= \frac{1}{\int_{\text{Wertebereich von } X} (x_{0\setminus 3})^{-\frac{1}{2}} (x_{1\setminus 3})^{-\frac{1}{2}} (x_{2\setminus 3})^{\frac{3}{2}} (x_{3\setminus 3})^{-\frac{1}{2}} dX} \\
 &\times (x_{0\setminus 3})^{-\frac{1}{2}} (x_{1\setminus 3})^{-\frac{1}{2}} (x_{2\setminus 3})^{\frac{3}{2}} (x_{3\setminus 3})^{-\frac{1}{2}} = \frac{8}{\pi^2} (x_{0\setminus 3})^{-\frac{1}{2}} (x_{1\setminus 3})^{-\frac{1}{2}} (x_{2\setminus 3})^{\frac{3}{2}} (x_{3\setminus 3})^{-\frac{1}{2}}
 \end{aligned} \tag{6.20}$$

Für die a posteriori-Verteilung gilt, wenn zwei (3 von 3)-Ereignisse beobachtet werden:

$$\begin{aligned}
 p(X|\text{zwei (3 von 3) – Ereignisse}) &= \\
 &= \frac{1}{\int_{\text{Wertebereich von } X} (x_{0\setminus 3})^{-\frac{1}{2}} (x_{1\setminus 3})^{-\frac{1}{2}} (x_{2\setminus 3})^{-\frac{1}{2}} (x_{3\setminus 3})^{\frac{3}{2}} dX} \\
 &\times (x_{0\setminus 3})^{-\frac{1}{2}} (x_{1\setminus 3})^{-\frac{1}{2}} (x_{2\setminus 3})^{-\frac{1}{2}} (x_{3\setminus 3})^{\frac{3}{2}} = \frac{8}{\pi^2} (x_{0\setminus 3})^{-\frac{1}{2}} (x_{1\setminus 3})^{-\frac{1}{2}} (x_{2\setminus 3})^{-\frac{1}{2}} (x_{3\setminus 3})^{\frac{3}{2}}
 \end{aligned} \tag{6.21}$$

Demgegenüber gilt für die a posteriori-Verteilung, wenn ein (2 von 3)-Ereignis und ein (3 von 3)-Ereignis beobachtet werden:

$$\begin{aligned}
 p(X|\text{ein (2 von 3) – Ereignis und ein (3 von 3) – Ereignis}) &= \\
 &= \frac{1}{\int_{\text{Wertebereich von } X} (x_{0\setminus 3})^{-\frac{1}{2}} (x_{1\setminus 3})^{-\frac{1}{2}} (x_{2\setminus 3})^{\frac{1}{2}} (x_{3\setminus 3})^{\frac{1}{2}} dX} \\
 &\times (x_{0\setminus 3})^{-\frac{1}{2}} (x_{1\setminus 3})^{-\frac{1}{2}} (x_{2\setminus 3})^{\frac{1}{2}} (x_{3\setminus 3})^{\frac{1}{2}} = \frac{24}{\pi^2} (x_{0\setminus 3})^{-\frac{1}{2}} (x_{1\setminus 3})^{-\frac{1}{2}} (x_{2\setminus 3})^{\frac{1}{2}} (x_{3\setminus 3})^{\frac{1}{2}}
 \end{aligned} \tag{6.22}$$

Damit folgt aus Gleichung (6.19):

$$\begin{aligned}
& p(X|M, z) \\
& \propto \frac{1}{12} p(X|\text{zwei (2 von 3) – Ereignisse}) \\
& + \frac{1}{6} p(X|\text{ein (2 von 3) – Ereignis und ein (3 von 3) – Ereignis}) \\
& + \frac{3}{4} p(X|\text{zwei (3 von 3) – Ereignisse})
\end{aligned} \tag{6.23}$$

Somit ist die a posteriori-Verteilung eine Mischverteilung derjenigen a posteriori-Verteilungen, die auftreten, wenn

- zwei (2 von 3)-Ereignisse beobachtet werden, mit Gewicht $1/12$,
- ein (2 von 3)-Ereignis und ein (3 von 3)-Ereignis beobachtet werden, mit Gewicht $1/6$,
- zwei (3 von 3)-Ereignisse beobachtet werden, mit Gewicht $3/4$.

Analog ergibt sich für größere z die Verteilung $p(X|M, z)$ als gewichtete Summe von $z + 1$ a posteriori-Verteilungen, die auftreten, wenn

- z (2 von 3)-Ereignisse beobachtet werden,
- $z - 1$ (2 von 3)-Ereignis und ein (3 von 3)-Ereignis beobachtet werden,
- ...,
- z (3 von 3)-Ereignisse beobachtet werden.

Hierbei wird das relative Gewicht der a posteriori-Verteilungen, die einer hohen Zahl von (3 von 3)-Ereignissen entsprechen, immer größer. Wie es auch der oben dargestellten Erwartung entspricht, wächst der Erwartungswert $\langle x_{3 \setminus 3} \rangle$ der Wahrscheinlichkeit, dass ein Ereignis ein (3 von 3)-Ausfall ist, stetig und konvergiert gegen 1:

$$\langle x_{3 \setminus 3} \rangle \rightarrow 1 \quad \text{für} \quad z \rightarrow \infty \tag{6.24}$$

Entsprechend konvergieren die Wahrscheinlichkeiten, dass ein Ausfall ein (i von 3)-Ausfall ist, für $i < 3$ gegen 0:

$$\forall_{i < 3} \langle x_{i \setminus 3} \rangle \rightarrow 0 \quad \text{für} \quad z \rightarrow \infty \quad (6.25)$$

Die Erwartungswerte können berechnet werden als

$$\langle x_{i \setminus 3} \rangle = \frac{\int_0^1 \int_0^{1-x_{1 \setminus 3}} \int_0^{1-x_{1 \setminus 3}-x_{2 \setminus 3}} x_{i \setminus 3} u(x_{1 \setminus 3}, x_{2 \setminus 3}, x_{3 \setminus 3}, z) dx_{3 \setminus 3} dx_{2 \setminus 3} dx_{1 \setminus 3}}{\int_0^1 \int_0^{1-x_{1 \setminus 3}} \int_0^{1-x_{1 \setminus 3}-x_{2 \setminus 3}} u(x_{1 \setminus 3}, x_{2 \setminus 3}, x_{3 \setminus 3}, z) dx_{3 \setminus 3} dx_{2 \setminus 3} dx_{1 \setminus 3}} \quad (6.26)$$

wobei $u(x_{1 \setminus 3}, x_{2 \setminus 3}, x_{3 \setminus 3}, z)$ nach Gleichung (6.17):

$$u(x_{1 \setminus 3}, x_{2 \setminus 3}, x_{3 \setminus 3}, z) = (1 - x_{1 \setminus 3} - x_{2 \setminus 3} - x_{3 \setminus 3})^{-\frac{1}{2}} \times (x_{1 \setminus 3})^{-\frac{1}{2}} (x_{2 \setminus 3})^{-\frac{1}{2}} (x_{3 \setminus 3})^{-\frac{1}{2}} \left(\frac{1}{3} x_{2 \setminus 3} + x_{3 \setminus 3} \right)^z \quad (6.27)$$

ist. Dabei wurde verwendet, dass $x_{0 \setminus 3} + x_{1 \setminus 3} + x_{2 \setminus 3} + x_{3 \setminus 3} = 1$ ist.

Diese Erwartungswerte sind in Abb. 6.4 dargestellt. Die Werte wurden durch analytische Auswertung von Gleichung (6.25) mittels von Mathematica /WOL 12/ ermittelt.

Es ist erkennbar, dass $\langle x_{3 \setminus 3} \rangle$ gegen 1 konvergiert, während $\langle x_{0 \setminus 3} \rangle$, $\langle x_{1 \setminus 3} \rangle$, und $\langle x_{2 \setminus 3} \rangle$ gegen 0 konvergieren. Hierbei gilt $\langle x_{2 \setminus 3} \rangle > \langle x_{0 \setminus 3} \rangle = \langle x_{1 \setminus 3} \rangle$, da bei Beobachtung von (2 von 2)-Ereignissen sicher ist, dass keine (0 von 3)-Ereignisse oder (1 von 3)-Ereignisse vorliegen, während (2 von 3)-Ereignisse möglich sind.

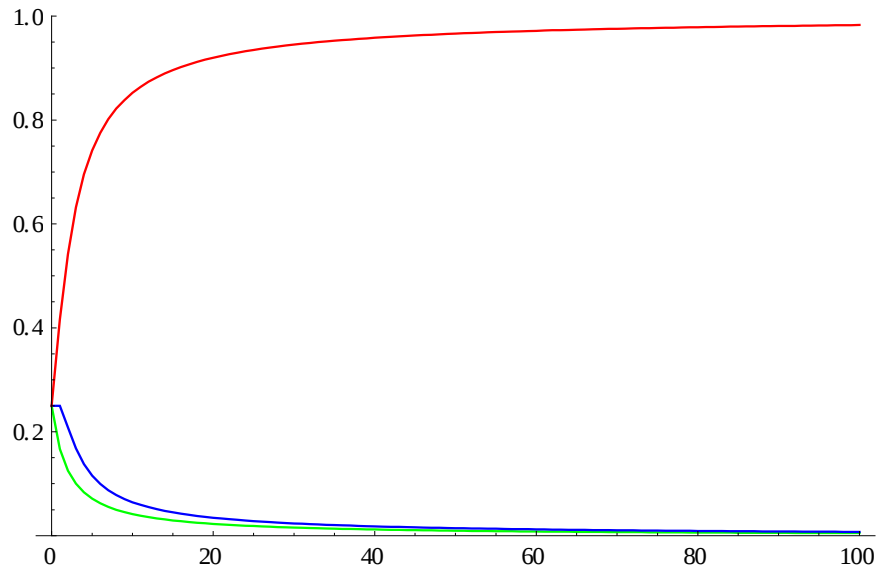


Abb. 6.4 Erwartungswerte $\langle x_{0 \setminus 3} \rangle = \langle x_{1 \setminus 3} \rangle$ (grün), $\langle x_{2 \setminus 3} \rangle$ (blau) und $\langle x_{3 \setminus 3} \rangle$ (rot) in Abhängigkeit von z für den Fall, dass ausschließlich (2 von 2)-Ereignisse beobachtet wurden

Wie schon aus der Konvergenz der Erwartungswerte gegen 1 bzw. 0 folgt, wird die Breite der Verteilungen für steigende Anzahl von Beobachtungen z immer kleiner. Dies ist explizit in Abb. 6.5 dargestellt, wo die Marginalverteilungen von $x_{3 \setminus 3}$ nach Beobachtung einer verschiedenen Anzahl von (2 von 2)-Ereignissen dargestellt sind.

Für $z = 0$ ist die Marginalverteilung von $x_{3 \setminus 3}$ eine Betaverteilung mit den Parametern $1/2$ und $3/2$, für $z > 0$ eine Mischverteilung aus Betaverteilung mit den Parametern $\frac{1}{2} + i$ und $\frac{3}{2} + z$, $i = 0, \dots, z$. Für $z = 0$ und $z = 1$ sind die Marginalverteilungsdichten monoton fallend, Für $z > 1$ haben sie einen mit größer werdenden z zunehmend höheren und schmalen Peak bei einem $x_{3 \setminus 3} > 0$. Das Gewicht des Peaks bei $x_{3 \setminus 3} = 0$ wird immer kleiner, da das Gewicht der Verteilung mit $i = 0$ bei wachsendem z abnimmt. Wie oben dargestellt, konvergiert der Erwartungswert von $x_{3 \setminus 3}$ gegen 1.

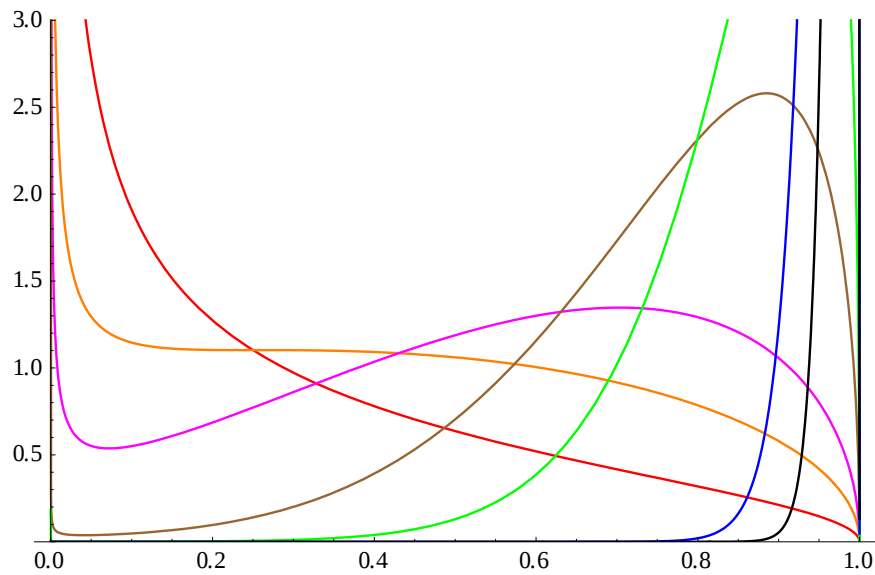


Abb. 6.5 Marginalverteilungsdichten von $x_{3\setminus 3}$ nach Beobachtung von $z = 0$ (rot), 1 (orange), 2 (magenta), 5 (braun), 10 (grün), 50 (blau) bzw. 100 (schwarz) (2 von 2)-Ereignissen

Der hier diskutierte Fall, dass ausschließlich (2 von 2)-Ereignisse beobachtet werden, ist ein Spezialfall. Aus den Beobachtungen in Komponentengruppe der Größe 2 kann für allgemeine Beobachtungen im Gegensatz zu diesem Spezialfall nicht eindeutig auf die Wahrscheinlichkeit der Ausfallkombinationen in einer Gruppe der Größe 3 geschlossen werden. Dies wird im nachfolgenden Abschnitt diskutiert.

Allgemeiner Fall

Der oben entwickelte Bayes'sche Formalismus ist auch im allgemeinen Fall anwendbar, bei dem alle möglichen Ausfallkombinationen beobachtet werden. In diesem Fall kann man im Allgemeinen auch im Grenzfall unendlich vieler Beobachtungen nicht eindeutig auf die Wahrscheinlichkeit der Ausfallkombinationen in einer Gruppe der Größe 3 aus Beobachtungen in Komponentengruppe der Größe 2 schließen. Dies ist anschaulich klar, da die Wahrscheinlichkeiten in der Komponentengruppe der Größe 3 durch drei unabhängige Variable (z. B. $x_{1\setminus 3}$, $x_{2\setminus 3}$ und $x_{3\setminus 3}$; $x_{0\setminus 3}$ ist durch $\sum_{i=0}^3 x_{i\setminus 3} = 1$ festgelegt.) bestimmt sind, während sie in der Komponentengruppe der Größe 2 durch zwei unabhängige Variable (z. B. $y_{1\setminus 2}$ und $y_{2\setminus 2}$; $y_{0\setminus 2}$ ist durch $\sum_{i=0}^2 y_{i\setminus 2} = 1$ festgelegt.) bestimmt sind. Auch bei Beobachtung beliebig vieler Ereignisse in der Komponentengruppe der Größe 2, die eine beliebig genaue Bestimmung von $y_{1\setminus 2}$ und $y_{2\setminus 2}$ ermögli-

chen, können $x_{1\setminus 3}$, $x_{2\setminus 3}$ und $x_{3\setminus 3}$ im Allgemeinen nicht eindeutig bestimmt werden, wie man an der Beziehung der $y_{i\setminus 2}$ zu den $x_{i\setminus 3}$ (Gleichung s.2) sieht:

$$\begin{aligned} y_{0\setminus 2} &= x_{0\setminus 3} + \frac{1}{3} x_{1\setminus 3} \\ y_{1\setminus 2} &= \frac{2}{3} x_{1\setminus 3} + \frac{2}{3} x_{2\setminus 3} \\ y_{2\setminus 2} &= \frac{1}{3} x_{2\setminus 3} + x_{3\setminus 3} \end{aligned} \tag{6.28}$$

Für den oben diskutierten Spezialfall, der im Grenzfall $z \rightarrow \infty$ $y_{0\setminus 2} = y_{1\setminus 2} = 0$ und $y_{2\setminus 2} = 1$ entspricht, folgt eindeutig $x_{0\setminus 3} = x_{1\setminus 3} = x_{2\setminus 3} = 0$ und $x_{3\setminus 3} = 1$. Für $y_{1\setminus 2} = 0$ lässt sich stets eine eindeutige Lösung finden, da dann $x_{1\setminus 3} = x_{2\setminus 3} = 0$ folgt, was zu $y_{0\setminus 2} = x_{0\setminus 3}$ und $y_{2\setminus 2} = x_{3\setminus 3}$ führt. Für beliebige $x_{i\setminus 3}$ gilt dies nicht.

Im allgemeinen Fall lassen sich mögliche Wertebereiche der $x_{i\setminus 3}$ durch folgende Überlegungen bestimmen: Gesucht sind die Lösungen der linearen Gleichung

$$Y = SX \tag{6.29}$$

mit

$$S = \begin{pmatrix} 1 & \frac{1}{3} & 0 & 0 \\ 0 & \frac{2}{3} & \frac{2}{3} & 0 \\ 0 & 0 & \frac{1}{3} & 1 \end{pmatrix} \tag{6.30}$$

Der Rang von S ist 3. Der Kern von S ist gegeben durch

$$\mathcal{K}(S) = \{V \in \mathbb{R}^4 | V = \psi (-1, \quad 3, \quad -3, \quad 1), \psi \in \mathbb{R}\} \tag{6.31}$$

Wenn man Y schreibt als

$$Y = \begin{pmatrix} 1 - y_{1\setminus 2} - y_{2\setminus 2} \\ y_{1\setminus 2} \\ y_{2\setminus 2} \end{pmatrix} \tag{6.32}$$

wodurch sichergestellt ist, dass die Summe der Elemente 1 ist, so hat die lineare Gleichung (6.29) die allgemeine Lösung

$$\begin{aligned} \mathcal{L}(y_{1\setminus 2}, y_{2\setminus 2}) = \{X \in \mathbb{R}^4 \mid X = & \left(1 - \frac{3y_{1\setminus 2}}{2}, \frac{3y_{1\setminus 2}}{2} - 3y_{2\setminus 2}, 3y_{2\setminus 2}, 0\right) \\ & + \psi (-1, 3, -3, 1), \psi \in \mathbb{R}\} \end{aligned} \quad (6.33)$$

da $\left(1 - \frac{3y_{1\setminus 2}}{2}, \frac{3y_{1\setminus 2}}{2} - 3y_{2\setminus 2}, 3y_{2\setminus 2}, 0\right)$ Gleichung (6.29) löst. Für die hier betrachtete Anwendung zulässige Lösungen unterliegen der zusätzlichen Bedingung, dass alle Komponenten von X aus dem Intervall $[0,1]$ sind, da sie Wahrscheinlichkeiten sind. Die Bedingung, dass die Summe der Elemente von X eins ergibt, ist durch die Konstruktion von Y und die Stochastizität von S automatisch erfüllt.

Aus (6.33) lassen sich für ψ folgende Bedingungen ablesen¹⁴:

- i. $\psi \geq 0$ (vierte Spalte)
- ii. $\psi \leq 1$ (vierte Spalte)
- iii. $\psi \leq y_{2\setminus 2}$ (dritte Spalte)
- iv. $\psi \geq y_{2\setminus 2} - \frac{1}{2}y_{1\setminus 2}$ (zweite Spalte)
- v. $\psi \leq \frac{1}{3} + y_{2\setminus 2} - \frac{1}{2}y_{1\setminus 2}$ (zweite Spalte)
- vi. $\psi \leq 1 - \frac{3}{2}y_{1\setminus 2}$ (erste Spalte)

Die Liste enthält redundante Bedingungen. Z. B. umfasst Bedingung iii. Bedingung ii.

Aus den Bedingungen folgt

$$\text{Max}\{0, y_{2\setminus 2} - \frac{1}{2}y_{1\setminus 2}\} \leq \psi \leq \text{Min}\{y_{2\setminus 2}, \frac{1}{3} + y_{2\setminus 2} - \frac{1}{2}y_{1\setminus 2}, 1 - \frac{3}{2}y_{1\setminus 2}\} \quad (6.34)$$

Für $y_{2\setminus 2} = 1, y_{1\setminus 2} = y_{0\setminus 2} = 0$, das dem oben diskutierten Beispiel zugrunde liegt, folgt

¹⁴ Hierbei wurden offensichtlich schwächeren Bedingungen als am Listenanfang nicht aufgenommen.

$$\text{Max}\{0,1\} \leq \psi \leq \text{Min}\{1, \frac{4}{3}, 1\} \quad (6.35)$$

Somit gilt $\psi = 1$, und damit ist $X = (0, 0, 0, 1)$ die einzige Lösung.

Für $y_{1 \setminus 2} = y_{2 \setminus 2} = 0$, $y_{0 \setminus 2} = 1$ folgt

$$\text{Max}\{0,0\} \leq \psi \leq \text{Min}\{0, \frac{1}{3}, 1\} \quad (6.36)$$

Deshalb gilt $\psi = 0$ und damit ist $X = (1, 0, 0, 0)$ die einzige Lösung.

Für $y_{0 \setminus 2} = y_{1 \setminus 2} = y_{2 \setminus 2} = 1/3$ folgt

$$\text{Max}\{0, \frac{1}{6}\} \leq \psi \leq \text{Min}\{\frac{1}{3}, \frac{1}{2}, \frac{1}{2}\} \quad (6.37)$$

Deshalb gilt $\psi \in [\frac{1}{6}, \frac{1}{3}]$.

Die daraus resultierende Wertebereiche für die $x_{i \setminus 3}$ sind:

$$x_{0 \setminus 3} \in [\frac{1}{6}, \frac{1}{3}] \quad x_{1 \setminus 3} \in [0, \frac{1}{2}] \quad x_{2 \setminus 3} \in [0, \frac{1}{2}] \quad x_{3 \setminus 3} \in [\frac{1}{6}, \frac{1}{3}] \quad (6.38)$$

Aus $Y = \{\frac{4}{10}, \frac{5}{10}, \frac{1}{10}\}$ folgt $\psi \in [0, \frac{1}{10}]$. Die daraus resultierende Wertebereiche für die $x_{i \setminus 3}$ sind:

$$x_{0 \setminus 3} \in [\frac{3}{20}, \frac{1}{4}] \quad x_{1 \setminus 3} \in [\frac{9}{20}, \frac{3}{4}] \quad x_{2 \setminus 3} \in [0, \frac{3}{10}] \quad x_{3 \setminus 3} \in [0, \frac{1}{10}] \quad (6.39)$$

Allgemein gilt, dass für diesen typischen Fall, bei dem $y_{2 \setminus 2}$ deutlich kleiner als $y_{1 \setminus 2}$ ist, $\psi \in [0, y_{2 \setminus 2}]$ ist. Demzufolge sind für $x_{3 \setminus 3}$ Werte zwischen 0 und $y_{2 \setminus 2}$ möglich, für $x_{2 \setminus 3}$ sogar zwischen 0 und $3 y_{2 \setminus 2}$. Die Unsicherheit für die in der PSA relevanten Ausfallkombinationen ist somit relativ groß.

Wie man an dem Vektor $(-1, 3, -3, 1)$, der den Kern von S aufspannt, erkennt, sind $x_{2 \setminus 3}$ und $x_{3 \setminus 3}$ antikorreliert. Das heißt, wenn $x_{3 \setminus 3}$ maximal ist, ist $x_{2 \setminus 3}$ minimal und

umgekehrt. Das zweite und dritte Element sind auch betragsmäßig größer. Das heißt, eine kleinen Variation von $x_{3\setminus 3}$ entspricht einer dreimal größere von $x_{2\setminus 3}$.

In einem Bayes'schen Verfahren werden natürlich alle Werte von ψ berücksichtigt; ihr jeweiliges Gewicht wird indirekt durch die a priori-Verteilung der Modellparameter X bestimmt.

Zunächst soll aber untersucht werden, ob die Unsicherheit bezüglich ψ durch eine konservative Festlegung von ψ berücksichtigt werden kann. Dann würde zumindest im Grenzfall einer sehr großen Anzahl von beobachteten Ereignissen in Komponentengruppen der Größe r die Abhängigkeit der Ergebnisse von dem a priori verschwinden.

In Sicherheitssystemen des Redundanzgrades 3 ist im Allgemeinen die Verfügbarkeit einer Redundante für die Systemfunktion ausreichend (sogenannte 3 * 100 %-Auslegung), so dass die Wahrscheinlichkeit des Ausfalls der Systemfunktion durch $x_{3\setminus 3}$ bestimmt wird. Allerdings kann auch die Kombination eines unabhängigen Ausfalls mit einem GVA hohe Relevanz haben; hierbei ist (neben der Wahrscheinlichkeit eines unabhängigen Ausfalls) $x_{2\setminus 3}$ maßgebend. Unübersichtlicher wird die Betrachtung bei Sicherheitssystemen mit Redundanzgrad 4: Hier ist eine so genannte 4 * 50 %-Auslegung typisch, d. h. die Systemfunktion ist auch für die ungünstigsten zu betrachtenden Bedingungen nachweisbar, wenn zwei der Redundanten verfügbar sind; allerdings ist typischerweise für einen großen Teil der Anforderungen schon eine Verfügbarkeit von einer Redundante ausreichend. Auch hier können Kombinationen mit unabhängigen Ausfällen relevant sein. Da auch im Fall des Mappings von Redundanzgrad 3 auf 4 der Vektor, der den Kern von S aufspannt, die oben gezeigte Eigenschaft hat (siehe Anhang A), die zu einer Antikorrelation von $x_{4\setminus 4}$ und $x_{3\setminus 4}$ führt, erscheint eine wohlbegründete Festlegung eines konservativen Wertes für ψ im Allgemeinen kaum möglich.

Im allgemeinen Fall wird also bei diesem Mappingverfahren, auch wenn die Anzahl der Beobachtungen gegen unendlich strebt, das Ergebnis durch die a priori-Verteilung mitbestimmt. Wie in Abschnitt 6.3 diskutiert wird, ist jedoch die Wahl eines nichtinformativen a prioris mit der dem Verfahren zugrundeliegenden Annahme, dass sich eine Komponentengruppe der Größe r verhält wie eine Untergruppe einer Gruppe der Größe $\hat{r}$, inkompatibel. Dort werden auch grundsätzliche Probleme beschrieben, die die Entwicklung eines kompatiblen a prioris als sehr schwierig erscheinen lassen. Da eine

Anwendung des in diesem Abschnitt dargestellten Ansatzes vor Entwicklung von geeigneten a priori-Verteilungen nicht sinnvoll möglich ist, wurde hier auf die weitere Ausarbeitung und formelmäßige Darstellung des allgemeinen Falles verzichtet.

6.2.3 Konservative Abwandlung des Probabilistisch-kombinatorischen Ansatzes

Bei den oben beschriebenen Ansätzen des Mapping Down werden zufällige Komponenten weggelassen, beim Mapping Up zufällige Komponenten „dupliziert“ und die möglichen Ergebnisse gemäß ihres statistischen Gewichtes für die Schätzung der GVA-Wahrscheinlichkeiten verwendet. Dem liegt, wie oben diskutiert, die Annahme zugrunde, dass sich alle Komponenten statistisch gleich verhalten und eine kleinere Gruppe von Komponenten sich gleich verhält wie eine Untergruppe einer größeren Komponentengruppe. Wenn diese Annahme nicht zutrifft, liefern die dargestellten Verfahren keine zutreffenden Ergebnisse.

Um diese Annahme zu vermeiden, können konservative Mappingalgorithmen entwickelt werden.

Konservatives Mapping Down

Statt wie oben zufällige Komponenten wegzulassen, werden die am schwächsten geschädigten Komponenten weggelassen. Bei dem in Abschnitt 6.2.2 verwendeten Beispiel, in dem in einer Komponentengruppe der Größe 3 ein Ausfall, eine starke Schädigung und eine sehr schwache Schädigung beobachtet wurden, ergibt sich dann als Ergebnis des Mappings der Schädigungsvektor (Ausfall, starke Schädigung). Hier fließt also nur die ungünstigste Möglichkeit in die weitere Berechnung ein, während beim Probabilistisch-kombinatorischen Mapping Down alle möglichen Kombinationen (Ausfall, starke Schädigung), (Ausfall, sehr schwache Schädigung) und (starke Schädigung, sehr schwache Schädigung) mit gleichem statistischen Gewicht verwendet werden. Die mit dem Mapping verknüpfte Unsicherheit wird beim konservativen Mapping Down nicht abgebildet.

Tab. 6.4 Schädigungsvektor und jeweilige Wahrscheinlichkeit bei probabilistisch-kombinatorischem und konservativem Mapping Down

Schädigungsvektor in der Komponentengruppe der Größe 2	Wahrscheinlichkeit bei probabilistisch-kombinatorischem Mapping Down (siehe Abschnitt 6.2.3)	Wahrscheinlichkeit bei konservativem Mapping Down
(Ausfall, starke Schädigung)	1/3	1
(Ausfall, sehr schwache Schädigung)	1/3	0
(starke Schädigung, sehr schwache Schädigung)	1/3	0

Konservatives Mapping Up

Analog kann ein konservatives Mapping Up entwickelt werden. Dazu wird die an stärksten geschädigte Komponente $\hat{r}$ – r -mal „dupliziert“.

Tab. 6.5 Schädigungsvektor und jeweilige Wahrscheinlichkeit bei heuristischem und bei konservativem Mapping Up

Schädigungsvektor in der Komponentengruppe der Größe 4	Wahrscheinlichkeit bei heuristischem Mapping Up (siehe Kap. 6.2.4)	Wahrscheinlichkeit bei konservativem Mapping Up
(Ausfall, Ausfall, starke Schädigung, sehr schwache Schädigung)	1/3	1
(Ausfall, starke Schädigung, starke Schädigung, sehr schwache Schädigung)	1/3	0
(Ausfall, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	1/3	0

Die mit dem Mapping verknüpfte Unsicherheit wird beim konservativen Mapping Up nicht abgebildet.

6.2.4 Heuristisches Mapping Up

Ein heuristisches Mapping Up-Verfahren kann in Analogie zum oben beschriebenen Mapping Down konstruiert werden. Dies besteht in einem „Vervielfachen“ der Komponenten. Im Fall von $\hat{r} = r + 1$ ist es eindeutig, dass eine beliebige zufällige Komponente dupliziert werden soll, da alle Komponenten statistisch äquivalent sind. Beim oben eingeführten Beispiel einer Komponentengruppe der Größe 3 mit einem Ausfall, einer starken Schädigung und einer sehr schwachen Schädigung ergibt sich das in Tab. 6.6 dargestellte Ergebnis.

Tab. 6.6 Schädigungsvektor der Größe 4 und jeweilige Wahrscheinlichkeit

Schädigungsvektor in der Komponentengruppe der Größe 4	Wahrscheinlichkeit
(Ausfall, Ausfall, starke Schädigung, sehr schwache Schädigung)	1/3
(Ausfall, starke Schädigung, starke Schädigung, sehr schwache Schädigung)	1/3
(Ausfall, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	1/3

Bei $\hat{r} > r + 1$ sind verschiedene Vorgehensweisen möglich. Diese entsprechen prinzipiell den statistischen Standardproblemen „Ziehen mit Zurücklegen“ und „Ziehen ohne Zurücklegen“.

Heuristisches Mapping Up: „Ziehen mit Zurücklegen“

Bei diesem Ansatz werden die zu duplizierenden Komponenten zufällig aus den vorhandenen gezogen, wobei Mehrfachziehungen derselben Komponente möglich sind. Beim oben eingeführten Beispiel einer Komponentengruppe der Größe 3 mit einem Ausfall, einer starken Schädigung und einer sehr schwachen Schädigung ergibt sich für ein Mapping auf die Zielkomponentengruppengröße 5 das in Tab. 6.7 dargestellte Ergebnis.

Tab. 6.7 Schädigungsvektor der Größe 5 und jeweilige Wahrscheinlichkeit

Schädigungsvektor in der Komponentengruppe der Größe 5	Wahrscheinlichkeit
(Ausfall, Ausfall, Ausfall, starke Schädigung, sehr schwache Schädigung)	1/9
(Ausfall, Ausfall, starke Schädigung, starke Schädigung, sehr schwache Schädigung)	2/9
(Ausfall, Ausfall, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	2/9
(Ausfall, starke Schädigung, starke Schädigung, starke Schädigung, sehr schwache Schädigung)	1/9
(Ausfall, starke Schädigung, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	2/9
(Ausfall, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	1/9

Insgesamt sind maximal $r^{\hat{r}-r}$ mögliche Kombinationen zu betrachten. Wenn mehrere der r Komponenten denselben Schädigungswert aufweisen, reduziert sich die Anzahl entsprechend.

Diese Vorgehensweise ist prinzipiell unabhängig von der Zielkomponentengruppengröße $\hat{r}$ möglich. Dabei werden wie oben erkennbar auch Schädigungsvektoren berücksichtigt, bei denen dieselbe Komponente $\hat{r} - r$ mal dupliziert wird, allerdings mit geringer Wahrscheinlichkeit $1/r^{\hat{r}-r}$. Beispielsweise wird bei Übertragung des oben beschriebenen Ereignisses auf eine Komponentengruppe der Größe 7 ein Schädigungsvektor mit maximal vier Ausfällen berücksichtigt. Dieser Schädigungsvektor hat die Wahrscheinlichkeit 1/81.

Heuristisches Mapping Up: „Ziehen ohne Zurücklegen“

Bei diesem Ansatz werden die zu duplizierenden Komponenten zufällig aus den vorhandenen gezogen, wobei Mehrfachziehungen derselben Komponente nicht möglich sind. Deshalb wird zunächst $\hat{r} \leq 2r$ vorausgesetzt. Beim oben eingeführten Beispielereignis in einer Komponentengruppe der Größe 3 mit einem Ausfall, einer starken Schädigung und einer sehr schwachen Schädigung ergibt sich für ein Mapping auf die Zielkomponentengruppengröße 5 das in Tab. 6.8 dargestellte Ergebnis.

Tab. 6.8 Schädigungsvektor der Größe 5 und jeweilige Wahrscheinlichkeit

Schädigungsvektor in der Komponentengruppe der Größe 5	Wahrscheinlichkeit
(Ausfall, Ausfall, starke Schädigung, starke Schädigung, sehr schwache Schädigung)	1/3
(Ausfall, Ausfall, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	1/3
(Ausfall, starke Schädigung, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	1/3

Insgesamt sind maximal $\binom{r}{\hat{r}-r}$ mögliche Kombinationen zu betrachten.

Dieser Ansatz kann auf $\hat{r} > 2r$ verallgemeinert werden. Hierbei wird von der Beobachtung ausgegangen, dass bei $\hat{r} = 2r$ alle Komponenten genau einmal dupliziert werden. Dies legt nahe, bei $\hat{r} > 2r$ zunächst alle Komponenten $\left\lfloor \frac{\hat{r}-r}{r} \right\rfloor$ mal zu duplizieren, wobei $\lfloor x \rfloor$ den ganzzahligen Anteil von x bezeichnet, und dann $\hat{r} - \left\lfloor \frac{\hat{r}-r}{r} \right\rfloor r$ zufällig ausgewählte Komponenten zu duplizieren. Beim oben eingeführten Beispiereignis ergibt sich bei einer Zielkomponentengruppe der Größe 7 das in Tab. 6.9 dargestellte Ergebnis.

Tab. 6.9 Schädigungsvektor der Größe 7 und jeweilige Wahrscheinlichkeit

Schädigungsvektor in der Komponentengruppe der Größe 7	Wahrscheinlichkeit
(Ausfall, Ausfall, Ausfall, starke Schädigung, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	1/3
(Ausfall, Ausfall, starke Schädigung, starke Schädigung, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	1/3
(Ausfall, Ausfall, starke Schädigung, starke Schädigung, sehr schwache Schädigung, sehr schwache Schädigung, sehr schwache Schädigung)	1/3

Insgesamt sind jeweils maximal $\binom{r}{\hat{r}-\left\lfloor \frac{\hat{r}-r}{r} \right\rfloor r}$ mögliche Kombinationen zu betrachten.

6.2.5 Modellbasiertes Mapping Up

Es ist ebenfalls möglich, statt heuristische oder probabilistisch-kombinatorische Ansätze zu verwenden, die „fehlenden“ $\hat{r} - r$ Komponentenschädigungen mit Hilfe eines Modells zu bestimmen. Hier bietet sich wiederum das Kopplungsmodell an. Der Kopplungsparameter η_j wird aus den Schädigungen der beobachteten Komponenten geschätzt (siehe Gleichung (5.1)). Mittels des Kopplungsparameters wird für die nicht beobachteten $\hat{r} - r$ Komponenten berechnet, mit welcher Wahrscheinlichkeit sie ausgefallen wären. Die Wahrscheinlichkeiten, dass k von $\hat{r} - r$ Komponenten ausgefallen wären, ist – analog zu Gleichung (6.39) – für $k = 0 \dots \hat{r} - r$ gegeben als

$$p(k \setminus (\hat{r} - r)) = \binom{\hat{r} - r}{k} \int_0^1 \eta_j^k (1 - \eta_j)^{\hat{r} - r - k} p(\eta_j) d\eta_j \quad (6.40)$$

Damit lassen sich die Elemente des Interpretationsvektors bestimmen zu

$$w_{k \setminus \hat{r}} = \sum_{i=0}^{\hat{r}-r} w_{(k-i) \setminus r} p(i \setminus (\hat{r} - r)) \quad (6.41)$$

Die rechte Seite von Gleichung (6.41) hat die Form einer Summe über alle möglichen Anzahlen von Ausfällen unbeobachteter Komponenten i von dem Produkt der Wahrscheinlichkeit eines $(k - i)$ von r Ausfalls der beobachteten Komponenten (diese Größen ergeben sich wie in Abschnitt 2.2.1 beschrieben) und der Wahrscheinlichkeit, dass i der nicht beobachteten $\hat{r} - r$ Komponenten ausfallen (Gleichung (5.3)). Der Unterschied zur in Abschnitt 6.2.1 beschriebenen Vorgehensweise besteht darin, dass hier nicht alle $\hat{r}$, sondern nur die $\hat{r} - r$ fehlenden Komponenten mittels des Modells berücksichtigt werden, während die beobachteten r Komponenten wie beobachtet berücksichtigt werden. Dies führt im Vergleich zu dem in Abschnitt 6.2.1 beschriebenen Verfahren zu einer geringeren Schätzunsicherheit, da die Komponentenschädigungen von r Komponenten als bekannt angesehen werden.

6.3 Kompatibilität des Mappings mit a priori-Annahmen

In diesem Abschnitt wird die Kompatibilität der Wahl von a priori-Verteilungen bei der Anwendung Bayes'scher Verfahren und der vielen Mappingverfahren zugrundeliegenden Annahme, dass eine Komponentengruppe der Größe r statistisch äquivalent einer

beliebigen Untergruppe der Größe r einer größeren Gruppe mit Größe $\hat{r} > r$ ist, untersucht.

Um dies zu untersuchen, soll zunächst der Fall betrachtet werden, dass GVA-Ereignisse in Komponentengruppen der Größe r beobachtet wurden, während die Zielkomponentengruppe die Größe $\check{r} < r$ hat ('Mapping Down'). Das Mapping soll auf der Annahme beruhen, dass sich die Zielkomponentengruppe im statistischen Sinne verhält wie eine Untergruppe der Größe $\check{r}$ einer GVA-Komponentengruppe der Größe r .

Aus dieser Annahme folgen zwei mögliche Vorgehensweisen für die Berechnung der GVA-Wahrscheinlichkeiten in der Komponentengruppe der Größe $\check{r}$, die mit den hebräischen Buchstaben Aleph (א) und Beth (ב) bezeichnet werden:

- א Erst die Ereignisse für die Größe $\check{r}$ der Zielkomponentengruppe mappen mittels der Annahme, dass sich die Zielkomponentengruppe im statistischen Sinne verhält wie eine Untergruppe der Größe $\check{r}$, dann die GVA-Wahrscheinlichkeiten mittels der Ereignisse der passenden Größe schätzen.
- ב Erst die GVA-Wahrscheinlichkeiten für GVA-Gruppen der Komponentengruppengröße r berechnen und aus diesen mittels der Annahme, dass sich die Zielkomponentengruppe im statistischen Sinne verhält wie eine Untergruppe der Größe $\check{r}$, die GVA-Wahrscheinlichkeiten für Gruppengröße $\check{r}$ berechnen.

Beide Vorgehensweisen sind gleichermaßen plausibel und sollten zu identischen Ergebnissen führen.

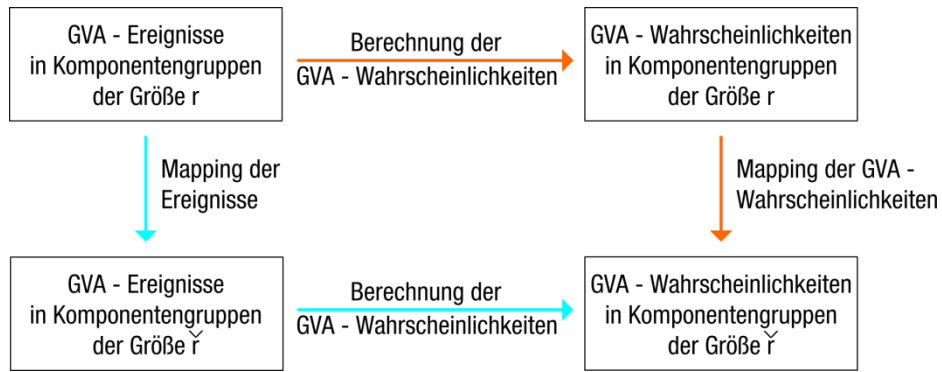


Abb. 6.6 Darstellung der Vorgehensweisen nach Verfahren 1 (cyan) und Verfahren 2 (orange)

Die in Abb. 6.6 dargestellte Vorgehensweise kann an einem einfachen Beispiel unter Verwendung des Maximum-Likelihood-Schätzers illustriert werden: Wenn zwei (2 von 3)-Ereignisse in einer Komponentengruppe der Größe 3 beobachtet wurden, so folgt für die Schätzwerte der Anteile der verschiedenen GVA-Ausfallkombinationen in der Komponentengruppe der Größe $r = 3$, $x_{i \setminus 3}^*$:

$$x_{2 \setminus 3}^* = 1 \quad (6.42)$$

$$x_{0 \setminus 3}^* = x_{1 \setminus 3}^* = x_{3 \setminus 3}^* = 0$$

Im Weiteren werden aus diesen Größen die Schätzungen der Anteile der verschiedenen GVA-Ausfallkombinationen in der Komponentengruppe der Größe $\check{r} = 2$, d. h. nach Vorgehensweise 2, berechnet.

Allgemein gilt unter der Annahme, dass eine Komponentengruppe der Größe $\check{r} = 2$ einer beliebigen Untergruppe der Größe 2 einer Komponentengruppe der Größe $r = 3$ äquivalent ist (siehe auch Anhang A):

$$\begin{aligned} y_{0 \setminus 2} &= x_{0 \setminus 3} + \frac{1}{3} x_{1 \setminus 3} \\ y_{1 \setminus 2} &= \frac{2}{3} x_{1 \setminus 3} + \frac{2}{3} x_{2 \setminus 3} \\ y_{2 \setminus 2} &= \frac{1}{3} x_{2 \setminus 3} + x_{3 \setminus 3} \end{aligned} \quad (6.43)$$

Daraus folgt für die Punktschätzer nach Vorgehensweise 2:

$$y_{0\setminus 2}^{*2} = x_{0\setminus 3}^* + \frac{1}{3} x_{1\setminus 3}^* = 0 \quad (6.44)$$

$$y_{1\setminus 2}^{*2} = \frac{2}{3} x_{1\setminus 3}^* + \frac{2}{3} x_{2\setminus 3}^* = \frac{2}{3}$$

$$y_{2\setminus 2}^{*2} = \frac{1}{3} x_{2\setminus 3}^* + x_{3\setminus 3}^* = \frac{1}{3}$$

Für die Vorgehensweise 3 wird zunächst bestimmt, welche Ereignisse in der Komponentengruppe der der Größe $\check{r} = 2$ aufgetreten wären:

- Mit Wahrscheinlichkeit $1/9$ wären zwei (2 von 2)-Ereignisse aufgetreten.
- Mit Wahrscheinlichkeit $4/9$ wären zwei (1 von 2)-Ereignisse aufgetreten.
- Mit Wahrscheinlichkeit $4/9$ wäre ein (2 von 2)-Ereignis und ein (1 von 2)-Ereignis aufgetreten.

Daraus folgt:

$$y_{0\setminus 2}^{*3} = 0 \quad (6.45)$$

$$y_{1\setminus 2}^{*3} = \frac{4}{9} + \frac{4}{9} \times \frac{1}{2} = \frac{2}{3}$$

$$y_{2\setminus 2}^{*3} = \frac{1}{9} + \frac{4}{9} \times \frac{1}{2} = \frac{1}{3}$$

Somit gilt $y_{i\setminus 2}^{*3} = y_{i\setminus 2}^{*2}$:

Eine Verallgemeinerung dieser Betrachtung auf unterschiedliche Ereignisse und andere Komponentengrößen ist trivial; es ergibt sich stets notwendigerweise dasselbe Ergebnis bei beiden Rechnungsmethoden.

Analoges sollte gelten, wenn die Unsicherheit einbezogen wird und Bayes'sche statistische Verfahren angewandt werden. Hier sollten sich unabhängig davon, ob nach Verfahren 3 oder 2 vorgegangen wird, dieselben a posteriori-Verteilungen der $y_{i\setminus 2}$ ergeben.

Dies wird im Folgenden untersucht.

Hierzu wird zunächst der Fall betrachtet, dass keine Ereignisse beobachtet wurden.

Bei Vorgehensweise 8 ist die a posteriori-Verteilung einfach die a priori-Verteilung für Komponentengruppengröße 3. Es gilt

$$p^8(y_{1\setminus 2}, y_{2\setminus 2}) = \pi(y_{1\setminus 2}, y_{2\setminus 2}) \quad (6.46)$$

Hierbei wurde die erste Variable $y_{0\setminus 2}$ nicht geschrieben, da $\sum_{i=0}^2 y_{i\setminus 2} = 1$ gilt.

Wählt man die a priori-Verteilung nach dem Verfahren von Jeffreys, so erhält man eine Dirichlet-Verteilung mit Parametern 1/2:

$$\pi(y_{1\setminus 2}, y_{2\setminus 2}) \propto (1 - y_{1\setminus 2} - y_{2\setminus 2})^{-\frac{1}{2}} (y_{1\setminus 2})^{-\frac{1}{2}} (y_{2\setminus 2})^{-\frac{1}{2}} \quad (6.47)$$

Bei Vorgehensweise 9 wird die a posteriori-Verteilung wie folgt berechnet: Die a posteriori-Verteilung für die $x_{i\setminus 3}$ ist die a priori-Verteilung für Komponentengruppengröße 3. Es gilt

$$p(x_{1\setminus 3}, x_{2\setminus 3}, x_{3\setminus 3}) = \pi(x_{1\setminus 3}, x_{2\setminus 3}, x_{3\setminus 3}) \quad (6.48)$$

wobei $\pi(x_{1\setminus 3}, x_{2\setminus 3}, x_{3\setminus 3})$ wieder nach dem Verfahren von Jeffreys eine Dirichlet-Verteilung mit den Parametern $\frac{1}{2}$ ist. $p^9(y_{1\setminus 2}, y_{2\setminus 2})$ ist dann durch die Beziehungen zwischen den $x_{i\setminus 3}$ und $y_{i\setminus 2}$ (Gleichung (s.(6.30)) festgelegt.

Am einfachsten lassen sich die Ergebnisse mittels einer Monte-Carlo-Rechnung vergleichen, indem Samples aus $p(x_{1\setminus 3}, x_{2\setminus 3}, x_{3\setminus 3})$ gezogen werden und daraus mit Gleichung (6.28) 2-Tupel $(y_{1\setminus 2}, y_{2\setminus 2})$ berechnet werden. Diese sind dann gemäß $p^9(y_{1\setminus 2}, y_{2\setminus 2})$ verteilt.

Wählt man die a priori-Verteilungen nach dem Verfahren von Jeffreys als Dirichlet-Verteilung mit Parametern $(1/2, 1/2, 1/2)$, so gilt

$$p^{\aleph}(y_{1\setminus 2}, y_{2\setminus 2}) = \pi(y_{1\setminus 2}, y_{2\setminus 2}) \propto \frac{1}{\sqrt{1 - y_{1\setminus 2} - y_{2\setminus 2}} \sqrt{y_{1\setminus 2}} \sqrt{y_{2\setminus 2}}} \quad (6.49)$$

und analog

$$\pi(x_{1\setminus 3}, x_{2\setminus 3}, x_{3\setminus 3}) \propto \frac{1}{\sqrt{1 - x_{1\setminus 3} - x_{2\setminus 3} - x_{3\setminus 3}} \sqrt{x_{1\setminus 3}} \sqrt{x_{2\setminus 3}} \sqrt{x_{3\setminus 3}}} \quad (6.50)$$

Aus den Samples von $\pi(x_{1\setminus 3}, x_{2\setminus 3}, x_{3\setminus 3})$ können, wie oben dargestellt, mittels mit Gleichung (6.28) gemäß $p^{\beth}(y_{1\setminus 2}, y_{2\setminus 2})$ verteilte Samples berechnet werden. In den Abbildungen Abb. 6.7 und Abb. 6.8 sind die Marginalverteilungen von $y_{2\setminus 2}$ bzw. $y_{1\setminus 2}$ nach Verfahren $\aleph$ und $\beth$ verglichen. Die Marginalverteilungen von $y_{0\setminus 2}$ entsprechen aus Symmetriegründen jeweils denjenigen von $y_{2\setminus 2}$. Die Marginalverteilungen von $p^{\aleph}$ sind Betaverteilungen mit den Parametern 1/2 und 1, da $p^{\aleph}$ eine Dirichlet-Verteilung mit den Parametern (1/2, 1/2, 1/2) ist:

$$p^{\aleph}(y_{i\setminus 2}) = \frac{1}{2 \sqrt{y_{i\setminus 2}}} \quad (6.51)$$

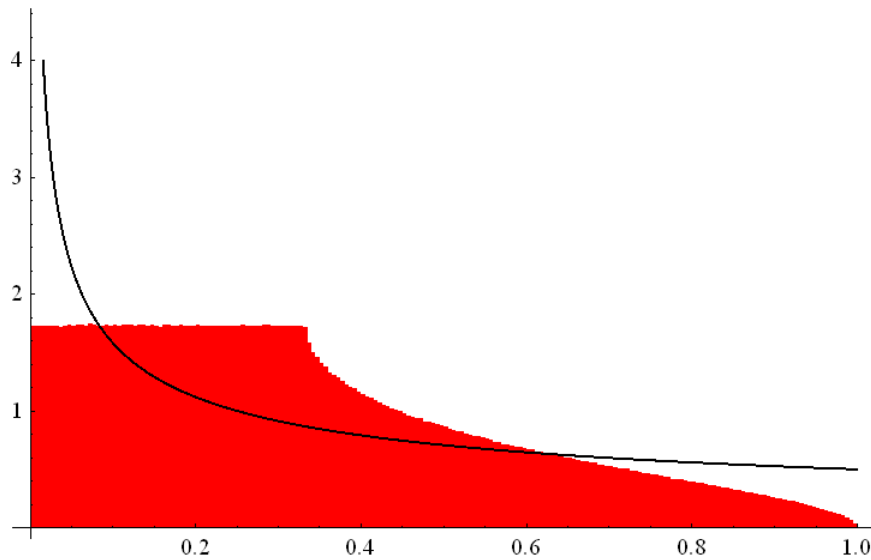


Abb. 6.7 Vergleich der Marginalverteilungen von $y_{2\setminus 2}$ nach Verfahren $\aleph$ (schwarze Kurve) und $\beth$ (rotes Histogramm) für eine Komponentengruppe der Größe $\check{r} = 2$, die als Teil einer Komponentengruppe der Größe $r = 3$ betrachtet wird

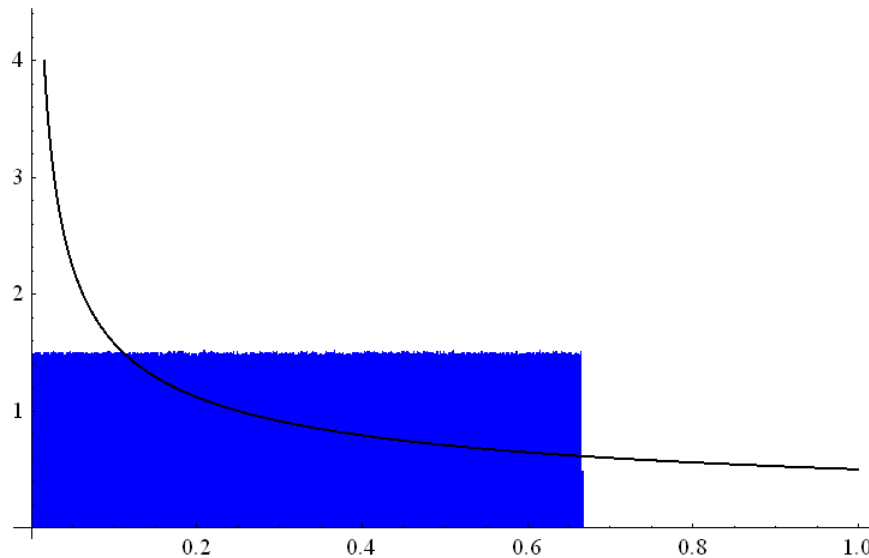


Abb. 6.8 Vergleich der Marginalverteilungen von $y_{1\setminus 2}$ nach Verfahren $\aleph$ (schwarze Kurve) und $\beth$ (blaues Histogramm) für eine Komponentengruppe der Größe $\check{r} = 2$, die als Teil einer Komponentengruppe der Größe $r = 3$ betrachtet wird

Es ist offensichtlich, dass die Verteilungen abweichen. An Abb. 6.8 ist z. B. erkennbar, dass $y_{1\setminus 2}$ nach Verfahren $\beth$ einen Maximalwert von $2/3$ hat. Dies ist auch aus Gleichung (6.28) ablesbar, da $y_{1\setminus 2} = \frac{2}{3}x_{1\setminus 3} + \frac{2}{3}x_{2\setminus 3}$ und $\sum_{i=0}^3 x_{i\setminus 3} = 1$, also insbesondere $x_{1\setminus 3} + x_{2\setminus 3} \leq 1$ ist. Dies folgt daraus, dass die Ereignisse in der Komponentengruppe der Größe 2 tatsächlich Ereignisse in einer Untergruppe einer Komponentengruppe der Größe 3 sind. Dieses Wissen ist bei der Festlegung der a priori-Verteilung in der Komponentengruppe der Größe 2 bei Verfahren $\aleph$ (Gleichung (6.51)) nicht eingeflossen. Somit ist die a priori-Verteilung nicht korrekt gewählt worden.

Diese Erkenntnis hat grundsätzliche Bedeutung:

Eine nichtinformative Wahl der a priori-Verteilung (z. B. nach dem Verfahren von Jeffreys) für alle Komponentengruppengrößen ist mit der Annahme, dass sich eine Komponentengruppe der Größe r verhält wie eine Untergruppe einer größeren Komponentengruppe der Größe $\hat{r} > r$, nicht kompatibel.

Die Unterschiede werden noch deutlicher, wenn ein Mapping zwischen Komponentengruppen mit größerem Unterschied im Redundanzgrad $|r - \check{r}|$ betrachtet wird. Als Beispiel wird hier das Mapping von einer Komponentengruppengröße $r = 4$ nach $\check{r} = 2$

betrachtet. Dabei wird analog oben vorgegangen. Für die a priori-Verteilung in der Komponentengruppe der Größe 4 (Verfahren 2) gilt statt Gleichung (6.51):

$$\pi(x_{1\setminus 4}, x_{2\setminus 4}, x_{3\setminus 4}, x_{4\setminus 4}) \propto \frac{1}{\sqrt{1 - x_{1\setminus 4} - x_{2\setminus 4} - x_{3\setminus 4} - x_{4\setminus 4}} \sqrt{x_{1\setminus 4}} \sqrt{x_{2\setminus 4}} \sqrt{x_{3\setminus 4}} \sqrt{x_{4\setminus 4}}} \quad (6.52)$$

Die Samples von $p^2(y_{1\setminus 2}, y_{2\setminus 2})$ werden aus $\pi(x_{1\setminus 4}, x_{2\setminus 4}, x_{3\setminus 4}, x_{4\setminus 4})$ mittels von Gleichung (A. 6) im Anhang A berechnet. Der Vergleich der Verteilungen ist in den Abbildungen Abb. 6.9 und Abb. 6.10 dargestellt.

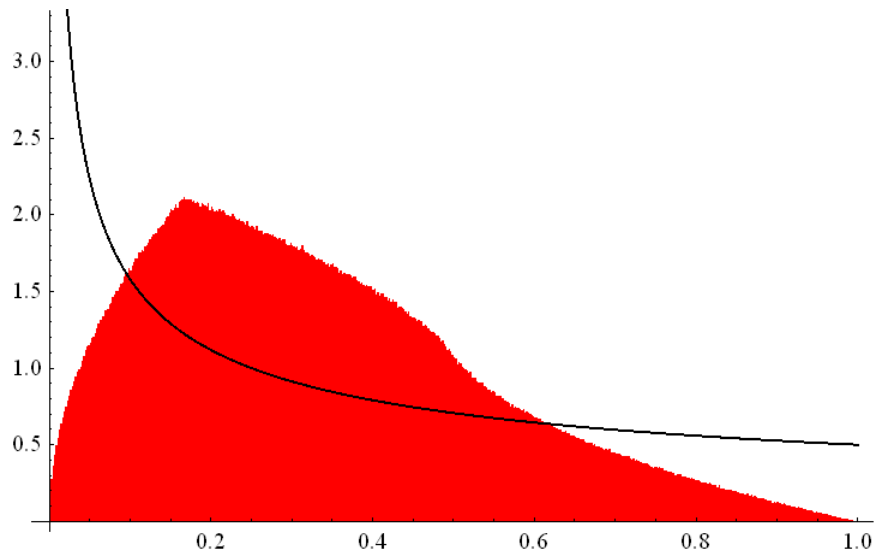


Abb. 6.9 Vergleich der Marginalverteilungen von $y_{2\setminus 2}$ nach Verfahren 8 (schwarze Kurve) und 2 (rotes Histogramm) für eine Komponentengruppe der Größe $\check{r} = 2$, die als Teil einer Komponentengruppe der Größe $r = 4$ betrachtet wird

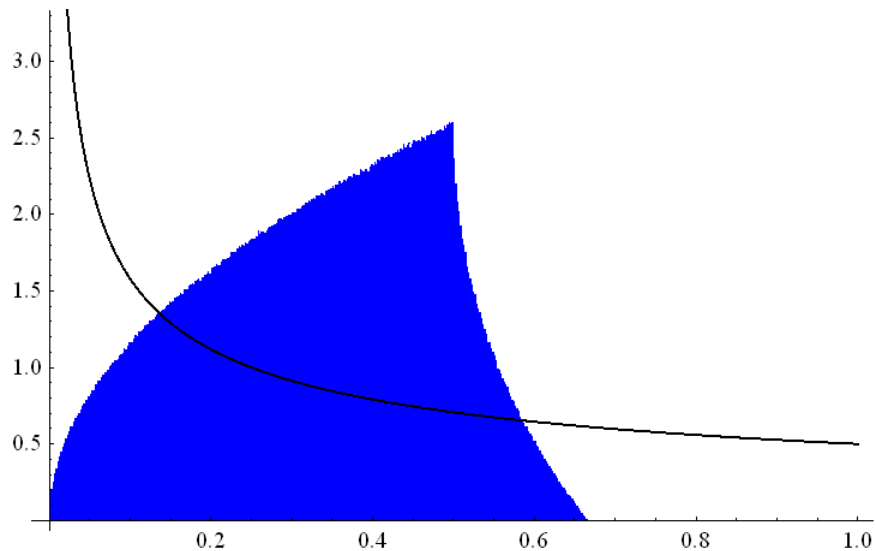


Abb. 6.10 Vergleich der Marginalverteilungen von $y_{1\setminus 2}$ nach Verfahren $\mathfrak{K}$ (schwarze Kurve) und $\mathfrak{L}$ (blaues Histogramm) für eine Komponentengruppe der Größe $\check{r} = 2$, die als Teil einer Komponentengruppe der Größe $r = 4$ betrachtet wird

Es ist erkennbar, dass die Verteilungsformen noch unähnlicher werden: Sie haben einen ausgeprägten Peak bei einem Wert $0 < y_{i\setminus 2} < 1$. Die $y_{i\setminus 2}$ sind hier gewichtete Summen jeweils dreier $x_{i\setminus 4}$, während sie oben gewichtete Summen jeweils zweier $x_{i\setminus 3}$ waren (Gleichung (6.43) bzw. Gleichung (6.28)). Es zeigen sich erste Ausprägungen des grundsätzlichen Phänomens, das dem zentralen Grenzwertsatz¹⁵ zugrunde liegt. Hier sind die $x_{i\setminus r}$ nicht vollständig unabhängig. Es gilt die Bedingung $\sum_{i=0}^r x_{i\setminus r}$. Trotzdem sind sie unabhängig genug, dass die Verteilung immer schmäler um den Mittelwert wird und die Form einer Normalverteilung annimmt, wenn ein wachsendes r betrachtet wird. Dies ist in den Abbildungen Abb. 6.11 und Abb. 6.12 für $r = 128$ dargestellt. Hier sind $y_{i\setminus 2}$ jeweils gewichtete Summen aus 125 $x_{i\setminus 4}$ (siehe Anhang A).

¹⁵ Der zentrale Grenzwertsatz besagt, dass die Verteilung der Summe von n unabhängigen identisch verteilten Zufallszahlen einer beliebigen Verteilung mit endlichem Mittelwert μ und endlicher Standardabweichung σ mit zunehmendem n gegen eine Normalverteilung mit Mittelwert $n\mu$ und Standardabweichung $\sigma\sqrt{n}$ konvergiert.

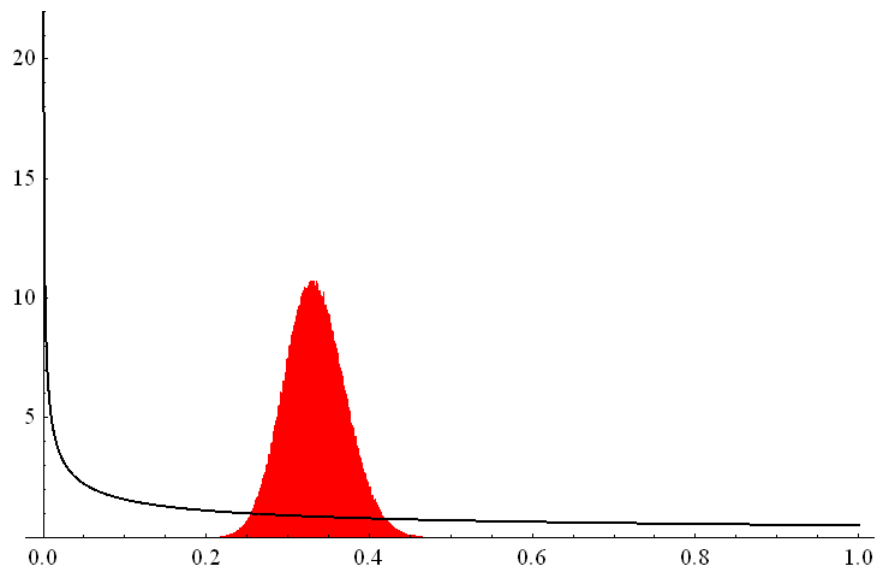


Abb. 6.11 Vergleich der Marginalverteilungen von $y_{2 \setminus 2}$ nach Verfahren $\mathfrak{K}$ (schwarze Kurve) und $\mathfrak{I}$ (rotes Histogramm) für eine Komponentengruppe der Größe $\check{r} = 2$, die als Teil einer Komponentengruppe der Größe $r = 128$ betrachtet wird

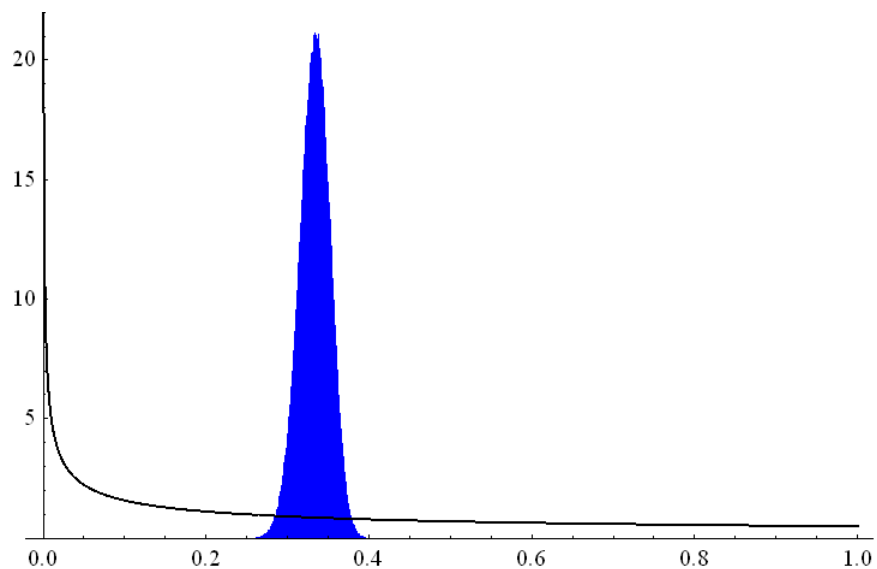


Abb. 6.12 Vergleich der Marginalverteilungen von $y_{2 \setminus 2}$ nach Verfahren $\mathfrak{K}$ (schwarze Kurve) und $\mathfrak{I}$ (rotes Histogramm) für eine Komponentengruppe der Größe $\check{r} = 2$, die als Teil einer Komponentengruppe der Größe $r = 128$ betrachtet wird

Es ist darauf hinzuweisen, dass die Marginalverteilungen nach Verfahren $\aleph$ in den obigen Abbildungen identisch sind und nur aufgrund der verschiedenen Skalierungen der Ordinate unterschiedlich erscheinen.

Charakteristika der Verteilung sind in Tab. 6.10 und Tab. 6.11 dargestellt, wobei jeweils zwei Dezimalstellen angegeben wurden.

Tab. 6.10 Vergleich der Charakteristika der Marginalverteilungen von $y_{2\setminus 2}$ nach Verfahren $\aleph$ und $\beth$

	r	Minimum	5 %-Quantil	Median	95 %-Quantil	Supremum	Mittelwert	Standardabweichung
$\aleph$	alle	0	0.0025	0.25	0.90	1	1/3	0.30
$\beth$	3	0	0.0029	0.29	0.80	1	0.33	0.24
$\beth$	4	0	0.059	0.30	0.73	1	0.33	0.20
$\beth$	128	0	0.27	0.33	0.40	nicht ermittelt	0.33	0.037

Tab. 6.11 Vergleich der Charakteristika der Marginalverteilungen von $y_{1\setminus 2}$ nach Verfahren $\aleph$ und $\beth$

	r	Minimum	5 %-Quantil	Median	95 %-Quantil	Supremum	Mittelwert	Standardabweichung
$\aleph$	alle	0	0.0025	0.25	0.90	1	1/3	0.30
$\beth$	3	0	0.0033	0.33	0.63	$\frac{2}{3}$	0.33	0.19
$\beth$	4	0	0.075	0.35	0.56	$\frac{2}{3}$	0.33	0.15
$\beth$	128	0	0.30	0.33	0.36	nicht ermittelt	0.33	0.018

Wie anhand der Tab. 6.10 und Tab. 6.11 erkennbar ist, wird die Breite der a priori-Verteilung nach Verfahren $\beth$ mit steigendem r sehr gering. Weiterhin existiert – unabhängig von r – für $y_{1\setminus 2}$ eine obere Schranke von $2/3$. d. h. größere Werte sind ausgeschlossen.

Die dargestellten Ergebnisse sind qualitativ nicht spezifisch für die spezielle gewählte a priori-Verteilung, d. h. die grundsätzlichen aufgezeigten problematischen Eigenschaften sind auch für anders gewählte a priori-Verteilungen (z. B. Gleichverteilung) vorhanden.

Die Ergebnisse sind auch nicht spezifisch für ein bestimmtes GVA-Modell. Vielmehr sind alle Modelle, die für die Komponentengruppengröße spezifische Parameter aufweisen, betroffen. Dies umfasst sowohl das Alpha-Faktor-Modell (Abschnitt 4.4) als auch die im Rahmen dieses Vorhabens entwickelten Modelle A, B und C (Abschnitt 5.2). Das Kopplungsmodell (Abschnitt 2.1) ist hingegen nicht betroffen, da es nur Parameter hat, die nicht komponentengruppengrößenspezifisch sind (Kopplungsparameter η_j und Rate λ_j).

Zusammenfassend ist das Ergebnis, dass die Annahme, dass eine Komponentengruppe als zufällige Untergruppe einer großen Komponentengruppe angesehen wird, sehr weitreichende Konsequenzen hat und mit der üblichen Vorgehensweise zur Schätzung von GVA-Wahrscheinlichkeiten unter Verwendung nicht-informativer a priori-Verteilungen, z. B. mit Modellen A - C oder dem Alpha-Faktor-Modell nicht kompatibel ist.

6.3.1 Diskussion der Lösungsmöglichkeiten und Konsequenzen

Das Problem der schmaler werdenden a priori-Verteilungen ließe sich womöglich prinzipiell umgehen, indem von der oben dargestellten a priori-Verteilung $\pi(y_{1\setminus 2}, y_{2\setminus 2})$ für $\check{r} = 2$ ausgegangen wird und für $r > 2$ solche a priori-Verteilungen $\pi(x_{1\setminus r}, \dots, x_{r\setminus r})$ gewählt werden, dass sich beim Berechnen der sich durch das Mapping Down ergebende Verteilungen aus diesen wieder $\pi(y_{1\setminus 2}, y_{2\setminus 2})$ ergibt. Die Existenz von solchen Verteilungen wurde hier nicht untersucht.

Die Existenz der oberen Schranke für $y_{1\setminus 2}$ ist allerdings untrennbar verbunden mit der Annahme, dass eine Komponentengruppe der Größe 2 als zufällige Untergruppe einer großen Komponentengruppe angesehen wird.

Somit erscheint es fraglich, ob die weitere Verwendung der Annahme, dass eine Komponentengruppe der Größe $\check{r}$ als zufällige Untergruppe einer großen Komponentengruppe $r > \check{r}$ angesehen werden kann, möglich ist.

6.4 Vergleich der Mappingansätze

In den vorausgehenden Abschnitten dieses Abschnitts wurden verschiedene Ansätze für das Mapping angegeben und diskutiert. Sie sind teilweise modellbasiert, teilweise basieren sie auf der Annahme, dass sich eine kleine Komponentengruppe identisch verhält wie eine Untergruppe einer größeren Komponentengruppe, teilweise stellen sie konservative Vorgehensweisen dar. Teilweise wird eine mit dem Mapping verbundene Unsicherheit berücksichtigt. Teilweise beziehen sich die Mappingalgorithmen unmittelbar auf die Komponentenschädigungen, teilweise werden die Interpretationen als mögliche Ausfälle (siehe Abschnitt 2.2.1) verwendet.

Um die Mappingverfahren sachgerecht zu bewerten, müssen Informationen vorhanden sein, wie die verschiedenen GVA-Phänomene in Komponentengruppen verschiedener Größe tatsächlich wirken. Allerdings lassen sich der Betriebserfahrung im Allgemeinen solche Informationen nicht entnehmen. Ein denkbarer Ausweg wäre, die verschiedenen Vorgehensweisen anhand von detaillierten Modellen der GVA-Entstehung und GVA-Entdeckung zu untersuchen. Derzeit erscheint dies allerdings ebenso problematisch, da die Modelle wiederum an der Betriebserfahrung oder anderen empirischen Informationen zu validieren wären, was beim jetzt vorhandenen Kenntnisstand ebenso schwierig erscheint wie ein direkter Vergleich. Eine Verwendung von nicht durch Betriebserfahrung gestützten Modellannahmen wäre ebenso wenig wohlbegründet wie das unmittelbare Treffen von Annahmen über das Mapping.

Zusammenfassend fehlt eine empirische Bewertungsbasis für Vorgehensweisen zum Mapping. Um eine solche zu entwickeln, erscheint ein umfassendes Forschungsvorhaben erforderlich mit den folgenden Elementen:

- Detaillierte Analyse der Betriebserfahrung. Hierbei ist es notwendig, nicht nur Meldepflichtige Ereignisse, sondern insbesondere auch bei Betreibern vorliegende detailliertere und umfassendere Informationen einzubeziehen (z. B. aus Wirkleistungsmessungen von Armaturen).
- Entwicklung und Validierung detaillierter Modelle der GVA-Entstehung und Entdeckung. Diese Modelle müssen nicht nur Ausfälle, sondern auch Schädigungen beschreiben.
- Gegebenenfalls Durchführung von Versuchen mit detaillierter Beobachtung von Komponentenschädigungen.

Bei solchen Untersuchungen wäre zu berücksichtigen, dass GVA-Phänomene sehr vielgestaltig sind und deshalb angemessene Mappingalgorithmen Komponentenart- und Phänomen-spezifisch sein können.

Da somit beim jetzigen Kenntnisstand nicht beurteilt werden kann, welches Verhalten der Mappingalgorithmen richtig ist, sind Vergleichsrechnungen anhand simulierter Daten oder der Ereignisse der Betriebserfahrung zum jetzigen Zeitpunkt nicht zielführend.

Im Folgenden wird daher betrachtet, in wieweit die Mappingverfahren

- konservativ und
- kompatibel mit den entwickelten Schätzalgorithmen

sind.

Der Begriff der Konservativität ist hier zunächst nicht eindeutig bestimmt, da wie oben beschrieben nicht bekannt ist, welche Ergebnisse des Mappings „richtig“ sind, sondern muss durch festzulegende Kriterien konkretisiert werden.

Dazu können folgende Kriterien verwendet werden:

1. Komplette GVA werden auf komplette GVA abgebildet.
Begründung: Ist dies nicht erfüllt, so kann die Wahrscheinlichkeit kompletter GVA unterschätzt werden.
2. GVA mit zwei Ausfällen werden beim Mapping Down auf ein Ereignis mit mindestens zwei Ausfällen abgebildet.
Begründung: In sehr vielen Fällen werden systematische Ursachen erst nach Feststellung des Ausfalls bzw. der Schädigung zweier Komponenten erkannt.
3. Beim Mapping Up ist die Anzahl ausgefallener Komponenten in der großen Komponentengruppe mindestens gleich hoch wie in der kleinen Komponentengruppe.
Begründung: Wenn eine beobachtete Komponentengruppe als Teil einer größeren, nicht vollständig beobachteten Komponentengruppe aufgefasst werden kann, ist dies stets der Fall.
4. Beim Mapping Up auf eine Komponentengruppe doppelter Größe ist die Anzahl ausgefallener Komponenten (mindestens) doppelt so hoch.
Begründung: Wenn eine weitere Komponentengruppe existieren würde, die voll-

ständig identisch in allen Details der technische Eigenschaften, Belastungen, Wartung usw. wäre, träfe dies zu.

Es sind keine sinnvollen noch konservativeren Kriterien unmittelbar erkennbar. Noch konservativer wäre es nur, beim Mapping – unabhängig von den beobachteten Komponentenschädigungen – alle Komponenten stets als ausgefallen zu betrachten. Dies stellt aber keine gut begründbare Vorgehensweise dar.

Auch für die praktische Anwendung eines Mappingverfahrens auf die deutsche Betriebserfahrung ist relevant, ob komplette GVA auf komplette GVA abgebildet werden (Kriterium 1), da dann eine separate Behandlung von Phänomenen mit notwendigem Ausfall aller Komponenten nicht erforderlich ist und somit die in der existierenden GVA-Datenbank vorliegenden Informationen ausreichend für eine Quantifizierung sind.

Neben der Konservativität stellen auch die Kompatibilität mit den a priori-Annahmen und die Berücksichtigung von mit dem Mapping verbundener Unsicherheit Kriterien für die Auswahl der Mappingverfahren dar. Alle diese Kriterien sind in Tab. 6.12 dargestellt.

Tab. 6.12 Vergleich der Mappingalgorithmen

Mappingalgorithmus	Konservativ nach Kriterien				Kompatibel mit a priori Modelle A - C	Unsicherheit des Mappings einbezogen
	1	2	3	4		
Modellbasierter Ansatz nach 6.2.1	nein	nein	nein	nein	nein	ja
Probabilistisch-kombinatorisches Mapping Down nach 6.2.2	ja	nein	n. a.	n. a.	nein	teilweise ¹⁶
Heuristisches Mapping Up (für $\hat{r} = r + 1$) nach 6.2.4	ja	n. a.	ja	n. a.	ja	teilweise
Heuristisches Mapping Up: „Ziehen mit Zurücklegen“ nach 6.2.4	ja	n. a.	ja	nein	ja	teilweise
Heuristisches Mapping Up: „Ziehen ohne Zurücklegen“ nach 6.2.4	ja	n. a.	ja	nein	ja	teilweise
Probabilistisch-kombinatorisches Mapping Up ¹⁷ nach 6.2.1	(nein)	n. a.	(ja)	(nein)	nein	ja
Konservatives Mapping Down nach 6.2.3	ja	ja	n. a.	n. a.	ja	nein
Konservatives Mapping Up nach 6.2.3	ja	n. a.	ja	ja	ja	nein
Modellbasiertes Mapping Up nach 6.2.5	nein	n. a.	ja	nein	nein	ja

¹⁶ „teilweise“ zeigt an, dass die mit dem Mapping verbundene Unsicherheit nur in Teilen berücksichtigt wird. Dies folgt z. B. aus der Erfüllung von Kriterium 1, das erfordert, dass ein kompletter GVA auf einen kompletten GVA (sicher) abgebildet wird.

¹⁷ Für das probabilistisch-kombinatorische Mapping up sind die Kriterien nicht unmittelbar anwendbar, da nicht einzelne Ereignisse, sondern die Gesamtmenge der Ereignisse verarbeitet wird. Deshalb wurden die Bewertungen in Klammern gesetzt. Es wurde mit „(ja)“ bewertet, wenn der Einfluss auf das Schätzergebnis Ereignissen entspricht, die die jeweiligen Kriterien erfüllen (Siehe dazu Gleichung (6.23) und den auf diese Gleichung folgenden Text).

Zusammenfassend lässt sich sagen, dass nur das konservative Mapping Down und konservative Mapping Up alle Kriterien der Konservativität erfüllen. Darüber hinaus sind das konservative Mapping Down und konservative Mapping Up kompatibel mit den entwickelten Schätzalgorithmen A - C und bilden komplette GVA auf komplette GVA ab.

Allerdings ist offensichtlich, dass sie für ein Mapping zwischen Komponentengruppen sehr unterschiedlicher Größe eine sehr konservative Extrapolation darstellen. Für ein Mapping zwischen Komponentengruppen ähnlicher Größe können sie allerdings ein brauchbares Mittel darstellen, um zu Quantifizierungen von GVA zu kommen, ohne nicht zwingend begründbare Annahmen zu treffen. Im Abschnitt 7.3 ist ein solches Beispiel für die Übertragung von Betriebserfahrung an Komponentengruppen der Größe 4 auf Komponentengruppen der Größe 3 dargestellt. Weiterhin kann eine solche Quantifizierung auch dann verwendet werden, wenn eine starke Konservativität akzeptabel ist (z. B. für Basisereignisse, die trotz konservativer Quantifizierung eine geringe Importanz aufweisen).

Allgemein lässt sich der Schluss ziehen, dass verschiedene Ansätze zum Mapping existieren, jedoch bei dem heutigen Kenntnisstand nicht entschieden werden kann, welche Vorgehensweise unter welchen Bedingungen eine realistische Abbildung ermöglicht und mit welchen Schätzabweichungen und Unsicherheiten sie verbunden ist. Für bestimmte Anwendungsfälle können jedoch die konservativen Verfahren eine gut begründbare und sachgerechte Vorgehensweise darstellen, wie in Abschnitt 7.3 gezeigt wird.

7 Anwendung auf die Betriebserfahrung

In diesem Abschnitt werden die entwickelten Verfahren auf repräsentative Datensätze der deutschen Betriebserfahrung angewandt. Allerdings sind die in der GVA-Datenbank enthaltenen Informationen zum Teil Betriebsgeheimnisse der Betreiber und sind deshalb vertraulich zu behandeln. Da eine Veröffentlichung der Schätzergebnisse mit den verschiedenen Verfahren möglicherweise Rückschlüsse auf die Daten erlauben könnte, wurden modifizierte Datensätze aus der deutschen Betriebserfahrung zusammengestellt.

7.1 Beispieldatensätze

Zum Vergleich der Schätzmethoden wurden zwei Datensätze verwendet, die jeweils nur Ereignisse an Komponentengruppen der Größe 4 enthalten. Diese Komponentengruppengröße wurde gewählt, da am meisten Betriebserfahrung an Komponentengruppen dieser Größe vorliegt und dieser Redundanzgrad insbesondere bei den noch im Betrieb befindlichen Anlagen am häufigsten vorkommt.

Datensatz 1 ist repräsentativ für Populationen mit einer relativ hohen Anzahl von Ereignissen. Er wurde wie folgt erzeugt:

- Ausgegangen wurde von einem GVA-Datensatz, in dem viele Ereignisse in Komponentengruppen der Größe 4 vorhanden sind.
- Zuerst wurden alle Ereignisse, die in Komponentengruppen einer abweichenden Größe $r \neq 4$ aufgetreten sind, entfernt.
- Dann wurde eine geringe Zahl zufällig ausgewählter Ereignisse gelöscht.
- Die Beobachtungszeit wurde mit einem Zufallswert multipliziert.

Der Datensatz weist 15 Ereignisse auf. In einigen Ereignissen haben Experten alle Komponenten als mindestens schwach geschädigt bewertet. Es wurden nie mehr als zwei Komponenten als ausgefallen bewertet.

Datensatz 2 ist repräsentativ für Populationen mit einer geringen Anzahl von Ereignissen. In diesen Datensätzen sind typischerweise wenige Ereignisse enthalten, bei de-

nen meist keine oder nur eine Komponente ausgefallen ist (potentieller GVA). Der Beispieldatensatz wurde wie folgt erzeugt:

- Ausgegangen wurde von einem GVA-Datensatz, in dem sehr wenige Ereignisse enthalten sind.
- Ein Ereignis in einer Komponentengruppe der Größe 4 wurde ausgewählt.
- Die Beobachtungszeit wurde willkürlich festgesetzt.

Der Datensatz weist somit ein einziges Ereignis auf. Von allen Experten wurde genau eine Komponente als ausgefallen bewertet. Drei Experten bewerteten die verbleibenden Komponenten als schwach geschädigt, ein Experte als nicht geschädigt.

Zur Untersuchung des konservativen Mappings wurden Datensatz 3 und 4 zusammengestellt. Datensatz 3 enthält Betriebserfahrung an Komponentengruppen der Größe 3 und 4, während Datensatz 4 nur Betriebserfahrung an Komponentengruppen der Größe 3 enthält.

Datensatz 3 wurde analog zu Datensatz 1 wie folgt erzeugt:

- Ausgegangen wurde von demselben GVA-Datensatz wie bei Datensatz 1.
- Zuerst wurden alle Ereignisse, die in Komponentengruppen einer abweichenden Größe $r \notin \{3,4\}$ aufgetreten sind, entfernt.
- Dann wurden eine geringe Zahl zufällig ausgewählte Ereignisse gelöscht.
- Die Beobachtungszeit wurde mit einem Zufallswert multipliziert.

Dieser Datensatz enthält 16 Ereignisse. Davon sind 15 Ereignisse in Komponentengruppen der Größe 4 und ein Ereignis in einer Komponentengruppen der Größe 3 aufgetreten.

Datensatz 4 basiert auf Datensatz 3.

- Es wurden alle Ereignisse, die in Komponentengruppen einer abweichenden Größe $r \neq 3$ aufgetreten sind, entfernt.
- Die Beobachtungszeit an Komponentengruppen der Größe 3 wurde anhand der in den deutschen Anlagen erfassten Beobachtungszeiten zu Komponentengruppen

der Größen 3 und 4 zu (maximal) 5 % der Beobachtungszeit von Datensatz 3 abgeschätzt.

Datensatz 4 enthält ein Ereignis.

7.2 Vergleich der GVA-Modelle und Schätzalgorithmen mit dem Kopplungsmodell

In diesem Abschnitt werden die Schätzergebnisse unter Verwendung der neu entwickelten GVA-Modelle A und B mit den Schätzergebnissen bei Verwendung des Kopplungsmodells verglichen.

Es wurden jeweils die Schätzergebnisse vor und nach der Berücksichtigung der weiteren Unsicherheitsquellen (siehe Abschnitt 2.2.2) angegeben und diskutiert.

In Abb. 7.1 sind die Schätzergebnisse für Datensatz 1 für die Modelle A, B und das Kopplungsmodell dargestellt. Dabei wurden die Ergebnisse vor der Berücksichtigung der weiteren Unsicherheitsquellen aufgeführt. Es sind jeweils die Erwartungswerte (gefüllte Kreise) und das symmetrische 95 %-Konfidenzintervall [2.5 %-Quantil, 97.5 %-Quantil] als Fehlerbalken dargestellt.

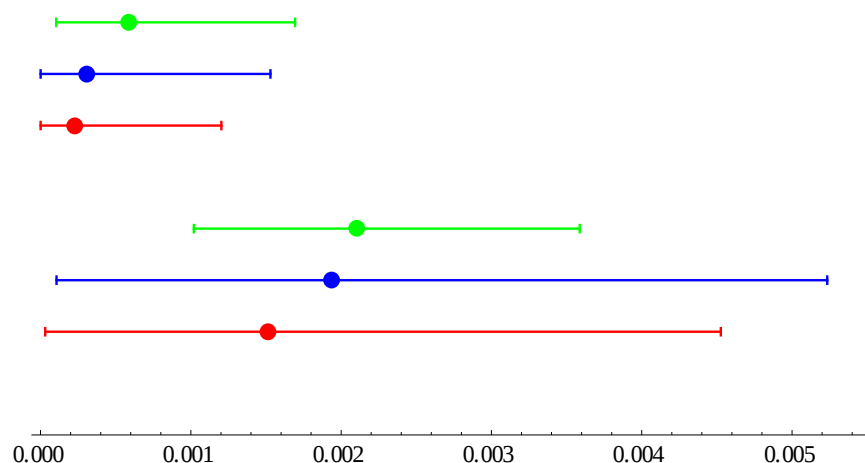


Abb. 7.1 Schätzergebnisse für $q_{4\setminus 4}$ (oben) und $q_{2\setminus 4}$ (unten) für Datensatz 1 für Modelle A (blau), B (rot) und das Kopplungsmodell (grün) ohne Berücksichtigung der weiteren Unsicherheitsquellen nach Abschnitt 2.2.2, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind

In Abb. 7.2 sind die Schätzergebnisse nach Berücksichtigung weiterer Unsicherheitsquellen dargestellt. Man beachte die gegenüber Abb. 7.1 abweichende Skala der Abszisse.

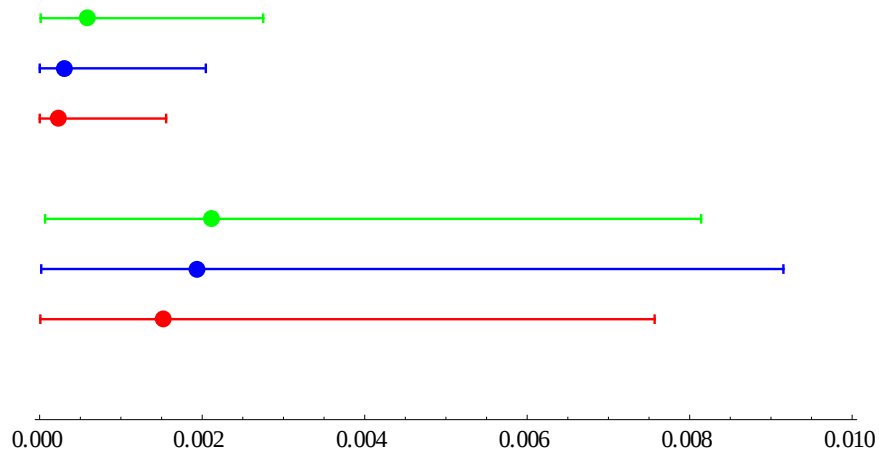


Abb. 7.2 Schätzergebnisse für $q_{4\setminus 4}$ (oben) und $q_{2\setminus 4}$ (unten) für Datensatz 1 für Modelle A (blau), B (rot) und das Kopplungsmodell (grün) mit Berücksichtigung der weiteren Unsicherheitsquellen nach Abschnitt 2.2.2, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind

Die Erwartungswerte aller Modelle liegen innerhalb der Konfidenzintervalle aller anderen Modelle.

Für den Fall $q_{2\setminus 4}$, wo viel empirische Information vorliegt, sind die Abweichungen der Ergebnisse recht gering. Die maximale Abweichung des Erwartungswerts ist weniger als 40 % zwischen Modell B und dem Kopplungsmodell und weniger als 10 % zwischen Modell A und dem Kopplungsmodell. Sie sind damit vergleichbar zwischen den Unterschieden zwischen Modell A und B von 30 %, die nur auf den verschiedenen a priori zurückzuführen sind. Die Abweichungen der oberen Grenze des Konfidenzintervalls sind auch vor Berücksichtigung weiterer Unsicherheitsquellen stets kleiner als 25 %. Nach Berücksichtigung weiterer Unsicherheitsquellen wird dieser Unterschied noch kleiner.

Für den Fall $q_{4\setminus 4}$, wo weniger empirische Information vorliegt, da keine Ereignisse mit mehr als zwei Ausfällen, sondern nur mit Schädigungen beobachtet wurden, sind die Abweichungen der Ergebnisse größer. Die maximale Abweichung des Erwartungswerts ist etwa 160 % zwischen dem Kopplungsmodell und Modell B und 95 % zwi-

schen dem Kopplungsmodell und Modell A. Die oberen Grenzen des Konfidenzintervalls unterscheiden sich um bis zu 60 %. Diese größeren Unterschiede lassen sich damit erklären, dass die Ergebnisse eine Extrapolation darstellen, die natürlich eine stärkere Abhängigkeit von den Modellannahmen und den a priori-Annahmen aufweist.

Nicht nur die dargestellten Charakteristika, sondern auch die Verteilungen selbst haben einen sehr großen Überlapp und insbesondere nach Berücksichtigung der weiteren Unsicherheitsquellen nach Abschnitt 2.2.2 eine ähnliche Gestalt, wie beispielhaft die Abb. 7.3 und Abb. 7.4 zeigen.

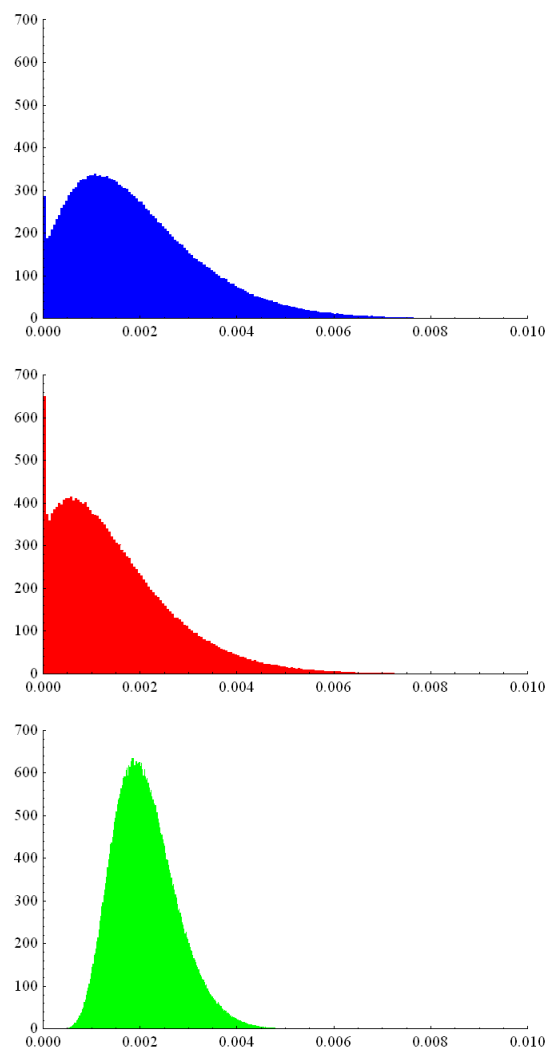


Abb. 7.3 Histogramme der Ergebnisse der Monte-Carlo-Rechnungen für $q_{2\backslash 4}$ für Datensatz 1 für Modell A (blau), B (rot) und das Kopplungsmodell (grün) ohne Berücksichtigung der weiteren Unsicherheitsquellen

Der Peak bei 0 der Schätzergebnisse bei Modell A und B in Abb. 7.3 resultiert aus der Interpretationsmöglichkeit, dass keine (2 von 4)-Ereignisse aufgetreten sind. Dies ist nach den Expertenschätzungen sehr unwahrscheinlich, aber möglich.

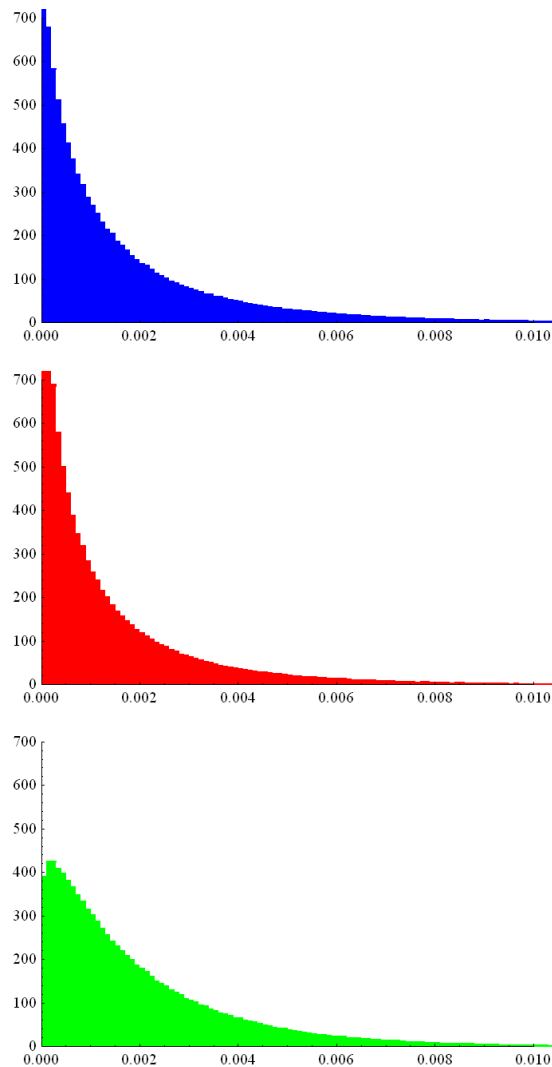


Abb. 7.4 Histogramme der Ergebnisse der Monte-Carlo-Rechnungen für $q_{2\backslash 4}$ für Datensatz 1 für Modell A (blau), B (rot) und das Kopplungsmodell (grün) mit Berücksichtigung der weiteren Unsicherheitsquellen nach Abschnitt 2.2.2

Im Folgenden werden die Schätzergebnisse für den Datensatz 2 dargestellt, der repräsentativ für Populationen mit einer geringen Anzahl von Ereignissen ist. In Abb. 7.5 sind die Ergebnisse vor der Berücksichtigung der weiteren Unsicherheitsquellen für die Modelle A, B und das Kopplungsmodell dargestellt. Es sind jeweils wiederum die Erwartungswerte (gefüllte Kreise) und die symmetrischen 95 %-Konfidenzintervalle dargestellt.

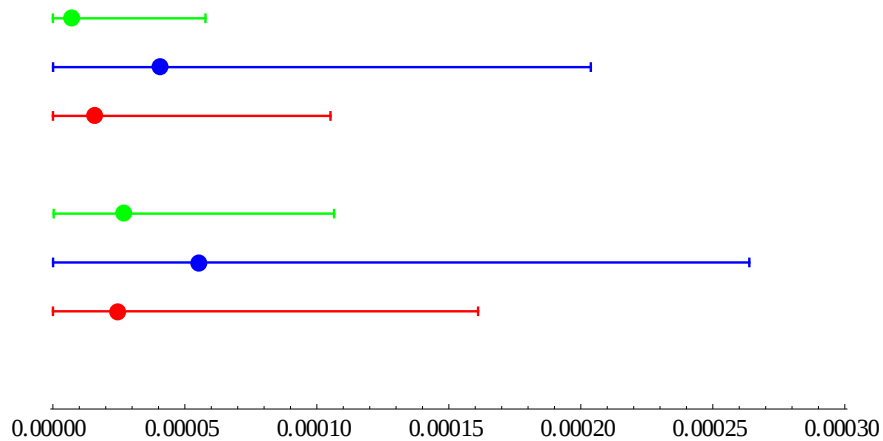


Abb. 7.5 Schätzergebnisse für $q_{4\setminus 4}$ (oben) und $q_{2\setminus 4}$ (unten) für Datensatz 2 für Modell A (blau), B (rot) und das Kopplungsmodell (grün) ohne Berücksichtigung der weiteren Unsicherheitsquellen, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind

In Abb. 7.6 sind die Schätzergebnisse nach der Berücksichtigung der weiteren Unsicherheitsquellen dargestellt.

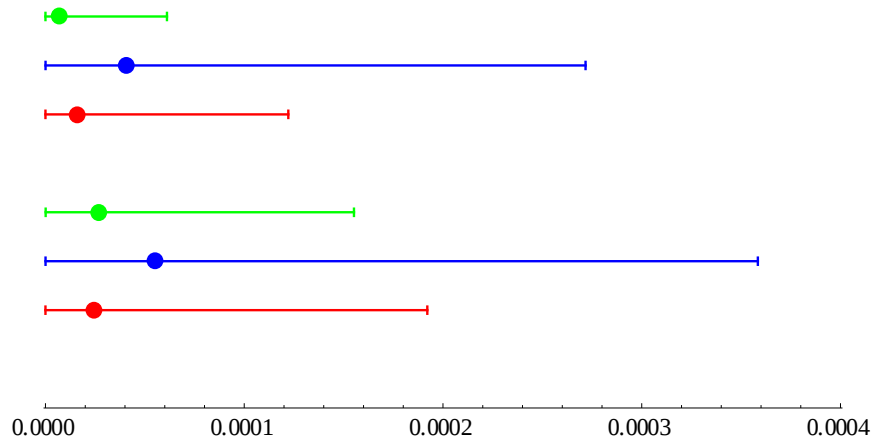


Abb. 7.6 Schätzergebnisse für $q_{4\setminus 4}$ (oben) und $q_{2\setminus 4}$ (unten) für Datensatz 2 für die Modelle A (blau), B (rot) und das Kopplungsmodell (grün) mit Berücksichtigung der weiteren Unsicherheitsquellen, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind

Hier ergeben sich die größten Unterschiede. Die Erwartungswerte der Schätzung mit dem Kopplungsmodell für $q_{4\setminus 4}$ sind relativ niedrig. Dies lässt sich mit der Annahme des Kopplungsmodells erklären, dass die beobachteten Ereignisse repräsentativ sind für alle möglichen GVA-Phänomene (siehe Abschnitt 2.1). Hier wurde nur ein einziges Ereignis beobachtet. Dass es auch noch andere GVA-Phänomene geben könnte, die andere Verteilungen von Ausfallkombinationen aufweisen, aber wegen ihrer geringen Rate nicht in der Beobachtungszeit aufgetreten sind, ist im Kopplungsmodell im Gegensatz zu Modellen A und B nicht berücksichtigt. Allerdings ist die Abweichung zwischen Kopplungsmodell und Modell B (Faktor von etwa 2.3) vergleichbar mit der Abweichung zwischen Modellen A und B (Faktor von etwa 2.6). Bei den Schätzungen von $q_{2\setminus 4}$ sind die Unterschiede noch geringer; hier stimmen die Erwartungswerte der Schätzungen von Modell B und dem Kopplungsmodell sehr gut überein. Die oberen Grenzen des Konfidenzintervalls unterscheiden sich zwischen Kopplungsmodell und Modell B und Modell B und Modell A um etwa den Faktor 2. Durch die Berücksichtigung der weiteren Unsicherheitsquellen wird dies nicht wesentlich verändert.

7.3 Mapping

In diesem Abschnitt wird das konservative Mapping untersucht. Dieses kann, wie oben diskutiert, Teil einer konservativen Vorgehensweise sein, wenn für eine in der PSA zu

modellierende Komponentengruppengröße keine oder nur sehr geringe Betriebserfahrung vorliegt.

Dazu wird Datensatz 4, der nur ein Ereignis in einer Komponentengruppe der Größe 3 enthält, verglichen mit Datensatz 4, der neben dem Ereignis in einer Komponentengruppe der Größe 3 auch 15 Ereignisse in einer Komponentengruppe der Größe 4 enthält und eine 20-mal größere Beobachtungszeit aufweist. Die Ereignisse mit einer Komponentengruppengröße 4 wurden dem konservativen Mapping Down unterworfen, d. h. in jeder Expertenbewertung wurde die am wenigsten geschädigte Komponente „gelöscht“ (siehe Abschnitt 6.2.3, „Konservatives Mapping down“).

Hier wurde nur Modell A betrachtet, da Modell B, wie oben diskutiert, nicht konservative Schätzabweichungen zeigen kann und sein Einsatz daher im Rahmen einer konservativen Vorgehensweise nicht sinnvoll ist.

Zum Vergleich wurde das Kopplungsmodell auf Datensatz 3 angewandt. Hier ist kein separates Mapping erforderlich, da es automatisch im Modell erfolgt. Deshalb ist es auch nicht sinnvoll, das Kopplungsmodell auf Datensatz 4 anzuwenden, da beim Kopplungsmodell keine Notwendigkeit besteht, Ereignisse mit nicht passenden Komponentengruppengrößen zu eliminieren.

In Abb. 7.7 sind die Ergebnisse ohne Berücksichtigung der weiteren Unsicherheitsquellen dargestellt.

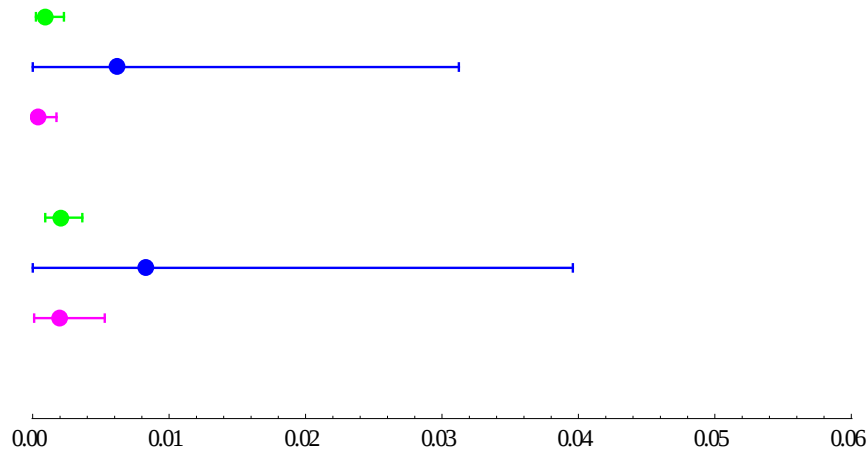


Abb. 7.7 Schätzergebnisse für $q_{3\setminus 3}$ (oben) und $q_{2\setminus 3}$ (unten) für Modell A angewandt auf Datensatz 4 (blau), für Modell A angewandt auf Datensatz 3 mit Mapping (magenta) und das Kopplungsmodell angewandt auf Datensatz 3 (grün) ohne Berücksichtigung der weiteren Unsicherheitsquellen, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind

Es ist erkennbar, dass trotz des konservativen Ansatzes die geschätzten GVA-Wahrscheinlichkeiten unter Verwendung der Betriebserfahrung an Komponentengruppen auch der Größe 4 deutlich kleiner sind, als wenn die Schätzung nur auf die Betriebserfahrung an Komponentengruppen der Größe 3 basiert. Hierfür ist entscheidend, dass Beobachtungszeit und Anzahl beobachteter Ereignisse wesentlich größer sind, was zu einer geringeren Schätzunsicherheit und zu einem geringen Einfluss der a priori-Annahmen führt. Die Ergebnisse von Kopplungsmodell und Modell A angewandt auf Datensatz 3 sind nicht signifikant verschieden; die Erwartungswerte liegen jeweils innerhalb des Konfidenzintervalls des anderen Modells. Modell A angewandt auf Datensatz 4 liefert demgegenüber signifikant größere Werte. Die obere Grenze des symmetrischen 95 %-Konfidenzintervalls ist bei der ausschließlichen Verwendung der Betriebserfahrung an Komponentengruppen der Größe 3 um ca. eine Zehnerpotenz größer als bei Verwendung von Mapping.

In Abb. 7.8 sind die Ergebnisse mit Berücksichtigung der weiteren Unsicherheitsquellen dargestellt.

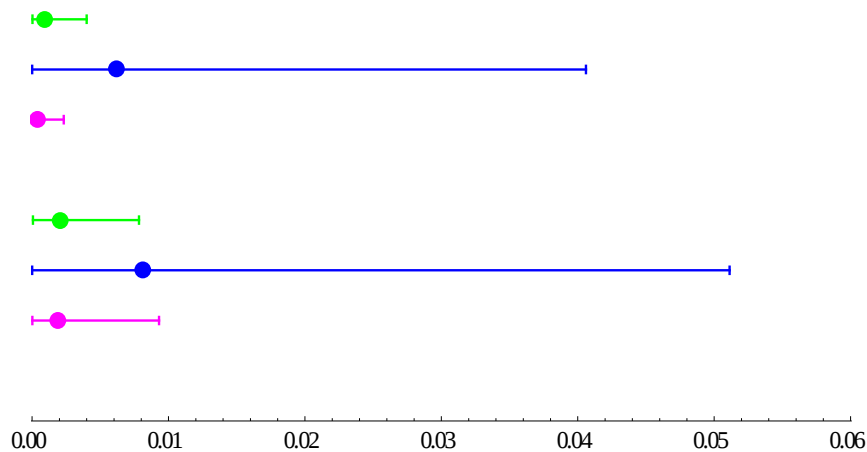


Abb. 7.8 Schätzergebnisse für $q_{3\setminus 3}$ (oben) und $q_{2\setminus 3}$ (unten) für Modell A angewandt auf Datensatz 4 (blau), für Modell A angewandt auf Datensatz 3 mit Mapping (magenta) und das Kopplungsmodell angewandt auf Datensatz 3 (grün) mit Berücksichtigung der weiteren Unsicherheitsquellen, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind

Durch die Berücksichtigung der weiteren Unsicherheitsquellen liegt nun der Erwartungswert für $q_{2\setminus 3}$ für Modell A angewandt auf Datensatz 4 innerhalb des symmetrischen 95 %-Konfidenzintervalls der Schätzung von für Modell A angewandt auf Datensatz 3, während die anderen Schätzergebnisse weiterhin signifikant verschieden sind. Qualitativ sind die Unterschiede aber vergleichbar mit denjenigen ohne Berücksichtigung der weiteren Unsicherheitsquellen.

Zusammenfassend zeigt dieses Beispiel aus der deutschen Betriebserfahrung, dass mittels des konservativen Mappings GVA-Wahrscheinlichkeiten geschätzt werden können, die deutlich weniger konservativ sind als Schätzungen, die nur auf Ereignissen an Komponentengruppen übereinstimmenden Redundanzgrades basieren. Dies setzt voraus, dass die einbezogene Zahl von Ereignissen erheblich größer ist, so dass die statistische Schätzunsicherheit wesentlich geringer ist.

8 Zusammenfassung

Im Rahmen des vom Bundesministerium für Wirtschaft und Energie (BMWi) geförderten Forschungs- und Entwicklungsvorhaben wurden Möglichkeiten zur Weiterentwicklung der Quantifizierung von GVA durch verbesserte GVA-Modelle entwickelt und diskutiert. Zuerst wurde der aktuelle Stand der GVA-Quantifizierung mit dem Kopplungsmodell unter Berücksichtigung aller Weiterentwicklungen, die seit Veröffentlichungen des Fachbands zu „Methoden zur probabilistischen Sicherheitsanalyse für Kernkraftwerke“ /FAK 05/ vorgenommen wurden, geschlossen dargestellt. Dabei wurden die Eigenschaften des Kopplungsmodells diskutiert. Diese umfassen einerseits die umfassende Berücksichtigung von Schätzunsicherheiten und die Möglichkeit, Betriebserfahrung von Komponentengruppen abweichender Größe zu verwenden, als auch unerwünschte Konvergenzeigenschaften: Bei Anwachsen der Anzahl von GVA-Ereignissen nimmt die Unsicherheit der Schätzungen der GVA-Wahrscheinlichkeiten nicht in dem Maße ab, wie es die abnehmende statistische Unsicherheit ermöglicht. Diese Eigenschaft ist mit der grundlegenden Annahme des Kopplungsmodells verbunden, dass die Komponenten bei Auftritt eines GVA-Phänomens mit einer gewissen Wahrscheinlichkeit ausfallen (dem Kopplungsparameter), dieser bei verschiedenen GVA-Phänomenen im Allgemeinen verschieden ist und deshalb für alle GVA-Ereignisse unabhängig geschätzt wird. Deshalb musste, um die unerwünschten Konvergenzeigenschaften zu vermeiden, auf diese zentrale Modellannahme bei der Entwicklung fortgeschrittener GVA-Modelle verzichtet werden.

Um einen neuen Modellansatz zu gewinnen, wurden zunächst die international üblichen Vorgehensweisen zum Schätzen von GVA (u. a. das Alpha-Faktor- und das Beta-Faktor-Modell) beschrieben. Verfahren zur Schätzung der Modellparameter unter Verwendung statistischer Methoden von Bayes wurden in einer einheitlichen Form dargestellt. Diese Verfahren lassen sich nicht unmittelbar auf die deutsche Betriebserfahrung übertragen, da einerseits für die Schätzung benötigte Informationen nicht vorliegen, andererseits bei der weiterentwickelten Modellierung die umfassende Einbeziehung der verschiedenen Unsicherheitsquellen, die die bisherige Vorgehensweise kennzeichnet, erhalten bleiben soll. Deshalb wurden Kriterien entwickelt, die der Entwicklung eines für die deutsche Betriebserfahrung geeigneten Modells zugrunde liegen sollen. Basierend auf diesen Kriterien wurden drei Modellansätze entwickelt. Im ersten Modellansatz (Modell A) werden GVA mit verschiedenen Ausfallkombinationen als unabhängige Basisereignisse angesehen. Die Raten dieser Ereignisse stellen die Modellpara-

meter dar. In den beiden weiteren Modellansätzen (Modelle B und C) werden generische Zustände „GVA“ bzw. „GVA-Phänomen“ postuliert, die mit einer Rate auftreten. Aus diesem Zustand geht das Modell in Endzustände über, die den verschiedenen Ausfallkombinationen entsprechen. Die entsprechenden bedingten Wahrscheinlichkeiten sind die weiteren Modellparameter. Diese Modellvorstellung ähnelt dem Alpha-Faktor-Modell. Ein wesentlicher Unterschied ist allerdings, dass nur Ausfälle mit systematischer Ursache und keine Einzelfehler beschrieben werden. Deshalb sind zum Schätzen der Modellparameter aus der Betriebserfahrung auch keine Einzelfehler erforderlich. Bayes'sche Schätzverfahren wurden für die drei Modelle hergeleitet. Bei ihnen werden die verschiedenen Unsicherheitsquellen in gleicher Qualität wie beim Kopplungsmodell berücksichtigt. Untersuchungen der Konvergenz der Modellparameter zeigen, dass die Modelle B und C unter bestimmten Randbedingungen eine starke Unterschätzung der Wahrscheinlichkeiten von einzelnen GVA-Kombinationen (z. B. komplette GVA) zeigen können. Als Ursache wurde die langsame Konvergenz der Modellparameter identifiziert, die die Verteilung der Ereignisse auf die verschiedenen Ausfallkombinationen beschreibt, während die Konvergenz der Gesamtrate schneller ist. Demgegenüber treten bei Modell A nur Überschätzungen auf. Deshalb ist nur mit Modell A eine konservative Schätzung der GVA-Wahrscheinlichkeiten möglich.

Die entwickelten GVA-Modelle sind so gestaltet, dass zur Quantifizierung unmittelbar nur GVA-Ereignisse verwendet werden können, die dieselbe Größe haben wie diejenige GVA-Gruppe, für die GVA quantifiziert werden sollen. Da nicht für alle Gruppengrößen ausreichend Betriebserfahrung vorliegt, sind separate Algorithmen erforderlich, um die Ereignisse zwischen Komponentengruppen verschiedener Größe zu übertragen. Für dieses Mapping wurden verschiedene, teilweise neue Ansätze, aufgeführt. Dabei ist ein neuer Ansatz für das Mapping Up besonders hervorzuheben, der nur auf der Annahme basiert, dass eine kleine Komponentengruppe als zufällige Untermenge der Komponenten einer großen angesehen werden kann, die nicht vollständig beobachtet wird. Mittels Bayes'scher statistischer Methoden können GVA-Wahrscheinlichkeiten in der großen Komponentengruppe berechnet werden. Die mathematischen Beziehungen lassen sich analytisch ausdrücken, d. h. man ist nicht auf Monte-Carlo-Verfahren zur Implementation angewiesen.

Für einen Spezialfall wurde die Konvergenz der geschätzten Parameter gegen ihre wahren Werte demonstriert. Im Zusammenhang mit diesem Verfahren wurde auch die Kompatibilität der Annahme, dass eine kleine Komponentengruppe als Teil einer großen angesehen werden kann, mit den für die Schätzalgorithmen verwendeten a priori-

Verteilungen untersucht mit dem Ergebnis, dass sie nicht kompatibel sind. Es wurden denkbare Lösungsmöglichkeiten diskutiert; jedoch ist nicht erkennbar, wie ein kompatibler a priori konstruiert werden könnte. Diese Inkompatibilität betrifft nicht nur das neu entwickelte Verfahren, sondern auch weitere Mappingverfahren, die auf dieser Grundannahme basieren (probabilistisch-kombinatorisches Mapping Down) und international häufig zusammen mit dem Alpha-Faktor-Modell angewandt werden. Bei dieser Vorgehensweise existiert ebenfalls die genannte Inkompatibilität, die zu einer inneren Widersprüchlichkeit führt. Das Kopplungsmodell ist nicht betroffen, da es keine komponentengruppengrößenspezifischen Parameter aufweist.

Es wurde diskutiert, wie – abgesehen von innerer Widerspruchsfreiheit – die verschiedenen Mappingalgorithmen bewertet werden können. Hierbei ist problematisch, dass keine ausreichende empirische Evidenz vorhanden ist, wie sich GVA-Phänomene in Komponentengruppen verschiedener Größe tatsächlich auswirken. Deshalb wurden Kriterien für die Konservativität entwickelt. Diese werden vollständig von den Verfahren „konservatives Mapping Down“ und „konservatives Mapping Up“ erfüllt, die sich als „weglassen“ der am schwächsten geschädigten Komponenten bzw. „duplizieren“ der am stärksten geschädigten Komponente charakterisieren lassen. Wegen der Konservativität wird die mit dem Mapping verbundene Unsicherheit in diesen Verfahren nicht explizit ausgewiesen.

Die entwickelten Verfahren wurden anhand von deutscher Betriebserfahrung erprobt. Dafür wurden zwei Datensätze zusammengestellt, die Populationen mit sehr wenigen beobachteten GVA-Ereignissen und Populationen mit vielen GVA-Ereignissen repräsentieren. In diesen Datensätzen sind nur Ereignisse an GVA-Gruppen einer Größe (vier) enthalten, um die Schätzverfahren unabhängig vom Mapping vergleichen zu können. Vergleiche der Schätzergebnisse der Modelle A und B mit dem Kopplungsmodell zeigen, dass die Ergebnisse sehr ähnlich sind. Die Mittelwerte der Ergebnisverteilungen liegen jeweils innerhalb der 95 %-Konfidenzintervalle aller anderen Verfahren. Dies gilt sowohl vor als auch nach der Einbeziehung der verbleibenden Unsicherheiten. Die Schätzungen mit dem Kopplungsmodell sind nicht signifikant von denjenigen mit den neuen Modellen verschieden. Insofern lässt sich aus den Ergebnissen keine Notwendigkeit ableiten, bei der zurzeit vorliegenden deutschen Betriebserfahrung die bisherige Vorgehensweise zur Quantifizierung mit dem Kopplungsmodell zu verändern.

Zusätzlich wurde das Verfahren zum konservativen Mapping in Verbindung mit Modell A angewandt, um die GVA-Wahrscheinlichkeiten in Komponentengruppen der Größe 3

konservativ zu schätzen. Die Ergebnisse wurden mit Schätzungen anhand der Betriebserfahrung in Komponentengruppen nur der Größe 3 verglichen, die nur ein Ereignis beinhaltet. Es zeigt sich, dass trotz der Konservativität die Schätzungen unter Verwendung des Mappings deutlich geringe Werte ergaben, da die Schätzungenauigkeit aufgrund der geringen Ereigniszahl bei Verwendung der Betriebserfahrung in Komponentengruppen nur der Größe 3 sehr hoch ist. Die Ergebnisse der konservativen Vorgehensweise sind auch vergleichbar zu den mit dem Kopplungsmodell erzielten Ergebnissen.

Modell A erlaubt es somit, eindeutig konservative Schätzungen von GVA-Wahrscheinlichkeiten zu berechnen. In Verbindung mit dem konservativen Mapping kann auch Betriebserfahrung in Komponentengruppen abweichender Größe einbezogen werden.

Aus den erzielten Ergebnissen kann zum Einen abgeleitet werden, dass das Kopplungsmodell für die gegenwärtig in Deutschland vorhandene Betriebserfahrung mit GVA als adäquat angesehen werden kann, und zum Anderen, dass weiterer Forschungsbedarf resultiert. Dies betrifft einerseits die vertiefte Untersuchung von GVA-Entstehung und Entdeckung, um Erkenntnisse zu erlangen, die eine empirisch fundierte Bewertung des Mappings erlauben. Andererseits sollte in Bezug auf die Kompatibilität von a priori-Verteilungen mit den dem Mapping zugrunde liegenden Annahmen mathematisch weiterführend untersucht werden, inwieweit ein innerer Widerspruch von probabilistisch-kombinatorischem Mapping und den modellbasierten Schätzverfahren vermieden werden kann.

Literaturverzeichnis

- /BER 80/ Berger, J. O.: Statistical Decision Theory and Bayesian Analysis, 2nd Edition, Springer, New York, New York, 1980.
- /BMU 05/ Bundesministerium für Umwelt, Naturschutz und Reaktorsicherheit (BMU): Bekanntmachung des Leitfadens zur Durchführung der „Sicherheitsüberprüfung gemäß § 19a des Atomgesetzes – Leitfaden Probabilistische Sicherheitsanalyse –“ für Kernkraftwerke in der Bundesrepublik Deutschland, Bundesanzeiger Jg. 57, Nr. 207a, 2005.
- /BOX 73/ Box, G. E. P., E. G. Tiao: Bayesian Inference in Statistical Analysis, Addison-Wesley, Reading, Massachusetts, 1973.
- /FAK 05/ Facharbeitskreis (FAK) Probabilistische Sicherheitsanalyse für Kernkraftwerke: Methoden zur probabilistischen Sicherheitsanalyse für Kernkraftwerke, Stand: August 2005, BfS-SCHR-37/05, Bundesamt für Strahlenschutz (BfS), Salzgitter, Germany, Oktober 2005, <http://doris.bfs.de/jspui/handle/urn:nbn:de:0221-201011243824>.
- /FAK 05a/ Facharbeitskreis (FAK) Probabilistische Sicherheitsanalyse für Kernkraftwerke: Daten zur Quantifizierung von Ereignisablaufdiagrammen und Fehlerbäumen, Stand: August 2005, BfS-SCHR-38/05, Bundesamt für Strahlenschutz (BfS), Salzgitter, Germany, Oktober 2005, <https://doris.bfs.de/jspui/handle/urn:nbn:de:0221-201011243838>.
- /KRE 01/ Kreuser, A., J. Peschke: Coupling Model: A Common-Cause-Failure Model with Consideration of Interpretation Uncertainties, Nuclear Technology Vol. 136, 2001.
- /KRE 06/ Kreuser, A., J. Peschke, J. C. Stiller: Further Development of the Coupling Model, Kerntechnik, Vol. 71, 2006.
- /MOS 88/ Mosleh, A., et al.: Procedures for Treating Common Cause Failures in Safety and Reliability Studies, NUREG/CR-4780, Vol. 1, 1988.

- /MOS 89/ Mosleh, A., et al.: Procedures for Treating Common Cause Failures in Safety and Reliability Studies, NUREG/CR-4780, Vol. 2, 1989.
- /MOS 98/ Mosleh, A., D. M. Rasmuson, F. M. Marshall: Guidelines on Modeling Common-Cause-Failures in Probabilistic Risk Assessment, NUREG/CR-5485, 1998.
- /STI 08/ Stiller, J. C., A. Kreuser, C. Versteegen: Consideration of Additional Uncertainties in the Coupling Model for the Estimation of Unavailabilities due to Common Cause Failures, in: Proceedings of the 9th International Conference on Probabilistic Safety Assessment and Management, Hong Kong, 2008.
- /STI 09/ Stiller, J. C., J. Peschke: Konsistente Berücksichtigung der Unsicherheit bezüglich der Rate von GVA-Ereignissen bei der Anwendung des Kopplungsmodells, GRS-A-3466, Gesellschaft für anlagen- und Reaktorsicherheit (GRS) mbH, Köln, 2009.
- /WIE 07/ Wierman, T. E., D. M. Rasmuson, A. Mosleh: Common-Cause Failure Database and Analysis System: Event Data Collection, Classification, and Coding, NUREG/CR-6268 Rev.1, 2007.
- /WOL 12/ Wolfram Research, Inc., Mathematica, Version 9.0.1.0, 2013.

Abbildungsverzeichnis

Abb. 2.1	Relative Überschätzung $\langle q_{r \setminus r} \rangle / q_{r \setminus r}^*$ für $r \in [2, 128]$	23
Abb. 3.1	Veranschaulichung der Parameter des Beta-Faktor Modells	28
Abb. 3.2	Veranschaulichung des MGL-Parameter	30
Abb. 4.1	Struktur des Alpha-Faktor-Modells.....	46
Abb. 5.1	Modellstruktur Modell A	52
Abb. 5.2	Modellstruktur Modell B	54
Abb. 5.3	Modellstruktur Modell C	56
Abb. 5.4	Erwartungswerte $\langle q_{4 \setminus 4} \rangle$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange).....	66
Abb. 5.5	Erwartungswerte $\langle q_{k \setminus 4} \rangle$ für $k \neq 4$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange)	67
Abb. 5.6	Erwartungswerte $\langle q_{8 \setminus 8} \rangle$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange).....	68
Abb. 5.7	Erwartungswerte $\langle q_{k \setminus 8} \rangle$ für $k \neq 8$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange)	68
Abb. 5.8	Erwartungswerte $\langle q_{128 \setminus 128} \rangle$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange); die Kurven für Modell B und Modell C liegen fast übereinander	69
Abb. 5.9	Erwartungswerte $\langle q_{k \setminus 128} \rangle$ für $k \neq 128$ in Abhängigkeit von T für $\iota = 1$ und $t = 0.0001$ für Modell A (blau), Modell B (rot) und Modell C (orange); die Kurven für Modell B und Modell C liegen fast übereinander	69
Abb. 5.10	Erwartungswerte $\langle \lambda \rangle = \langle \zeta \rangle$ in Abhängigkeit von T für $\iota = 1$	71
Abb. 5.11	Erwartungswerte $\langle \omega_{128 \setminus 128} \rangle$ bzw. $\langle x_{128 \setminus 128} \rangle$ in Abhängigkeit von T für $\iota = 1$ für Modell B (rot) und Modell C (orange); die Kurven liegen fast übereinander.....	71

Abb. 5.12	Erwartungswerte $\langle \omega_{k \setminus 128} \rangle$ bzw. $\langle x_{k \setminus 128} \rangle$ für $k \neq 128$ in Abhängigkeit von T für $\iota = 1$ für Modell B (rot) und Modell C (orange); die Kurven liegen fast übereinander	72
Abb. 6.1	Indirekte Abhängigkeit der Beobachtungen M in der „kleinen“ Komponentengruppe von den Modellparametern X in der „großen“ Komponentengruppe über nicht beobachtete Ereignisse M in der „großen“ Komponentengruppe	93
Abb. 6.2	Betrachtete indirekte Abhängigkeit der Beobachtungen M in der „kleinen“ Komponentengruppe von den Modellparametern X in der „großen“ Komponentengruppe über die Modellparameter Y in der „kleinen“ Komponentengruppe	94
Abb. 6.3	Abhängigkeit der Beobachtungen M in der „kleinen“ Komponentengruppe von den Beobachtungen M in der „großen“ Komponentengruppe	95
Abb. 6.4	Erwartungswerte $x_{0 \setminus 3} = x_{1 \setminus 3}$ (grün), $x_{2 \setminus 3}$ (blau) und $x_{3 \setminus 3}$ (rot) in Abhängigkeit von z für den Fall, dass ausschließlich (2 von 2)-Ereignisse beobachtet wurden	103
Abb. 6.5	Marginalverteilungsdichten von $x_{3 \setminus 3}$ nach Beobachtung von $z = 0$ (rot), 1 (orange), 2 (magenta), 5 (braun), 10 (grün), 50 (blau) bzw. 100 (schwarz) (2 von 2)-Ereignissen	104
Abb. 6.6	Darstellung der Vorgehensweisen nach Verfahren $\aleph$ (cyan) und Verfahren $\beth$ (orange).	116
Abb. 6.7	Vergleich der Marginalverteilungen von $y_{2 \setminus 2}$ nach Verfahren $\aleph$ (schwarze Kurve) und $\beth$ (rotes Histogramm) für eine Komponentengruppe der Größe $r = 2$, die als Teil einer Komponentengruppe der Größe $r = 3$ betrachtet wird.	119
Abb. 6.8	Vergleich der Marginalverteilungen von $y_{1 \setminus 2}$ nach Verfahren $\aleph$ (schwarze Kurve) und $\beth$ (blaues Histogramm) für eine Komponentengruppe der Größe $r = 2$, die als Teil einer Komponentengruppe der Größe $r = 3$ betrachtet wird.	120
Abb. 6.9	Vergleich der Marginalverteilungen von $y_{2 \setminus 2}$ nach Verfahren $\aleph$ (schwarze Kurve) und $\beth$ (rotes Histogramm) für eine Komponentengruppe der Größe $r = 2$, die als Teil einer Komponentengruppe der Größe $r = 4$ betrachtet wird	121

Abb. 6.10	Vergleich der Marginalverteilungen von $y_{1\backslash 2}$ nach Verfahren $\mathfrak{K}$ (schwarze Kurve) und $\mathfrak{L}$ (blaues Histogramm) für eine Komponentengruppe der Größe $r = 2$, die als Teil einer Komponentengruppe der Größe $r = 4$ betrachtet wird	122
Abb. 6.11	Vergleich der Marginalverteilungen von $y_{2\backslash 2}$ nach Verfahren $\mathfrak{K}$ (schwarze Kurve) und $\mathfrak{L}$ (rotes Histogramm) für eine Komponentengruppe der Größe $r = 2$, die als Teil einer Komponentengruppe der Größe $r = 128$ betrachtet wird	123
Abb. 6.12	Vergleich der Marginalverteilungen von $y_{2\backslash 2}$ nach Verfahren $\mathfrak{K}$ (schwarze Kurve) und $\mathfrak{L}$ (rotes Histogramm) für eine Komponentengruppe der Größe $r = 2$, die als Teil einer Komponentengruppe der Größe $r = 128$ betrachtet wird	123
Abb. 7.1	Schätzergebnisse für $q_{4\backslash 4}$ (oben) und $q_{2\backslash 4}$ (unten) für Datensatz 1 für Modelle A (blau), B (rot) und das Kopplungsmodell (grün) ohne Berücksichtigung der weiteren Unsicherheitsquellen nach Abschnitt 2.2.2, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken)dargestellt sind.....	133
Abb. 7.2	Schätzergebnisse für $q_{4\backslash 4}$ (oben) und $q_{2\backslash 4}$ (unten) für Datensatz 1 für Modelle A (blau), B (rot) und das Kopplungsmodell (grün) mit Berücksichtigung der weiteren Unsicherheitsquellen nach Abschnitt 2.2.2, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind.....	134
Abb. 7.3	Histogramme der Ergebnisse der Monte-Carlo-Rechnungen für $q_{2\backslash 4}$ für Datensatz 1 für Modell A (blau), B (rot) und das Kopplungsmodell (grün) ohne Berücksichtigung der weiteren Unsicherheitsquellen	135
Abb. 7.4	Histogramme der Ergebnisse der Monte-Carlo-Rechnungen für $q_{2\backslash 4}$ für Datensatz 1 für Modell A (blau), B (rot) und das Kopplungsmodell (grün) mit Berücksichtigung der weiteren Unsicherheitsquellen nach Abschnitt 2.2.2	136
Abb. 7.5	Schätzergebnisse für $q_{4\backslash 4}$ (oben) und $q_{2\backslash 4}$ (unten) für Datensatz 2 für Modell A (blau), B (rot) und das Kopplungsmodell (grün) ohne Berücksichtigung der weiteren Unsicherheitsquellen, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind.....	137

Abb. 7.6	Schätzergebnisse für $q_{4\setminus 4}$ (oben) und $q_{2\setminus 4}$ (unten) für Datensatz 2 für die Modelle A (blau), B (rot) und das Kopplungsmodell (grün) mit Berücksichtigung der weiteren Unsicherheitsquellen, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind.....	138
Abb. 7.7	Schätzergebnisse für $q_{3\setminus 3}$ (oben) und $q_{2\setminus 3}$ (unten) für Modell A angewandt auf Datensatz 4 (blau), für Modell A angewandt auf Datensatz 3 mit Mapping (magenta) und das Kopplungsmodell angewandt auf Datensatz 3 (grün) ohne Berücksichtigung der weiteren Unsicherheitsquellen, wobei. jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken) dargestellt sind	140
Abb. 7.8	Schätzergebnisse für $q_{3\setminus 3}$ (oben) und $q_{2\setminus 3}$ (unten) für Modell A angewandt auf Datensatz 4 (blau), für Modell A angewandt auf Datensatz 3 mit Mapping (magenta) und das Kopplungsmodell angewandt auf Datensatz 3 (grün) mit Berücksichtigung der weiteren Unsicherheitsquellen, wobei jeweils Erwartungswert (Kreis) und symmetrisches 95 %-Konfidenzintervall (Fehlerbalken)dargestellt sind.....	141

Tabellenverzeichnis

Tab. 2.1	Schädigungskategorien und Schädigungswerte	14
Tab. 2.2	Berechnungsvorschriften für eine Komponentengruppe mit drei Komponenten	15
Tab. 3.1	Übersicht der Modelle und ihrer Parameter.....	34
Tab. 5.1	Erwartete GVA-Wahrscheinlichkeit bei Nullfehlerstatistik.....	64
Tab. 5.2	Zu schätzende Modellparameter.....	80
Tab. 6.1	Schädigungsvektor und jeweilige Wahrscheinlichkeit	89
Tab. 6.2	Berechnungsvorschriften für das Mapping von einer Komponentengruppe mit drei Komponenten auf eine Komponentengruppe mit zwei Komponenten.....	90
Tab. 6.3	Ergebnisse für das Mapping eines Ereignisses in einer Komponentengruppe mit drei Komponenten, bei der ein Ausfall ($d_1 = 1$), eine starke Schädigung ($d_2 = 1/2$) und eine sehr schwache Schädigung ($d_3 = 1/100$) beobachtet wurde, auf eine Komponentengruppe mit zwei Komponenten.....	91
Tab. 6.4	Schädigungsvektor und jeweilige Wahrscheinlichkeit bei probabilistisch-kombinatorischem und konservativem Mapping Down ...	110
Tab. 6.5	Schädigungsvektor und jeweilige Wahrscheinlichkeit bei heuristischem und bei konservativem Mapping Up	110
Tab. 6.6	Schädigungsvektor der Größe 4 und jeweilige Wahrscheinlichkeit	111
Tab. 6.7	Schädigungsvektor der Größe 5 und jeweilige Wahrscheinlichkeit	112
Tab. 6.8	Schädigungsvektor der Größe 5 und jeweilige Wahrscheinlichkeit	113
Tab. 6.9	Schädigungsvektor der Größe 7 und jeweilige Wahrscheinlichkeit	113
Tab. 6.10	Vergleich der Charakteristika der Marginalverteilungen von $y_{2 \setminus 2}$ nach Verfahren $\aleph$ und $\beth$	124

Tab. 6.11	Vergleich der Charakteristika der Marginalverteilungen von $y_{1\setminus 2}$ nach Verfahren $\aleph$ und $\beth$	124
Tab. 6.12	Vergleich der Mappingalgorithmen	129

A Anhang: Kombinatorische Formeln für das Mapping

In diesem Anhang sind die kombinatorischen Formeln angegeben, die dem Mapping unter der Annahme zugrunde liegen, dass eine Komponentengruppe der Größe r sich verhält wie eine Untergruppe einer größeren Komponentengruppe der Größe $\hat{r} > r$. Im Folgenden sind die Wahrscheinlichkeiten in der Komponentengruppe der Größe $\hat{r}$ als X und die Wahrscheinlichkeiten in der Komponentengruppe der Größe r als Y bezeichnet. Es gilt $Y = (y_{0 \setminus 2}, \dots, y_{r \setminus r})$ und $X = (x_{0 \setminus \hat{r}}, \dots, x_{\hat{r} \setminus \hat{r}})$. Die Beziehung zwischen X und Y lässt sich allgemein schreiben als

$$Y = SX \quad (\text{A. 1})$$

Dabei ist S von $\hat{r}$ und r abhängig.

Fall: $r = 2$ und $\hat{r} = 3$

Es gibt $\binom{3}{2} = 3$ Möglichkeiten, $r = 2$ von $\hat{r} = 3$ Komponenten auszuwählen. Wie bereits in Abschnitt 6.2.2 dargestellt, gilt:

$$S = \begin{pmatrix} 1 & \frac{1}{3} & 0 & 0 \\ 0 & \frac{2}{3} & \frac{2}{3} & 0 \\ 0 & 0 & \frac{1}{3} & 1 \end{pmatrix} \quad (\text{A. 2})$$

Der Rang von S ist 3. Der Kern von S ist gegeben durch

$$\mathcal{K}(S) = \{V \in \mathbb{R}^4 | V = \psi (-1, \quad 3, \quad -3, \quad 1), \psi \in \mathbb{R}\} \quad (\text{A. 3})$$

Fall: $r = 3$ und $\hat{r} = 4$

Es gibt $\binom{4}{3} = 4$ Möglichkeiten, $r = 3$ von $\hat{r} = 4$ Komponenten auszuwählen. Es gilt:

$$S = \begin{pmatrix} 1 & \frac{1}{4} & 0 & 0 & 0 \\ 0 & \frac{3}{4} & \frac{1}{2} & 0 & 0 \\ 0 & 0 & \frac{1}{2} & \frac{3}{4} & 0 \\ 0 & 0 & 0 & \frac{1}{4} & 1 \end{pmatrix} \quad (\text{A. 4})$$

Der Rang von S ist 4. Der Kern von S ist gegeben durch

$$\mathcal{K}(S) = \{V \in \mathbb{R}^5 | V = \psi (1, -4, 6, 4, -1), \psi \in \mathbb{R}\} \quad (\text{A. 5})$$

Fall: $r = 2$ und $\hat{r} = 4$

Es gibt $\binom{4}{2} = 6$ Möglichkeiten, $r = 2$ von $\hat{r} = 4$ Komponenten auszuwählen. Es gilt:

$$S = \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{6} & 0 & 0 \\ 0 & \frac{1}{2} & \frac{2}{3} & \frac{1}{2} & 0 \\ 0 & 0 & \frac{1}{2} & \frac{1}{6} & 1 \end{pmatrix} \quad (\text{A. 6})$$

Der Rang von S ist 3. Der Kern von S ist gegeben durch

$$\mathcal{K}(S) = \{V \in \mathbb{R}^5 | V = \psi (-3, 8, -6, 0, 3) + \Psi(87, -114, -180, 354, 31), \psi \in \mathbb{R}, \Psi \in \mathbb{R}\} \quad (\text{A. 7})$$

wobei orthogonale Basisvektoren des Kerns mit ganzzahligen Komponenten gewählt wurden.

Fall: $r = 2$ und beliebiges $\hat{r}$

Allgemeine Formeln lassen sich einfach begründen für den Fall $r = 2$, der in Abschnitt 6.3 von Bedeutung ist.

Es gibt $\binom{\hat{r}}{2}$ Möglichkeiten, zwei von $\hat{r}$ Komponenten auszuwählen. Wenn k Komponenten ausgefallen sind, gibt es $\binom{k}{2}$ Möglichkeiten, zwei ausgefallene Komponenten aus-

zuwählen, $\binom{\hat{r}-k}{2}$ Möglichkeiten, zwei nicht ausgefallene Komponenten auszuwählen und $k(\hat{r}-k)$ Möglichkeiten, eine ausgefallene und eine nicht ausgefallene Komponente auszuwählen. Alle Möglichkeiten sind gleich wahrscheinlich. Daraus folgt für die Elemente der Matrix S :

$$\begin{aligned}
 S_{0,k} &= \frac{\binom{\hat{r}-k}{2}}{\binom{\hat{r}}{2}} = \frac{(\hat{r}-k)^2 - \hat{r} + k}{\hat{r}^2 - \hat{r}} \\
 S_{1,k} &= \frac{k(\hat{r}-k)}{\binom{\hat{r}}{2}} = \frac{2k(\hat{r}-k)}{\hat{r}^2 - \hat{r}} \quad (\text{A. 8}) \\
 S_{2,k} &= \frac{\binom{k}{2}}{\binom{\hat{r}}{2}} = \frac{k^2 - k}{\hat{r}^2 - \hat{r}}
 \end{aligned}$$

Wie man leicht nachvollziehen kann, gilt $S_{0,k} + S_{1,k} + S_{2,k} = 1$.

Fall: Beliebige r und $\hat{r}$

Der obigen Überlegungen lassen sich wie folgt verallgemeinern: Es gibt $\binom{\hat{r}}{r}$ Möglichkeiten, r von $\hat{r}$ Komponenten auszuwählen. Wenn k der $\hat{r}$ Komponenten ausgefallen sind, gibt es $\binom{k}{r}$ Möglichkeiten, r ausgefallene Komponenten auszuwählen, $\binom{\hat{r}-k}{r}$ Möglichkeiten, r nicht ausgefallene Komponenten auszuwählen und es gibt allgemein $\binom{k}{w} \binom{\hat{r}-k}{r-w}$ Möglichkeiten, w ausgefallene Komponenten und $r-w$ nicht ausgefallene Komponenten auszuwählen ($0 \leq w \leq r$). Somit folgt:

$$S_{w,k} = \frac{\binom{k}{w} \binom{\hat{r}-k}{r-w}}{\binom{\hat{r}}{r}} \quad (\text{A. 9})$$

**Gesellschaft für Anlagen-
und Reaktorsicherheit
(GRS) mbH**

Schwertnergasse 1
50667 Köln

Telefon +49 221 2068-0

Telefax +49 221 2068-888

Forschungszentrum

85748 Garching b. München

Telefon +49 89 32004-0

Telefax +49 89 32004-300

Kurfürstendamm 200

10719 Berlin

Telefon +49 30 88589-0

Telefax +49 30 88589-111

Theodor-Heuss-Straße 4

38122 Braunschweig

Telefon +49 531 8012-0

Telefax +49 531 8012-200

www.grs.de

ISBN 978-3-944161-02-0