

Berechnungen zur Funktionalen Sicherheit

Größen, Formeln und Methoden

Thomas Brunnengräber
tbrunnengraeber@thomas-brunnengraeber.de

29. November 2025

Inhaltsverzeichnis

1	Einführung	5
2	Zuverlässigkeit und verwandte Größen	6
2.1	Zuverlässigkeit und Unzuverlässigkeit	6
2.2	Ausfalldichte und Ausfallrate	7
2.3	Badewannenkurve	8
2.4	Mittlere Zeit bis zum Ausfall (MTTF)	9
2.5	Verteilungsfunktionen	9
2.5.1	Exponentialverteilung	9
2.5.2	Weibull-Verteilung	11
2.5.3	Sterblichkeit	12
3	Mittlere Ausfallrate und mittlere Zeit bis zum Ausfall	14
3.1	MTTF bei langer Nutzung ohne vorbeugenden Tausch	15
3.2	Vorbeugender Tausch und unvollständige MTTF	16
3.3	Gefährliche und ungefährliche Ausfallmodi, gefährliche MTTF	18
4	Wiederherstellung und Verfügbarkeit	21
4.1	Reparierbarkeit	21
4.2	Diagnose, Test, Wiederherstellung	21
4.3	Verfügbarkeit und Nichtverfügbarkeit	21
4.4	Zeit zur Reparatur, MRT	24
4.5	Kontinuierliche Diagnose	25
4.6	Betriebliche und sicherheitsbezogene Verfügbarkeit	26
4.7	Ausfallrate bei Tests	26
5	Nichtverfügbarkeit von komplexen Funktionen	28
5.1	Berechnung mit Fehlerbäumen	29
5.1.1	Nichtverfügbarkeit einer UND-Verknüpfung	30
5.1.2	Nichtverfügbarkeit einer ODER-Verknüpfung	32
5.1.3	Nichtverfügbarkeit von Kombinationen von UND- und ODER-Verknüpfungen	34
5.1.4	Transiente und stationäre Berechnung, Rechnen mit Mittelwerten	35
5.2	Berechnung mit Markov Modellen	37
5.2.1	Stationäre Berechnung	39
5.2.2	Transiente Berechnung	40
6	Ausfallrate von komplexen Funktionen	43
6.1	Berechnung mit Fehlerbäumen	45
6.1.1	System ohne Redundanzen	45
6.1.2	System mit Redundanzen	46
6.1.3	Einkanalig Fail-Safe	49
6.1.4	Modellierung von regelmäßigen Tests und Diagnose	51
6.1.5	Zweikanalig Fail-Safe	55

6.1.6	Fail-Operational Systeme	58
6.1.7	Transiente und stationäre Berechnung	62
6.2	Berechnung mit Markov-Modellen	62
6.3	Erwartungswert der Ausfälle	65
7	Unzuverlässigkeit von komplexen Funktionen	66
7.1	Berechnung mit Fehlerbäumen	66
7.2	Berechnung mit Markov Modellen	69
A	Modelle für Basis-Ereignisse	70
A.1	Modell Wiederherstellbar	70
A.2	Modell Nicht-Wiederherstellbar	71
A.3	Modell Konstant	72
A.4	Allgemeine Empfehlungen zu Basismodellen	72
B	Ergänzungen zur Berechnung der System-Ausfallrate mit Fehlerbäumen	74
B.1	Verbesserung bei sehr hohen Nichtverfügbarkeiten	74
C	Weitere Verteilungsfunktionen	75
C.1	Normal-Verteilung	75
C.2	Gleichverteilung	75
C.3	Delta-Verteilung	76
D	Importanzen	78
D.1	Allgemeine Hinweise	78
D.2	Partielle Ableitung (PD) und Birnbaum-Importanz (BI)	78
D.2.1	Partielle Ableitung für die System-Nichtverfügbarkeit	79
D.2.2	Partielle Ableitung für die System-Unzuverlässigkeit	80
D.2.3	Partielle Ableitung für die System-Ausfallrate	80
D.2.4	Berechnung für Markov-Modelle	82
D.3	Risk-Reduction (RR)	82
D.4	Risk-Reduction-Worth (RRW)	82
D.5	Fussell-Vesely-Importanz (FV)	83
D.6	Risk-Achievement (RA)	84
D.7	Risk-Achievement-Worth (RAW)	84
D.8	Kritikalitäts-Importanz (CRI)	85
D.9	Importanzen für generische Basis-Ereignisse	85
D.10	Beispielhafte Importanzen für die System-Nichtverfügbarkeit	86
D.11	Beispielhafte Importanzen für die System-Ausfallrate	88
E	Glossar und Abkürzungsverzeichnis	91
F	Literaturverzeichnis	94
G	Änderungsverzeichnis	95

Vorwort und Motivation

Während in der Vergangenheit die funktionale Sicherheit in vielen Branchen kaum eine Rolle spielte, und in den übrigen im Wesentlichen durch detaillierte Design-Regeln gewährleistet war, getrieben durch (negative) Erfahrungen ¹, geht heute der Trend weg von festen Design-Regeln hin zu quantitativen Forderungen und Nachweisen. Dies fördert zweifellos die Innovation und den Wettbewerb, birgt jedoch das Risiko, dass unsichere Systeme auf den Markt kommen.

Die Praxis des Autors als Gutachter für funktionale Sicherheit zeigt immer wieder, dass es selbst erfahrenen Sicherheitsingenieuren schwer fällt, korrekte Berechnungen anzustellen. Oft ist dafür mangelndes Verständnis der unterschiedlichen Größen ursächlich, genauso oft aber auch mangelnde Kenntnisse über die verwendeten Berechnungswerkzeuge und -methoden (insbesondere FTA-Tools), gepaart mit ungerechtfertigt großem Vertrauen in selbige.

Diese Einführung richtet sich in erster Linie an angehende und erfahrende Sicherheitsingenieure, aber auch an Mathematiker oder Informatiker, welche mit der Entwicklung von Berechnungswerkzeugen betraut sind. Es wird gelegentlich Bezug auf Normen genommen, jedoch wird die Kenntnis dieser Normen nicht vorausgesetzt.

Vorbemerkungen

Im Folgenden wird der Begriff System verwendet, da dies in diesem Zusammenhang üblich ist. Tatsächlich aber wäre der Begriff Funktion oft korrekter, da sich ein Ausfall und somit sämtliche Berechnungen generell auf eine Funktion beziehen, welche von einem technischen System ausgeführt werden soll. Der Begriff System ist ohne Nennung der betrachteten (Fehl-) Funktion bedeutungslos, denn ein System wird in der Regel mehrere Funktionen ausführen, welche aufgrund der im Allgemeinen unterschiedlichen beteiligten Komponenten unterschiedliches Ausfallverhalten haben werden, und deren unterschiedliche Fehlfunktionen im Allgemeinen unterschiedliche Folgen haben werden. Der (unpräzise) Begriff System wird im Folgenden auch verwendet, um eine Verwechslung mit den mathematischen Funktionen zu vermeiden und somit das Lesen zu erleichtern.

Im Bereich der Zuverlässigkeitsrechnung wird sehr häufig die sogenannte wissenschaftliche Notation von Zahlenwerten verwendet, etwa $0,0123=1,23\text{e-}2=1,23\text{E-}2$ oder $1000=1\text{E}3=1,0\text{E}3$ ($1,0\cdot 10\text{E}3=10\text{E}3$ ist tatsächlich $1\text{E}4=10000$ – und nicht 1000, wie oft vermutet!).

Als Dezimal-Trennzeichen wird in Graphiken immer der Punkt verwendet, im Text das Komma. Tausender-Trennzeichen werden nicht verwendet.

Das Verständnis der Abschnitte 2 und 4 ist Voraussetzung für alle weiteren Abschnitte, sie sollten daher gelesen werden, bevor die in den Abschnitten 5 und 6 gegebenen Beispielen näher betrachtet werden. Wer sich nur für Fehlerbäume interessiert, kann die Abschnitte zur Markov-Modellierung übergehen.

¹ein ehemaliger Kollege sagte immer: „Sicherheit wurde mit Blut bezahlt“

1 Einführung

Für alle Arten von technischen Systemen, die bei Fehlfunktionen Schaden verursachen können, muss die Sicherheit nachgewiesen werden, bevor sie in Betrieb genommen oder in Verkehr gebracht werden können. Beispiele sind Werkzeugmaschinen, Roboter, Straßen- oder Schienenfahrzeuge, Flugzeuge oder Kraftwerke. Für diese Systeme kann die Sicherheit in Form von Gefährdungsraten, Ausfallraten oder allgemein Eintrittsraten für bestimmte unerwünschte Ereignisse ausgedrückt werden, welche von Sicherheitsingenieuren berechnet werden müssen.

Daneben gibt es eine weitere Klasse von technischen Systemen, für die eine Sicherheit nachgewiesen werden muss: Systeme, die im Gefahrenfall vor Schaden bewahren sollen. Beispiele sind Brandmelder, Notrufsysteme, Notentlastungsventile oder Notpumpen. Diese Systeme können selbst keinen Schaden verursachen, jedoch kann eine Fehlfunktion im Anforderungsfall den Schaden vergrößern oder erst ermöglichen. Die Sicherheit dieser Systeme ist durch ihre Verfügbarkeit bestimmt, welche ein Mindestmaß erfüllen muss.

In dieser Einführung werden alle für die Beschreibung der Sicherheit relevanten Größen benannt und erklärt, und die mathematischen Zusammenhänge erläutert. Dabei werden insbesondere auch instandhaltbare (oder reparierbare) Komponenten und Systeme intensiv betrachtet. Ebenso werden Methoden der Berechnung vorgestellt, insbesondere die Fehlerbaumanalyse und die Markov-Modellierung. Dabei wird auch auf die mathematischen Hintergründe eingegangen, die für die korrekte Modellierung nach Meinung des Autors unerlässlich sind.

2 Zuverlässigkeit und verwandte Größen

2.1 Zuverlässigkeit und Unzuverlässigkeit

Zuverlässigkeit $R(t_1, t_2)$ ist die Wahrscheinlichkeit, dass eine Komponente/ein System/eine Funktion im Zeitintervall t_1 bis t_2 nicht ausfällt, unabhängig davon, ob sie/es zum Zeitpunkt t_1 funktionierte.

Unzuverlässigkeit $F(t_1, t_2)$ ist die Wahrscheinlichkeit, dass eine Komponente/ein System/eine Funktion im Zeitintervall t_1 bis t_2 ausfällt, unabhängig davon, ob sie/es zum Zeitpunkt t_1 funktionierte. Sie ist folglich die Gegenwahrscheinlichkeit zur Zuverlässigkeit:

$$F(t_1, t_2) = 1 - R(t_1, t_2) \quad \text{bzw.} \quad R(t_1, t_2) = 1 - F(t_1, t_2) \quad (1)$$

Bei praktisch allen Fragestellungen wählt man $t_1 = 0$ und erhält somit die ein-parametrischen Funktionen $R(t_1 = 0, t_2 = t) = R(t)$ und $F(t_1 = 0, t_2 = t) = F(t)$. Für die Beziehung zwischen Zuverlässigkeit und Unzuverlässigkeit gilt entsprechend:

$$F(t) = 1 - R(t) \quad \text{bzw.} \quad R(t) = 1 - F(t) \quad (2)$$

Sie wird manchmal auch als Überlebenswahrscheinlichkeit bezeichnet.

Anmerkung: Die Unzuverlässigkeit wird oft auch Ausfallwahrscheinlichkeit genannt. Allerdings wird auch die Nichtverfügbarkeit oft mit Ausfallwahrscheinlichkeit bezeichnet, obwohl sie eine völlig andere Größe ist. Um Missverständnisse zu vermeiden, sollte der Begriff Ausfallwahrscheinlichkeit daher nicht verwendet werden.

Anmerkung: Insbesondere bei Fahrzeugen bezieht man die Zuverlässigkeit und auch die nachfolgenden Größen manchmal auf die zurückgelegte Strecke. Aus $F(t_1, t_2)$ wird dann $F(s_1, s_2)$ etc. In allen angegebenen Formeln muss dann die Zeit durch die Strecke ersetzt werden.

Beispiel 2.1 *Gefragt sei die Wahrscheinlichkeit p , dass eine Komponente mit dem Alter t innerhalb der nächsten Stunde ausfällt. Dabei soll nicht vorausgesetzt werden, dass sie zum Zeitpunkt t noch funktioniert:*

$$p = F(t, t + 1 \text{ h}) = F(t + 1 \text{ h}) - F(t)$$

Beispiel 2.2 *Gefragt sei die Wahrscheinlichkeit p , dass eine Komponente mit dem Alter t , die zum Zeitpunkt t noch funktioniert, innerhalb der nächsten Stunde ausfällt:*

$$p = \frac{F(t, t + 1 \text{ h})}{R(t)} = \frac{F(t + 1 \text{ h}) - F(t)}{R(t)} = \frac{R(t) - R(t + 1 \text{ h})}{R(t)}$$

Beispiel 2.3 *Aus langjähriger Erfahrung sei bekannt, dass die Zuverlässigkeit bzw. Unzuverlässigkeit den in Abbildung 1 dargestellten (kumulierten) Verteilungsfunktionen folge.*

Frage 1: Wie groß ist die Wahrscheinlichkeit, dass die Komponent mehr als 40000 Stunden funktioniert?

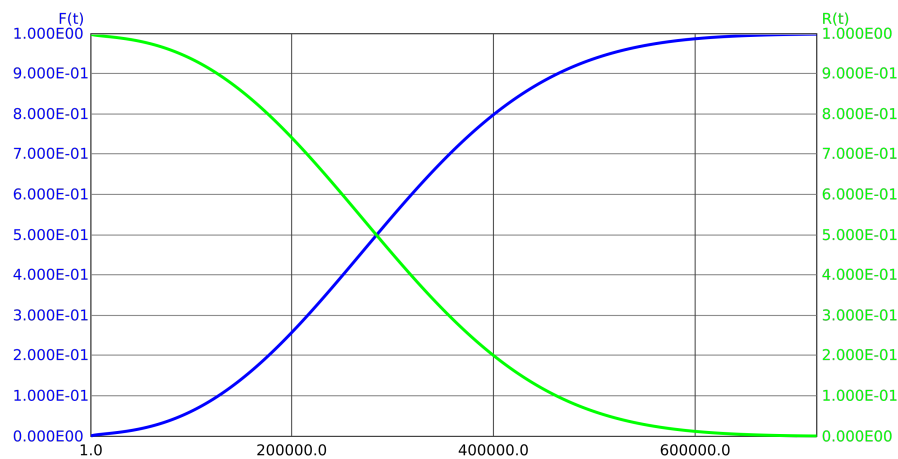


Abbildung 1: Beispielhafte Zuverlässigkeits- und Unzuverlässigkeitsfunktion

Antwort: $p = R(40\,000\text{ h}) \approx 0,2$. Folglich wird die Komponente mit 20% Wahrscheinlichkeit länger als 40000 Stunden funktionieren.

Frage 2: Wie groß ist die Wahrscheinlichkeit, dass die Komponente zwischen 40000 und 50000 Betriebsstunden ausfällt?

Antwort: $p = F(50\,000\text{ h}) - F(40\,000\text{ h}) = R(40\,000\text{ h}) - R(50\,000\text{ h}) \approx 0,15$. Die Komponente wird also mit 15% Wahrscheinlichkeit nach 40000 bis 50000 Stunden ausfallen.

Frage 3: Wie groß ist die Wahrscheinlichkeit, dass die Komponente zwischen 40000 und 50000 Betriebsstunden ausfällt, wenn sie bei 40000 Stunden noch funktioniert hat?

Antwort: $p = \frac{F(50\,000\text{ h}) - F(40\,000\text{ h})}{R(40\,000\text{ h})} \approx \frac{0,15}{0,2} = 0,75$.

2.2 Ausfalldichte und Ausfallrate

Die Änderung der Unzuverlässigkeit pro Zeit ist die Ausfalldichte. Für beliebige Ausfalldichtefunktionen $f(t)$ gilt:

$$F(t) = \int_0^t f(\tau) d\tau \quad \text{bzw.} \quad f(t) = \frac{dF(t)}{dt} = -\frac{dR(t)}{dt} \quad (3)$$

Im Gegensatz zur Unzuverlässigkeit ist die Ausfalldichte keine Wahrscheinlichkeit. Sie kann beliebige positive Werte annehmen und hat eine Dimension (meist 1/Zeit oder 1/Strecke).

Da die Unzuverlässigkeit für $t \rightarrow \infty$ gegen 1 geht, muss für jede Dichtefunktion gelten:

$$\int_0^{\infty} f(\tau) d\tau = 1 \quad (4)$$

Die Ausfallrate $h(t)$ ist für beliebige Ausfallichtefunktionen gegeben als

$$h(t) = \frac{f(t)}{R(t)} = \frac{f(t)}{1 - F(t)} \quad (5)$$

Mit $f(t) = -\frac{dR(t)}{dt}$ ergibt sich:

$$h(t) = \frac{f(t)}{R(t)} = \frac{-dR(t)/dt}{R(t)} = -\frac{\dot{R}}{R} \quad (6)$$

$$R(t) = e^{-\int_0^t h(\tau) d\tau} \quad (7)$$

$$F(t) = 1 - e^{-\int_0^t h(\tau) d\tau} \quad (8)$$

Eine Ausfallverteilung ist durch $f(t)$ oder $F(t)$ oder $R(t)$ oder $h(t)$ vollständig beschrieben.

Abbildung 2 zeigt eine beispielhafte Ausfallverteilung und die sie beschreibenden Größen.

Anmerkung: Während die Symbole $R(t)$, $F(t)$ und auch $f(t)$ weitgehend einheitlich verwendet werden, hat sich für die Ausfallrate noch kein Symbol durchgesetzt. Anstelle von $h(t)$ findet man auch $\Lambda(t)$ oder $\lambda(t)$.

Die Ausfallrate $h(t)$ gibt die Wahrscheinlichkeit eines Ausfalles pro Zeitintervall an, unter der Bedingung, dass die Funktion zu Beginn des Zeitintervalls nicht ausgefallen ist:

$$h(t) = \frac{f(t)}{R(t)} = \frac{dF(t)/dt}{R(t)} = \frac{1}{R(t)} \lim_{\Delta t \rightarrow 0} \frac{F(t + \Delta t) - F(t)}{\Delta t} \quad (9)$$

Hierbei muss das Zeitintervall Δt hinreichend klein sein! Daher darf diese Gleichung keinesfalls herangezogen werden, um eine mittlere Ausfallrate über einen längeren Zeitraum zu berechnen. Die hierfür erforderlichen Formeln sind in Abschnitt 3 erwähnt.

2.3 Badewannenkurve

Jedes nicht extrem einfache Bauteil kann auf unterschiedliche Weise ausfallen. Jeder dieser Ausfallmodi hat eine eigene Ausfallverteilungsfunktion. Fast immer gibt es Ausfallmodi mit abnehmender Ausfallrate, diese sind meist auf Produktionsfehler zurückzuführen. Bei mechanischen oder auch stark belasteten elektronischen Bauteilen gibt es auch Ausfallmodi mit steigender Ausfallrate, dies sind insbesondere die auf Verschleiß oder Alterung zurückzuführenden Ausfälle. Die Gesamt-Ausfallrate ergibt sich durch Addition der Einzel-Ausfallraten aller n Fehlermodi:

$$h(t) = \sum_{i=1}^n h_i(t) \quad (10)$$

Sobald es mindestens je einen Ausfallmodus mit fallender und einen mit steigender Ausfallrate gibt, ähnelt der Graph der Gesamt-Ausfallrate $h(t)$ einer Badewanne, siehe Abbildung 2.

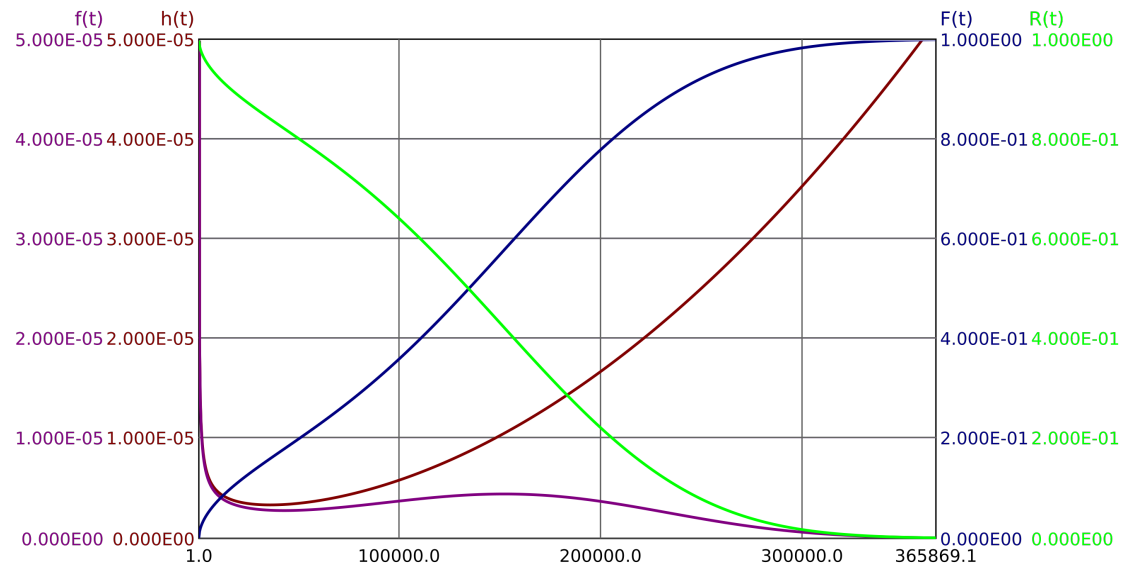


Abbildung 2: Verteilungsfunktion, bei der $h(t)$ einer Badewanne ähnelt

2.4 Mittlere Zeit bis zum Ausfall (MTTF)

Die mittlere Zeit (genauer: Betriebsdauer) bis zum Ausfall (engl.: Mean Time To Failure, kurz MTTF) ist der Erwartungswert der Zeit bis zum Ausfall. Sie berechnet sich für beliebige Ausfallverteilungen wie folgt:

$$\text{MTTF} = \int_0^{\infty} t \cdot f(t) dt \quad (11)$$

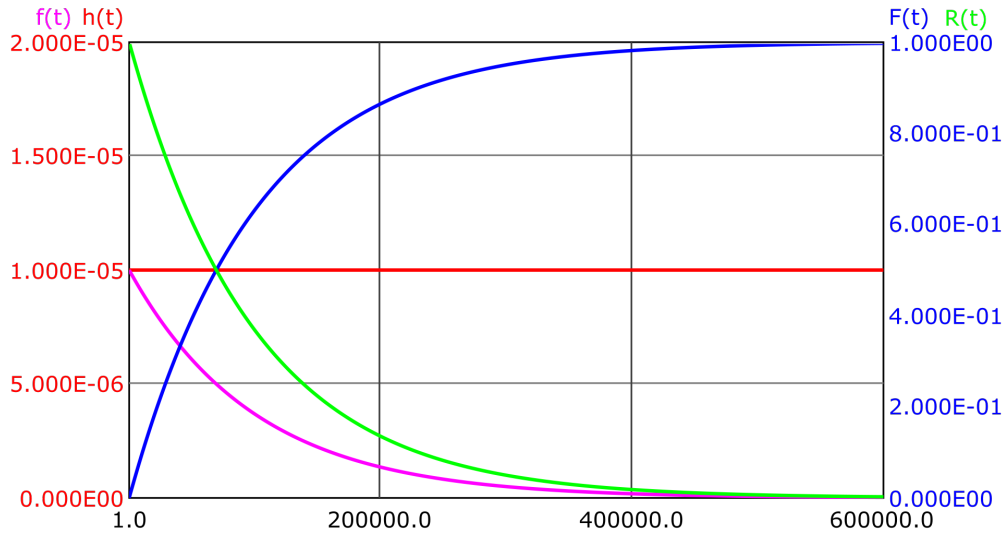
Dieser Wert wird in Abschnitt 3 als natürliche MTTF bezeichnet, da es der Wert ist, der sich experimentell ergibt, wenn die Komponente immer bis zum Ausfall betrieben wird und dann ersetzt wird. In der Praxis ist dies jedoch oft nicht der Fall, so dass insbesondere zur Bestimmung von mittleren Ausfallraten andere Formeln gelten, siehe Abschnitt 3.

2.5 Verteilungsfunktionen

In diesem Abschnitt werden einige Verteilungsfunktionen vorgestellt, die in der Praxis oder für nachfolgende Betrachtungen relevant sind. Weitere Verteilungen sind in Anhang C beschrieben.

2.5.1 Exponentialverteilung

Die Exponentialverteilung zeichnet sich durch eine konstante Ausfallrate $h(t) = \text{const}$ aus. Die Ausfallrate wird also durch einen einzigen Parameter beschrieben, welcher mit λ bezeichnet wird. Die Exponentialverteilung ist die einfachste und zugleich wichtigste Verteilungsfunktion. Sie trifft für gedächtnislose Komponenten zu, also dann, wenn das Alter der Komponente keinen nennenswerten Einfluss auf die Ausfallrate hat (auch als ergodisches Verhalten bezeichnet). Der zeitliche Verlauf von Zuverlässigkeit $R(t)$, Unzuverlässigkeit $F(t)$, Ausfallrate $h(t)$ und Ausfalldichte $f(t)$ ist in Abbildung 3 dargestellt.

Abbildung 3: Exponential-Verteilung mit $\lambda = 1,0\text{E}-5/\text{h}$

Die Exponential-Verteilung ist durch die folgenden Gleichungen beschrieben:

$$f(t) = \lambda \cdot e^{-\lambda \cdot t} \quad (12)$$

$$F(t) = 1 - e^{-\lambda \cdot t} \quad (13)$$

$$R(t) = e^{-\lambda \cdot t} \quad (14)$$

Die mittlere Zeit bis zum Ausfall ist:

$$\text{MTTF} = \int_0^{\infty} t \cdot \lambda \cdot e^{-\lambda \cdot t} dt = -\frac{(\lambda \cdot t + 1) e^{-\lambda \cdot t}}{\lambda} \Big|_0^{\infty} = \frac{1}{\lambda} \quad (15)$$

Beispiel 2.4 Gefragt ist die Wahrscheinlichkeit p , dass eine Komponente mit dem Alter t , die zum Zeitpunkt t noch funktioniert, innerhalb der nächsten Stunde ausfällt. Es sei bekannt, dass die Komponente sich durch eine konstante Ausfallrate beschreiben lässt.

$$\begin{aligned} p &= \frac{F(t, t + 1 \text{ h})}{R(t)} = \frac{F(t + 1 \text{ h}) - F(t)}{R(t)} = \frac{e^{-\lambda \cdot t} - e^{-\lambda \cdot (t+1 \text{ h})}}{e^{-\lambda \cdot t}} \\ &= \frac{e^{-\lambda \cdot t} - e^{-\lambda \cdot t} \cdot e^{-\lambda \cdot 1 \text{ h}}}{e^{-\lambda \cdot t}} = 1 - e^{-\lambda \cdot 1 \text{ h}} \end{aligned}$$

Wie zu erwarten war, taucht das Alter der Komponente t im Ergebnis nicht auf.

Viele Elemente lassen sich durch eine konstante Ausfallrate hinreichend genau beschreiben. Insbesondere kann bei Elementen eines Systems, deren Lebensdauer kürzer ist als die System-Einsatzdauer, und die nicht zu bestimmten Zeitpunkten präventiv getauscht werden, im Rahmen von System-Berechnungen ohnehin nur eine konstante mittlere Ausfallrate $h(t) = \bar{h} = \lambda$ angegeben werden, da die tatsächliche Ausfallratenfunktion selbst zufällig ist.

2.5.2 Weibull-Verteilung

Die Weibull-Verteilung ist eine Verallgemeinerung der Exponentialverteilung. Durch einen zusätzlichen Parameter $k > 0$ im Exponenten können fallende oder steigende Ausfallraten modelliert werden. Für $0 < k < 1$ ergibt sich eine fallende, für $k > 1$ eine steigende Ausfallrate. Für $k = 1$ ergibt sich die Exponentialverteilung.

$$h(t) = \lambda \cdot k \cdot (\lambda \cdot t)^{k-1} \quad (16)$$

$$f(t) = \lambda \cdot k \cdot (\lambda \cdot t)^{k-1} e^{-(\lambda \cdot t)^k} \quad (17)$$

$$F(t) = 1 - e^{-(\lambda \cdot t)^k} \quad (18)$$

$$\text{MTTF} = \frac{1}{\lambda} \cdot \Gamma\left(1 + \frac{1}{k}\right) \quad (19)$$

Fehlermodi, die ausschließlich auf Abnutzung zurückzuführen sind, können in der Regel für eine bestimmte Zeit t_0 völlig ausgeschlossen werden. Diese Fehlermodi kann man meist gut mit einer Weibull-Verteilung mit $k > 1$ (steigende Ausfallrate) modellieren, die zusätzlich noch um $t_0 > 0$ nach rechts verschoben ist:

$$h(t) = \lambda \cdot k \cdot (\lambda \cdot (t - t_0))^{k-1} \quad \text{für } t > t_0 \quad (20)$$

Hinweis: Insbesondere im englischsprachigen Raum wird die Weibull-Verteilung oft mit $\mu = 1/\lambda$ parametrisiert.

Die Abbildungen 4 und 5 und zeigen eine Weibull-Verteilung mit fallender Ausfallrate ($k=0,5$) und eine mit steigender Ausfallrate ($k=3$).

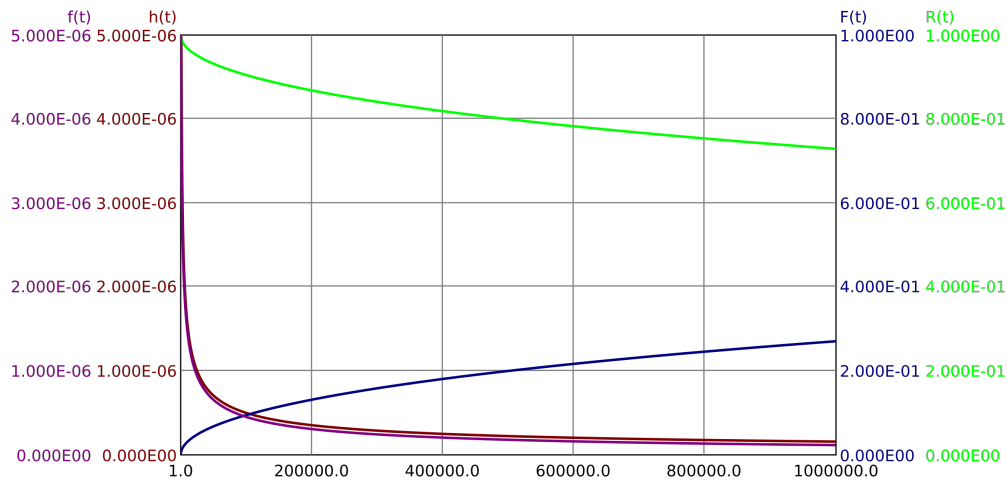


Abbildung 4: Weibull-Verteilung mit $\lambda = 1,0\text{E}-7/h$ und $k = 0,5$

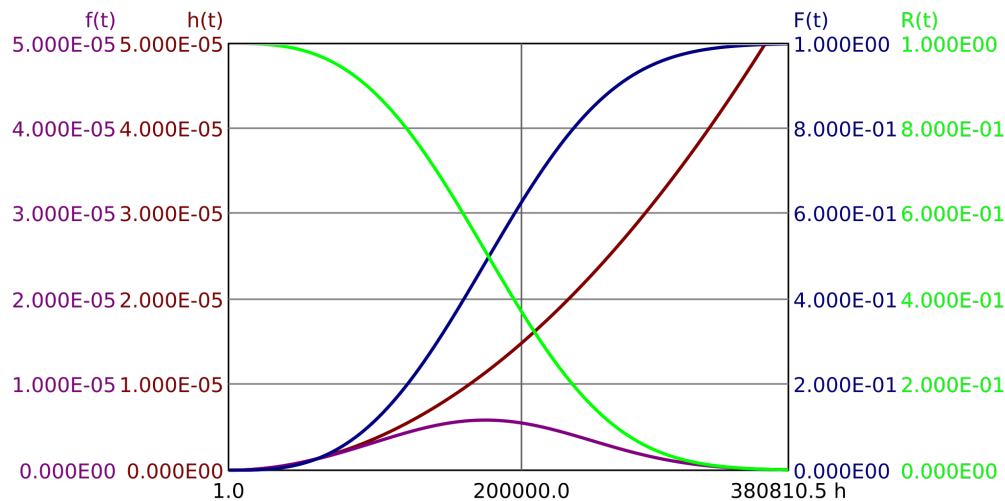


Abbildung 5: Weibull-Verteilung mit $\lambda = 5,0\text{E}-6/h$ und $k = 3,0$

2.5.3 Sterblichkeit

Auch die Sterblichkeit des Menschen ist eine Verteilungsfunktion, wenngleich diese nicht direkt einer mathematische Funktion folgt. In Abbildung 6 ist die Querschnittsbetrachtung der Sterblichkeit der westdeutschen Bevölkerung in den Jahren 1960-1962 dargestellt.

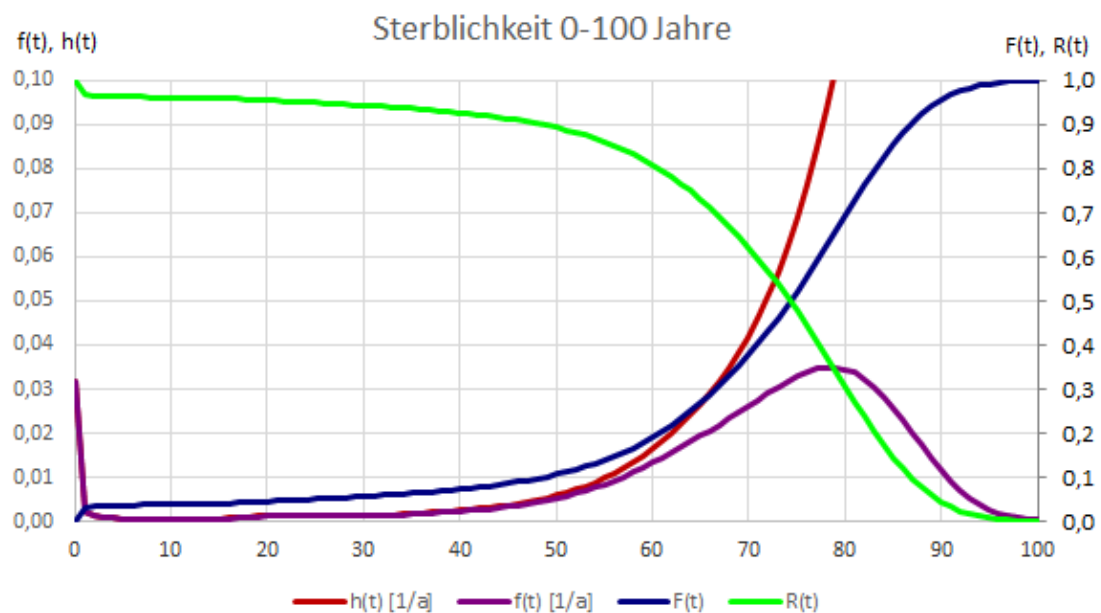


Abbildung 6: Sterblichkeit 1962 (Querschnittsbetrachtung)

Die Querschnittsbetrachtung basiert auf der Statistik der Todesfälle in dem genannten Zeitraum. Sie macht damit insbesondere eine Aussage über das tatsächliche mittlere Sterbealter der Menschen, die in diesem Zeitraum gestorben sind. Die Sterberate $h(t)$ gibt hier also die Wahrscheinlichkeit an, dass eine Person, die das Alter t erreicht hatte,

innerhalb der Zeitspanne Δt stirbt, dividiert durch diese Zeitspanne Δt .

Im Gegensatz zur Querschnittsbetrachtung macht die sogenannte Längsschnittbetrachtung eine Aussage über die Sterbeverteilung der in dem jeweiligen Zeitraum geborenen Menschen. Statistische Längsschnittbetrachtungen können also nur für Geburtsjahrgänge gemacht werden, von denen kein Mensch mehr am Leben ist. Für spätere Jahrgänge stellen sie ganz oder teilweise Prognosen dar. Wäre die Sterblichkeitsverteilung unabhängig vom Jahrgang, wären Querschnitts- und Längsschnittverteilung identisch. Für die Längsschnittbetrachtung sind die Größen unmittelbar anschaulich:

- Die Zuverlässigkeit (Überlebenswahrscheinlichkeit) $R(t)$ ist die Wahrscheinlichkeit, das Alter t zu erreichen.
- Die Unzuverlässigkeit (Ausfallwahrscheinlichkeit) $F(t)$ ist die Wahrscheinlichkeit, vor Erreichen des Alters t zu sterben.
- Die Ausfalldichte $f(t)$ ist die Wahrscheinlichkeit, im Alter zwischen t und $t + \Delta t$ zu sterben, dividiert durch den Zeitraum Δt , mit $\Delta t \rightarrow 0$.
- Die Ausfallrate $h(t)$ ist die Wahrscheinlichkeit, im Alter zwischen t und $t + \Delta t$ zu sterben, unter der Bedingung, das Alter t erreicht zu haben, dividiert durch den Zeitraum Δt , mit $\Delta t \rightarrow 0$.
- Die MTTF ist die Lebenserwartung eines Neugeborenen.

3 Mittlere Ausfallrate und mittlere Zeit bis zum Ausfall

Bei vielen Komponenten ist die Ausfallrate stark zeitabhängig. Auch für solche Komponenten wird oft eine mittlere Ausfallrate benötigt, unter anderem aus folgenden Gründen:

- Sollen Systemgrößen ($\overline{h_{\text{sys}}}$, $\overline{Q_{\text{sys}}}$) stationär berechnet werden, können prinzipbedingt nur mittlere Ausfallraten verwendet werden.
- Wenn eine Komponente während der Einsatzzeit des Systems wahrscheinlich mehrfach ausfallen (und ersetzt oder repariert werden) wird, ist ab dem ersten Ausfall auch die Ausfallratenfunktion selbst eine Zufallsgröße, die mit zunehmender Anzahl von Ausfällen immer unschärfer wird. Selbst bei transienten, also zeitlich kontinuierlich aufgelösten Betrachtungen, kann also keine Ausfallratenfunktion angegeben werden.
- Falls die Komponente regelmäßig präventiv getauscht werden soll, damit sie möglichst nicht ausfallen wird, stellt sich die Frage, in welchem Intervall die Komponente getauscht werden sollte, um die effektive (Rest-)Ausfallrate und damit die Wahrscheinlichkeit eines Ausfalls möglichst gering zu halten.

In diesem Abschnitt sollen daher die folgenden Aufgaben behandelt werden:

1. Wie groß sind MTTF und mittlere Ausfallrate für den Fall, dass die Komponente bis zum Ausfall betrieben wird und dann getauscht wird? (Beispiel: Glühbirne)
2. Was groß sind MTTF und effektive mittlere Ausfallrate für den Fall, dass die Komponente regelmäßig präventiv getauscht wird? (Beispiel: Steuerriemen eines Verbrennungsmotors)
3. Wenn es gefährliche und ungefährliche Ausfallmodi einer Komponente gibt: Wie können die gefährliche MTTF und die gefährliche mittlere Ausfallrate für die beiden vorgenannten Fälle berechnet werden?

Häufig wird behauptet, man solle bei derlei Fragestellungen die Ausfallrate im flachen Bereich der „Badewannenkurve“ verwenden. Dies ist jedoch nur dann richtig, wenn die Komponente auch wirklich nur in diesem Bereich betrieben wird, also Frühausfälle (insbesondere Produktionsfehler) ebenso absolut ausgeschlossen werden können wie Ausfälle durch Alterung und Verschleiß. Diese Bedingungen sind sehr häufig nicht erfüllt – und zudem weist die Badewannenkurve oft auch gar keinen richtig flachen Bereich auf, sondern sogenannte Frühausfälle überlagern sich mit Spätausfällen.

Anmerkung: Mancher Leser wird sich fragen, warum die Abkürzung MTTF auch für die Zeit zwischen zwei Ausfällen verwendet wird, und nicht etwa MTBF (Mean Time Between Failures). Die Antwort ist einfach, dass fast immer, wenn von MTBF die Rede ist, tatsächlich die MTTF gemeint ist. In Anwendungen oder Berechnungen, in denen Fehlerdetektionszeiten oder Reparaturzeiten nicht unerheblich sind, müssen diese ohnehin ausdrücklich erwähnt werden. Es gibt somit weder in der Sicherheits- noch in der Zuverlässigkeitstheorie einen stichhaltigen Grund, eine Größe MTBF einzuführen oder

zu verwenden.²

3.1 MTTF bei langer Nutzung ohne vorbeugenden Tausch

Zunächst sollen die MTTF und die mittlere Ausfallrate für eine Komponente berechnet werden, welche eine deutlich geringere Lebenserwartung hat als die nominale Einsatzzeit des Gesamtsystems. In dem Fall ist davon auszugehen, dass die Komponente mehrfach ausfallen wird und daher mehrfach getauscht werden muss.

Für beliebige Ausfallverteilungsfunktionen gilt:

$$\text{MTTF} = \int_0^{\infty} t \cdot f(t) dt \quad (21)$$

Liegen ausreichende Test- oder Felddaten vor, so kann dieses Integral leicht mit einer Tabellenkalkulation berechnet werden. Falls anstelle von $f(t)$ die Daten $F(t)$ oder $R(t)$ vorliegen, kann $f(t)$ leicht durch numerisches Differenzieren gewonnen werden.

Ganz allgemein gilt für die Ausfalldichtenfunktion

$$f(t) = h(t) \cdot R(t) = h(t) \cdot e^{-\int_0^t h(\tau) d\tau} \quad (22)$$

und damit für die MTTF

$$\text{MTTF} = \int_0^{\infty} t \cdot h(t) \cdot e^{-\int_0^t h(\tau) d\tau} dt \quad (23)$$

Fast alle Komponenten haben mehrere Fehlermodi, die unterschiedlichen Ausfallverteilungsfunktionen gehorchen. Angenommen die Ausfallverteilungsfunktionen der einzelnen Fehlermodi seien bekannt, wie kann dann die MTTF berechnet werden? Dazu muss vorausgesetzt werden, dass die Fehlermodi voneinander unabhängig sind, also sich nicht gegenseitig beeinflussen. Damit diese Voraussetzung erfüllt ist, ist insbesondere notwendig, dass die Komponente im Fall eines jeden Ausfalls ausgetauscht wird. Dann gilt für die Gesamt-Ausfallratenfunktion $h(t)$ der Komponente, dass sie durch die Summe der Ausfallratenfunktionen der einzelnen Fehlermodi gegeben ist:

$$h(t) = \sum_{i=1}^n h_i(t) \quad (24)$$

Falls sowohl alle $h_i(t)$ als auch die jeweils zugehörigen Zuverlässigkeitsfunktionen $R_i(t)$ durch mathematische Formeln gegeben sind, ist es hilfreich, das doppelte Integral zu

²Tatsächlich bin ich nicht sicher, ob ich jemals ein Dokument gesehen habe, in dem der Begriff MTBF korrekt verwendet worden wäre – abgesehen von Lehrbüchern.

vereinfachen:

$$\begin{aligned} \text{MTTF} &= \int_0^{\infty} t \cdot h(t) \cdot e^{-\int_0^t \sum_{i=1}^n h_i(\tau) d\tau} dt = \int_0^{\infty} t \cdot h(t) \cdot e^{-\sum_{i=1}^n \int_0^t h_i(\tau) d\tau} dt \\ &= \int_0^{\infty} t \cdot h(t) \cdot \prod_{i=1}^n R_i(t) dt \end{aligned} \quad (25)$$

Die so ermittelte MTTF wird im folgenden vollständige MTTF oder natürliche MTTF genannt, denn es ist die mittlere Zeit bis zum Ausfall, die sich ergibt, wenn die Komponente bis zum Ausfall eingesetzt und dann getauscht wird (wie dies in langlebigen Systemen wie Maschinen, Flugzeugen oder Eisenbahnfahrzeugen für viele Komponenten der Fall ist).

Die mittlere Ausfallrate der Komponente für den Fall, dass die Komponente wahrscheinlich (mehrfach) ausfallen wird und dann jeweils ersetzt wird, ist der Kehrwert der vollständigen MTTF:

$$\lambda = \frac{1}{\text{MTTF}} \quad (26)$$

3.2 Vorbeugender Tausch und unvollständige MTTF

Im Falle von vorbeugendem Austausch nach einem Zeitintervall T wird die unvollständige MTTF(T) für die Zeitspanne 0 bis T benötigt, ebenso wenn die natürliche MTTF der Komponente (in der Anwendung) wesentlich größer ist als die Lebenszeit des Systems, in dem sie eingesetzt wird. Aus Überlegungen, welche hier nicht wiedergegeben werden können, folgt:

$$\text{MTTF}(T) = \frac{\int_0^T t \cdot f(t) dt + T \cdot R(T)}{F(T)} \quad (27)$$

Mit $R(T) = 1 - F(T)$ und den bereits bekannten Formeln für die Beziehung zwischen Zuverlässigkeit und Ausfallrate erhält man eine Formel zur Berechnung von $\text{MTTF}(T)$ bei gegebener oder experimentell ermittelter Ausfallratenfunktion $h(t)$:

$$\begin{aligned} \text{MTTF}(T) &= \frac{\int_0^T t \cdot f(t) dt + T \cdot (1 - F(T))}{F(T)} = \frac{\int_0^T t \cdot f(t) dt + T}{F(T)} - T \\ &= \frac{\int_0^T t \cdot h(t) \cdot e^{-\int_0^t h(\tau) d\tau} dt + T}{1 - e^{-\int_0^T h(t) dt}} - T \end{aligned} \quad (28)$$

Für $T \rightarrow \infty$ geht die unvollständige $\text{MTTF}(T)$ in die vollständige MTTF über.

Der Kehrwert der unvollständigen MTTF zur Zeit T ist die effektive Ausfallrate λ_{eff} . Sie gibt ganz praktisch an, wie oft die Komponente trotz regelmäßiger präventiver Tauschs

bei einer (sehr langen) Einsatzzeit des Gesamtsystems $T_{\text{Life,sys}}$ ausfallen würde:

$$\lambda_{\text{eff}}(T) = \frac{1}{\text{MTTF}(T)} = \frac{N(T_{\text{Life,sys}})}{T_{\text{Life,sys}}} \quad (29)$$

Dabei meint $N(T_{\text{Life,sys}})$ die zählbaren Ausfälle der Komponente im System. Man kann $\text{MTTF}(T)$ daher auch als die effektive MTTF für ein bestimmtes Tauschintervall T bezeichnen.

Beispiel 3.1 Die Ausfallrate des Steuerriemens eines Motors eines PKW lasse sich durch zwei überlagerte Weibull-Verteilungen beschreiben:

$$\lambda_1 = 1\text{E}-9/\text{h}; k_1 = 0,3 \Rightarrow h_1(t) = 1\text{E}-9/\text{h} \cdot 0,3 \cdot (1\text{E}-9/\text{h} \cdot t)^{0,3-1}$$

$$\lambda_2 = 2\text{E}-4/\text{h}; k_2 = 4,0 \Rightarrow h_2(t) = 2\text{E}-4/\text{h} \cdot 4,0 \cdot (2\text{E}-4/\text{h} \cdot t)^{4,0-1}$$

Dabei beschreibt $h_1(t)$ sogenannte Frühausfälle, wie sie etwa durch fehlerhafte Komponenten oder fehlerhafte Montage entstehen, und $h_2(t)$ die verschleißbedingten Ausfälle des Riemens.

Für die Ausfallrate des Steuerriemens gilt gemäß Formel (24):

$$h_{\text{ges}}(t) = h_1(t) + h_2(t)$$

Weiter sei angenommen, dass ein PKW mindestens 5000 Betriebsstunden wirtschaftlich betrieben werden können soll.

Mit dieser Ausfallratenfunktion errechnet sich die Unzuverlässigkeit zum Zeitpunkt $T = 5000\text{h}$ zu $F(5000\text{h}) \approx 0,64$. Es wäre also bei mindestens jedem zweiten Fahrzeug zu erwarten, dass der Steuerriemens reißt, bevor 5000 Betriebsstunden erreicht sind. Da der Riss des Steuerriemens eines Motors meist einen Totalschaden des Motors und somit oft einen wirtschaftlichen Totalschaden des Fahrzeugs mit sich bringt, stellt sich die Frage, ob nicht ein präventiver Tausch nach einer bestimmten Zeit (oder Fahrstrecke) sinnvoll ist.

In Abbildung 7 sind neben Ausfalldichte $f(t)$, Ausfallrate $h(t)$, Zuverlässigkeit $R(t)$ und Unzuverlässigkeit $F(t)$ auch die effektive $\text{MTTF}(T)$ und (gestrichelt) die effektive Ausfallrate $\lambda_{\text{eff}}(T)$ in Abhängigkeit der Zeit bis zum präventiven Tausch T dargestellt.

Man erkennt, dass bei einer Einsatzzeit T von etwa 1200 Stunden die $\text{MTTF}(T)$ mit etwa 61000h ein Maximum annimmt. Wechselt man also nach etwa 1200 Stunden den Steuerriemen, beträgt die effektive Ausfallrate $\lambda_{\text{eff}} \approx 1,6\text{E}-5/\text{h}$. Wird der Riemen häufiger gewechselt, sinkt die effektive (unvollständige) $\text{MTTF}(T)$, da Frühausfälle noch einen relativ starken Einfluss haben. Wird der Riemen länger betrieben, sinkt die effektive $\text{MTTF}(T)$ ebenfalls, da sich Ausfälle aufgrund von Verschleiß stärker bemerkbar machen. Der präventive Tausch sollte also nach etwa 1200 Stunden (oder einer entsprechenden Strecke) vorgeschrieben werden.

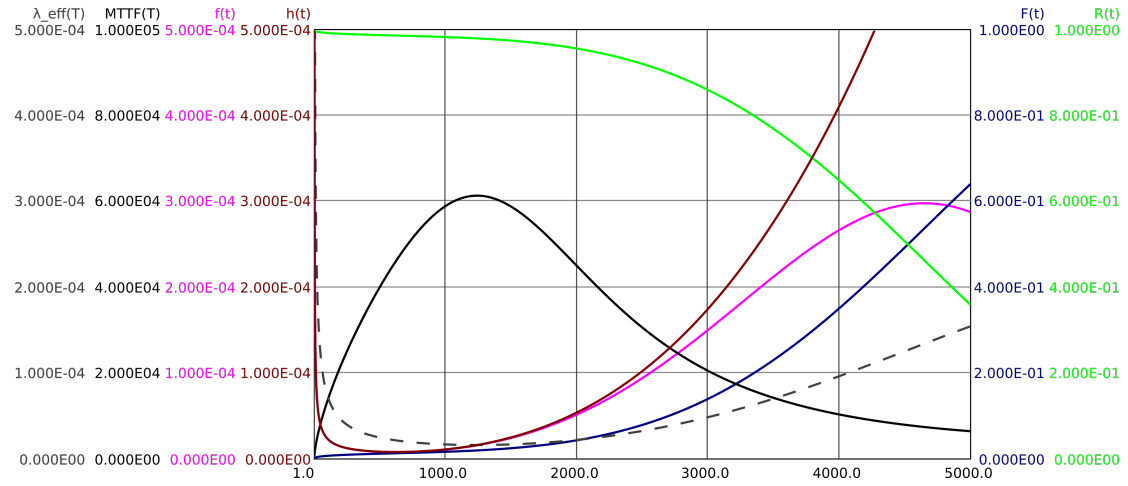


Abbildung 7: Zuverlässigkeit eines Steuerriemens

Die effektive MTTF(T) von 61000 h bedeutet bei einem Tauschintervall von 1200 h praktisch, dass etwa jeder fünfzigste ($61000\text{ h}/1200\text{ h} \approx 50$) Riemen im Betrieb reißen wird.
3

3.3 Gefährliche und ungefährliche Ausfallmodi, gefährliche MTTF

Viele insbesondere komplexere Komponenten können verschiedenartig ausfallen. Dabei sind oft einige Ausfallarten sicherheitskritisch, andere nicht. Für Sicherheitsbetrachtungen ist es daher oft sinnvoll oder erforderlich, zwischen gefährlichen (dangerous, d) und ungefährlichen (safe, s) Ausfallarten zu unterscheiden. Die Gesamt-Ausfallrate zu jeder Zeit t ist die Summe zweier Teil-Ausfallraten für gefährliche und ungefährliche Ausfälle:

$$h(t) = h_d(t) + h_s(t) \quad (30)$$

Die Dichte kann man über

$$\begin{aligned} f(t) &= h(t) \cdot R(t) = (h_d(t) + h_s(t)) \cdot R(t) \\ &= h_d(t) \cdot R(t) + h_s(t) \cdot R(t) \\ &= \varphi_d(t) + \varphi_s(t) \end{aligned} \quad (31)$$

in zwei Teil-Ausfalldichten φ_d und φ_s zerlegen (φ_d und φ_s sind selbst keine Dichten, da ihre einzelnen Integrale kleiner 1 sind).

Entsprechend kann man die Verteilungsfunktion $F(t)$ in zwei Teilfunktionen zerlegen:

$$F(t) = \Phi_d(t) + \Phi_s(t) = \int_0^t \varphi_d(\tau) d\tau + \int_0^t \varphi_s(\tau) d\tau \quad (32)$$

³Die Ausfallratenfunktionen sind natürlich erfunden und statistische Unsicherheiten wie Umweltbedingungen, Straßenarten, Fahrstil etc. außer Acht gelassen.

Aus ähnlichen Überlegungen wie für die unvollständige MTTF(T) erhält man für die effektive gefährliche MTTF_d(T):

$$\text{MTTF}_d(T) = \frac{\int_0^T t \cdot f(t) dt + T \cdot R(T)}{\Phi_d(T)} \quad (33)$$

Mit den bereits bekannten Formeln für die Beziehung zwischen Zuverlässigkeit und Ausfallrate erhält man eine Formel zur Berechnung von MTTF_d(T) bei gegebenen oder experimentell ermittelten Ausfallratenfunktionen $h(t)$ und $h_d(t)$:

$$\begin{aligned} \text{MTTF}_d(T) &= \frac{\int_0^T t \cdot f(t) dt + T \cdot R(T)}{\int_0^T \varphi_d(t) dt} = \frac{\int_0^T t \cdot h(t) \cdot R(t) dt + T \cdot R(T)}{\int_0^T h_d(t) \cdot R(t) dt} \\ &= \frac{\int_0^T t \cdot h(t) \cdot e^{-\int_0^t h(\tau) d\tau} dt + T \cdot e^{-\int_0^T h(t) dt}}{\int_0^T h_d(t) \cdot e^{-\int_0^t h(\tau) d\tau} dt} \end{aligned} \quad (34)$$

Ein einfacher Formelvergleich liefert ferner die Beziehung:

$$\text{MTTF}_d(T) = \text{MTTF}(T) \frac{F(T)}{\Phi_d(T)} \quad (35)$$

Auch hier kann man die effektive mittlere gefährliche Ausfallrate λ_d als Kehrwert berechnen:

$$\lambda_d(T) = \frac{1}{\text{MTTF}_d(T)} \quad (36)$$

In der folgenden Abbildung 8 sind die für eine Komponente mit je drei gefährlichen und drei ungefährlichen Ausfallarten relevanten Zuverlässigkeitsgrößen dargestellt (durchgezogene Linien für gefährliche Ausfälle, gestrichelte Linien für gesamte Ausfälle):

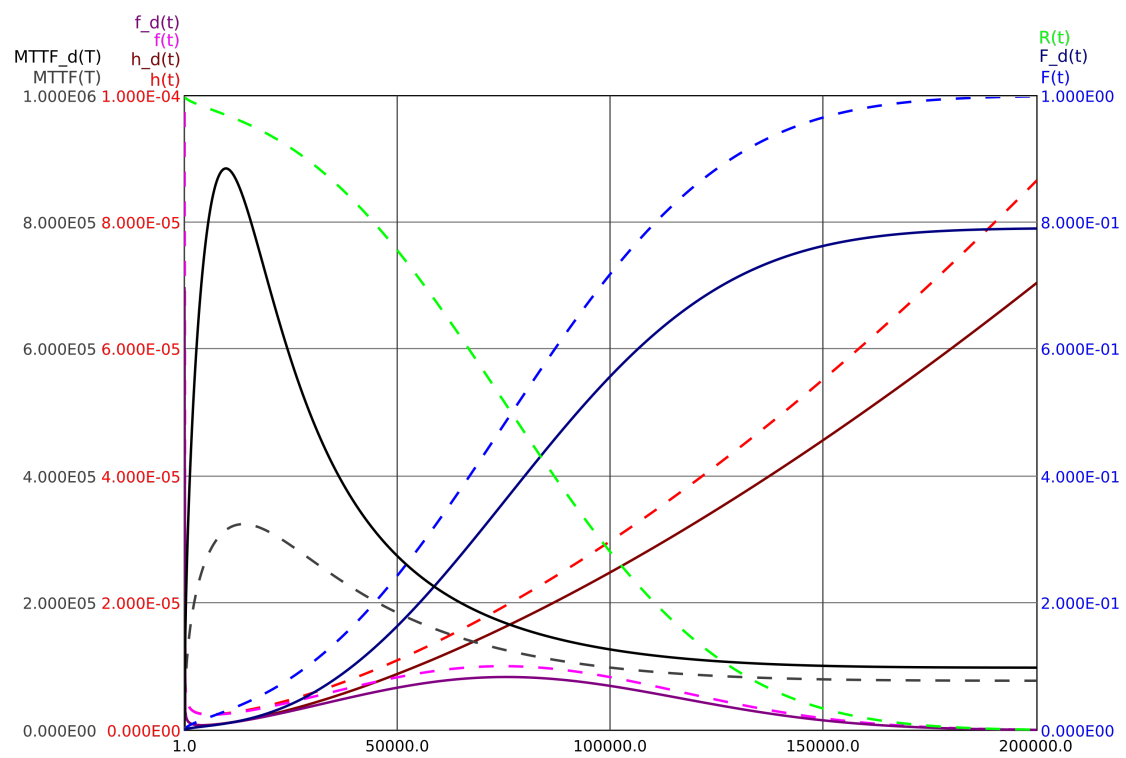


Abbildung 8: Badewannenkurve und Größen bei gefährlichen und ungefährlichen Ausfällen

4 Wiederherstellung und Verfügbarkeit

Bei Komponenten und Systemen, die im Fall eines Ausfalls repariert oder ersetzt werden, sind weitere Betrachtungen und Größen erforderlich.

4.1 Reparierbarkeit

Wenn davon ausgegangen werden muss, dass es während der spezifizierten Einsatzdauer des Gesamtsystems zu mehreren Ausfällen der Funktion kommen kann, ist ein System (bzw. eine Funktion) reparierbar. Dabei spielt es keine Rolle, ob das System im wörtlichen Sinne repariert wird, oder einzelne Komponenten oder sogar das ganze System durch ein gleiches oder anderes ersetzt wird. Reparierbarkeit ist also keine Definitionssache (wie die Festlegung einer kleinsten tauschbaren Einheit oder des Totalschadens), sondern ergibt sich durch die Zuverlässigkeiten der Komponenten und die geplante Einsatzdauer zwangsläufig. Die viel einfachere Modellierung einer Funktion als nicht-reparierbar darf nur dann gewählt werden, wenn die Unzuverlässigkeit sämtlicher Komponenten über die angedachte, spezifizierte Einsatzdauer klein ist (etwa kleiner 0,1). Dies ist bei langlebigen Systemen wie Flugzeugen, Eisenbahnen, Maschinen oder Industrieanlagen praktisch nie erfüllt, bei kurzlebigen wie PKW nur bei manchen Funktionen.

4.2 Diagnose, Test, Wiederherstellung

Als Diagnose bezeichnet man gemäß [IEC 61508] und anderen Normen die Maßnahmen, die einen Fehler innerhalb der Prozessfehlertoleranzzeit (PFTZ, auch als Prozess-Sicherheitszeit bezeichnet) entdecken. Das ist die Zeit, die ein physikalischer Prozess (z. B. ein Motor oder ein Ventil in einer Maschine) falsch angesteuert werden darf, ohne dass hierdurch ein unkontrollierbarer oder gefährlicher Systemzustand resultiert.

Als Tests werden die Maßnahmen bezeichnet, die Fehler erst nach einer mehr oder weniger genau definierten Fehleroffenbarungszeit offenbaren, beispielsweise im Rahmen eines Neustarts (Power-on-Self-Test), einer regelmäßig durchzuführenden Testroutine (Prüflauf) oder bei einer Wartungsmaßnahme (Werkstattinspektion). Die mittlere Zeit, in der ein vorhandener Fehler detektiert wird, wird meist mit MTTD (engl. Mean Time To Detect) abgekürzt. Wird ein Test in regelmäßigen Abständen T_{test} durchgeführt, so beträgt die $\text{MTTD} = T_{\text{test}}/2$.

Im Fall eines erkannten Defekts wird entweder die defekte Komponente repariert oder ersetzt, oder ein ganzes Modul ersetzt oder gar das ganze System (Maschine, Fahrzeug,...) außer Betrieb genommen und durch ein neues ersetzt. In jedem Fall wird die Funktion wieder hergestellt, denn diese wird ja in der Regel weiterhin benötigt. Mit welcher Maßnahme die Funktion konkret wieder hergestellt wird, spielt für die weiteren Betrachtungen keine Rolle.

4.3 Verfügbarkeit und Nichtverfügbarkeit

Verfügbarkeit $A(t)$ ist die Wahrscheinlichkeit, dass eine Komponente/ein System/eine

Funktion zum Zeitpunkt t funktioniert.

Nichtverfügbarkeit $Q(t)$ ist die Wahrscheinlichkeit, dass eine Komponente/ein System/eine Funktion zum Zeitpunkt t nicht funktioniert. Sie ist folglich die Gegenwahrscheinlichkeit zur Verfügbarkeit:

$$A(t) = 1 - Q(t) \quad \text{bzw.} \quad Q(t) = 1 - A(t) \quad (37)$$

Die Verfügbarkeit ist die entscheidende Größe bei Systemen/Funktionen, die nur gelegentlich benötigt werden, insbesondere also Funktionen, die nur in Ausnahme- oder Notfällen benötigt werden (z.B. Alarmer oder Löscheinrichtungen). Sie ist auch eine wesentliche Größe bei Systemen/Funktionen, die Redundanzen aufweisen, sogenannte mehrkanalige Systeme, hierzu später mehr.

Für eine Komponente oder ein System, welches nie getestet und somit auch nie repariert oder ersetzt wird, gilt:

$$Q(t) = F(0, t) \quad (38)$$

Die Verfügbarkeit kann durch kontinuierliche Diagnose oder regelmäßige Tests und falls nötig Wiederherstellung maßgeblich erhöht werden. Dies ist der Grund, warum praktisch alle Notfallsysteme regelmäßig getestet werden.

Wenn eine Komponente getestet und im Fall eines Defekts repariert oder ausgetauscht wird, gilt $Q(t) = F(t)$ nur bis zum ersten Test. Im Fall von regelmäßigen Tests im Abstand von T_{test} gilt bis zur ersten Wiederherstellung theoretisch:

$$Q(t) = \frac{F(t - t \bmod T_{\text{test}}, t)}{R(0, t - t \bmod T_{\text{test}})} = \frac{F(t) - F(t - t \bmod T_{\text{test}})}{R(t - t \bmod T_{\text{test}})}$$

Da aber nicht bekannt ist, wann der erste Defekt auftritt und somit die erste Reparatur oder der erste Austausch erforderlich ist, ist diese Formel praktisch bedeutungslos. Aus demselben Grund ist auch eine veränderliche Ausfallrate praktisch bedeutungslos, denn man kann ja nie sagen, wie lange zum Zeitpunkt t eine Komponente bereits im Einsatz gewesen ist, da man nicht weiß, wann sie eingebaut wurde – es könnte sich ja bereits um einen Ersatz handeln. Folglich muss immer eine mittlere Ausfallrate $\bar{h} = \lambda = 1/\text{MTTF}$ gemäß Abschnitt 3 bestimmt und verwendet werden.

Die Nichtverfügbarkeit ist für Komponenten und Systeme, welche regelmäßig (und vollständig) getestet werden, daher eine periodische Aneinanderreihung von Anfangsstücken der Exponentialverteilung:

$$Q(t) = F(0, t \bmod T_{\text{test}}) = 1 - e^{-\lambda(t \bmod T_{\text{test}})} \quad (39)$$

Ist die Zeit t klein im Verhältnis zu $MTTF = 1/\lambda$, so gilt geringfügig konservativ

$$Q(t) \lesssim \lambda \cdot (t \bmod T_{\text{test}}) \quad (40)$$

Für eine mittlere Ausfallrate $h(t) = \lambda = 1\text{E}-5/\text{h}$ und ein Testintervall $T_{\text{test}} = 1000\text{ h}$ sind Nichtverfügbarkeit $Q(t)$ und Unzuverlässigkeit $F(t)$ in Abbildung 9 dargestellt.

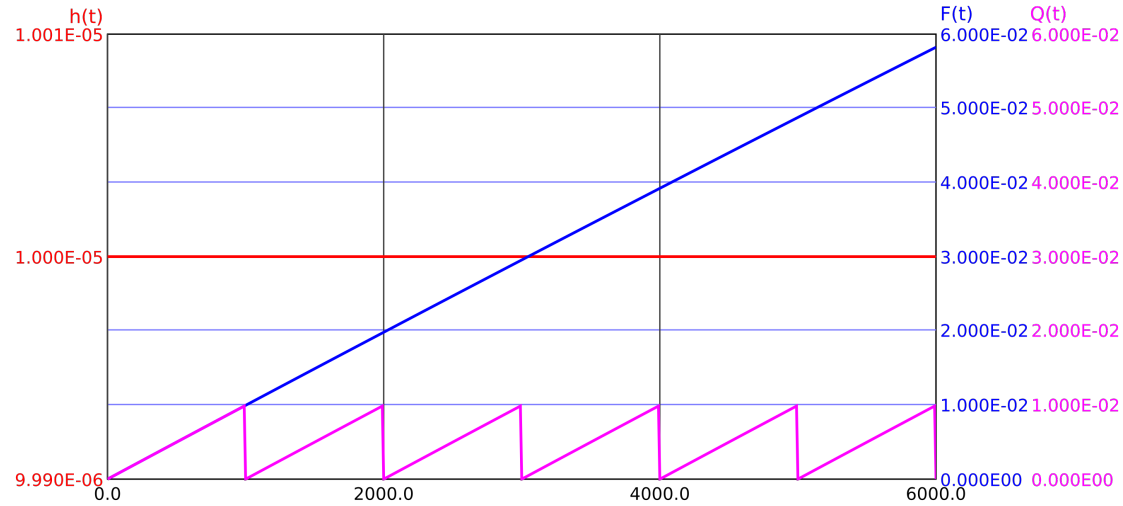


Abbildung 9: Unzuverlässigkeit und Nichtverfügbarkeit mit Tests

Die Nichtverfügbarkeit sinkt bei jedem Test auf null, um dann erneut anzusteigen. Dabei wird vorausgesetzt, dass der regelmäßige Test vollständig ist, also alle relevanten Fehler der Komponente offenbart. Ist dies nicht der Fall, muss dieser verbleibende Anteil als weitere Nichtverfügbarkeit mit einer entsprechend kleineren Ausfallrate, aber längeren Testzeit (dann meist die System-Lebenszeit) hinzugefügt werden. Überschlagsmäßig kann man diese zweite Nichtverfügbarkeit einfach addieren, in speziellen Software-Werkzeugen (FTA- oder Markov-Tools) ist auch eine exakte Behandlung möglich.

Beispiel 4.1 Ein Rauchmelder habe eine mittlere Ausfallrate von $\overline{h(t)} = \lambda = 1/100\,000\text{ h} = 1\text{E}-5/\text{h}$. Er werde jährlich (alle $T = 8760\text{ h}$) getestet und falls erforderlich ersetzt. Wie groß ist die Wahrscheinlichkeit, dass er im Fall eines Brandes nicht funktioniert?

Es muss die mittlere Nichtverfügbarkeit über das Testintervall T_{test} bestimmt werden:

$$\overline{Q(0..T)} = \frac{1}{T} \int_0^T 1 - e^{-\lambda \cdot t} dt = \frac{1}{T} \left(t + \frac{1}{\lambda} e^{-\lambda \cdot t} \right) \Bigg|_0^T = 1 + \frac{e^{-\lambda \cdot T} - 1}{\lambda \cdot T} = 0,0426$$

Wenn die Nichtverfügbarkeit klein ist (etwa $Q < 0,1$), gilt sehr gut die konservative Näherung:

$$\overline{Q(0..T_{\text{test}})} \lesssim 0,5 \cdot \lambda \cdot T_{\text{test}} \quad (41)$$

Im vorherigen Beispiel ergäbe sich $\overline{Q(0..T_{\text{test}})} \approx 0,0438$ statt 0,0426.

Die Nichtverfügbarkeit ist zwar wie die Unzuverlässigkeit eine Wahrscheinlichkeit, kann also nur Werte von 0 bis 1 annehmen, sie ist jedoch im Gegensatz zur Unzuverlässigkeit nur im Sonderfall eines nicht-testbaren Systems monoton steigend. Nach jedem Test hingegen fällt die Nichtverfügbarkeit $Q(t)$ im Gegensatz zur Unzuverlässigkeit $F(t)$ wieder auf null (bzw. im Fall eines nicht vollständigen Tests zumindest auf einen Wert nahe null).

4.4 Zeit zur Reparatur, MRT

Falls durch den Defekt nur eine Teilfunktion betroffen ist, der Weiterbetrieb des umgebenden größeren Systems aber (ggf. mit Einschränkungen) möglich ist, muss auch die für die Reparatur und/oder Ersatzbeschaffung benötigte Zeit (MRT, Mean Repair Time) in der Verfügbarkeit berücksichtigt werden. Zusammen mit der Zeit zur Entdeckung (MTTD) ergibt sich die mittlere Wiederherstellungszeit (Mean Time To Restore, MTTR):

$$\text{MTTR} = \text{MTTD} + \text{MRT} \quad (42)$$

Zur exakten Berechnung der mittleren Nichtverfügbarkeit kann man von deren Definition ausgehen:

$$\begin{aligned} \bar{Q} &= \frac{T_{\text{defekt}}}{T_{\text{gesamt}}} = \frac{T_{\text{ud}} + p_{\text{def}} \cdot \text{MRT}}{(1 - p_{\text{def}}) \cdot T_{\text{test}} + p_{\text{def}} \cdot (T_{\text{test}} + \text{MRT})} \\ &= \frac{T_{\text{ud}} + p_{\text{def}} \cdot \text{MRT}}{T_{\text{test}} + p_{\text{def}} \cdot \text{MRT}} \end{aligned} \quad (43)$$

Dabei bezeichnet p_{def} die Wahrscheinlichkeit, die Funktion zum regelmäßigen Testzeitpunkt T_{test} als defekt vorzufinden

$$p_{\text{def}} = F(T_{\text{test}}) = 1 - e^{-\lambda \cdot T_{\text{test}}}$$

und T_{ud} die mittlere Zeit in jedem Testintervall, während der die Funktion aufgrund des Defekts unerkannt nicht verfügbar ist (also quasi die MTTD aufgeteilt auf alle Testintervalle bis zum Ausfall)

$$\begin{aligned} T_{\text{ud}} &= \int_0^{T_{\text{test}}} f(t) \cdot (T_{\text{test}} - t) dt = \int_0^{T_{\text{test}}} \lambda \cdot e^{-\lambda \cdot T_{\text{test}}} \cdot (T_{\text{test}} - t) dt \\ &= \frac{e^{-\lambda \cdot T_{\text{test}}} - 1}{\lambda} + T_{\text{test}} \end{aligned}$$

Durch Einsetzen in Formel (43) ergibt sich für die mittlere Nichtverfügbarkeit

$$\begin{aligned} \bar{Q} &= \frac{\frac{e^{-\lambda \cdot T_{\text{test}}} - 1}{\lambda} + T_{\text{test}} + \text{MRT} \cdot (1 - e^{-\lambda \cdot T_{\text{test}}})}{T_{\text{test}} + \text{MRT} \cdot (1 - e^{-\lambda \cdot T_{\text{test}}})} \\ &= \frac{e^{-\lambda \cdot T_{\text{test}}} - 1}{\lambda \cdot T_{\text{test}} + \lambda \cdot \text{MRT} \cdot (1 - e^{-\lambda \cdot T_{\text{test}}})} + 1 \end{aligned} \quad (44)$$

Für vernachlässigbare Detektionszeit $T_{\text{test}} \rightarrow 0$ (kontinuierliche Diagnose, kleine Testintervalle oder Fehleroffenbarung durch unmittelbar erkennbare Fehlfunktion) geht die mittlere Nichtverfügbarkeit gegen

$$\bar{Q} = \frac{\lambda \cdot \text{MRT}}{\lambda \cdot \text{MRT} + 1} \quad (45)$$

wie man durch einmalige Anwendung der Regel von de l'Hospital auf Formel (44) erhält.

Für vernachlässigbare Reparaturzeit $\text{MRT} \rightarrow 0$ (z. B. Außerbetriebnahme während der Reparatur) vereinfacht sich Formel (44) unmittelbar zu

$$\bar{Q} = \frac{e^{-\lambda \cdot T_{\text{test}}} - 1}{\lambda \cdot T_{\text{test}}} + 1 \quad (46)$$

Wenn sowohl Testintervall T_{test} als auch Reparaturzeit MRT klein gegen $MTTF$ sind, gilt mit ausreichender Genauigkeit

$$\bar{Q} \lesssim \lambda \cdot (0,5 \cdot T_{\text{test}} + \text{MRT}) \quad (47)$$

Schwieriger ist die Herleitung einer (exakten) Formel für die Nichtverfügbarkeit zu einem bestimmten Zeitpunkt, wenn die Reparaturzeit nicht vernachlässigbar klein ist. Eine sehr gute Näherung ist durch

$$Q(t) = 1 - \frac{e^{-\frac{\lambda \cdot (t \bmod T_{\text{test}})}{\lambda \cdot \text{MRT} + 1}}}{\lambda \cdot \text{MRT} + 1} \quad (48)$$

gegeben (ohne Herleitung). Für Reparaturzeit $\text{MRT} \rightarrow 0$ geht Formel (48) unmittelbar in Formel (39) über:

$$Q(t) = 1 - e^{-\lambda(t \bmod T_{\text{test}})}$$

Für vernachlässigbare Detektionszeit $T_{\text{test}} \rightarrow 0$ ergibt sich unmittelbar

$$Q(t) = 1 - \frac{e^{-\frac{0}{\lambda \cdot \text{MRT} + 1}}}{\lambda \cdot \text{MRT} + 1} = 1 - \frac{1}{\lambda \cdot \text{MRT} + 1} = \frac{\lambda \cdot \text{MRT}}{\lambda \cdot \text{MRT} + 1} = \bar{Q}$$

also Formel (45).

4.5 Kontinuierliche Diagnose

Im Fall von kontinuierlicher vollständiger Diagnose (d. h. jeder Fehler wird sofort entdeckt) hat die Nichtverfügbarkeit gar nichts mit der (Un-)zuverlässigkeit zu tun, es gilt also immer $Q(t) \neq F(t)$. Die Nichtverfügbarkeit ist dann nur von der (mittleren) Ausfallrate $\bar{h} = \lambda$ und der für die Wiederherstellung nötigen Zeit MRT abhängig:

$$Q(t) = \bar{Q} = \frac{\lambda \cdot \text{MRT}}{\lambda \cdot \text{MRT} + 1} = \text{const} \quad (49)$$

Wird die Komponente nie getestet und ggf. repariert oder ersetzt, so ist $Q(t) = F(t)$. Dies mag in der (unbemannten) Raumfahrt der Fall sein, in der funktionalen Sicherheit jedoch nicht, da das regelmäßige Testen und ggf. Reparieren oder Ersetzen zentrale Maßnahmen der funktionalen Sicherheit darstellen. Daher dürfen Unzuverlässigkeit und Nichtverfügbarkeit niemals verwechselt werden. Des Weiteren darf niemals eine Unzuverlässigkeit und eine Nichtverfügbarkeit addiert oder multipliziert oder sonst mathematisch verknüpft werden!

Beispiel 4.2 Eine Komponente mit der konstanten Ausfallrate $h(t) = \lambda = 1/10\,000\text{ h} = 1\text{E}-4/\text{h}$ werde alle 5000 Stunden getestet und erforderlichenfalls umgehend repariert oder ausgetauscht. Wie hoch sind Nichtverfügbarkeit und Unzuverlässigkeit zu den Zeitpunkten $T=1999$ und $T=2001$ Stunden?

$$Q(1999\text{ h}) = 1 - e^{-\lambda \cdot (1999 - 15000\text{ h})} \approx 0,3935$$

$$Q(2001\text{ h}) = 1 - e^{-\lambda \cdot (2001 - 20000\text{ h})} \approx 0,0001$$

$$F(1999\text{ h}) = 1 - e^{-\lambda \cdot 1999\text{ h}} \approx 0,8646$$

$$F(2001\text{ h}) = 1 - e^{-\lambda \cdot 2001\text{ h}} \approx 0,8647$$

4.6 Betriebliche und sicherheitsbezogene Verfügbarkeit

Häufig hört man als Sicherheitsingenieur den Satz: „Der Ausfall ist unkritisch, der geht nur in die Verfügbarkeit.“ Dieser Satz basiert auf dem Unverständnis von Zuverlässigkeit und Verfügbarkeit. Beides können sicherheitsbezogene Größen sein, müssen aber nicht.

Beispiel 4.3 Die Verfügbarkeit eines Rauchmelders gibt an, mit welcher Wahrscheinlichkeit er eine Rauchentwicklung melden wird. Dies ist offensichtlich eine sicherheitsrelevante Größe, in vielen Anwendungen gibt es daher vorgeschriebene Mindestwerte (bzw. Maximalwerte für die Nichtverfügbarkeit). Je häufiger man ihn testet, umso größer wird die Verfügbarkeit. Je größer die Ausfallrate, um so häufiger muss man testen (und reparieren), um die geforderte Verfügbarkeit zu erreichen. Die Zuverlässigkeit (oder Ausfallrate) alleine erlaubt hier keine Aussage über die Sicherheit, da es für die Gebäudesicherheit irrelevant ist, wie oft der Rauchmelder kaputt geht – wenn der Fehler nur schnell detektiert und behoben wird. Vielmehr ist die Zuverlässigkeit hier eine betrieblich relevante Größe: Je schlechter die Zuverlässigkeit (also je größer die Ausfallrate) des Rauchmelders, desto häufiger muss man ihn testen und ersetzen, um die aus Sicherheitsgründen vorgegebene Verfügbarkeit zu erreichen.

4.7 Ausfallrate bei Tests

Aufgrund der Tests und ggf. Reparatur verliert die Dichtefunktion $f(t)$ ihre Bedeutung. An ihre Stelle tritt eine neue Größe, die meist als „Ausfallhäufigkeit“ bezeichnet wird. In [NUREG] wird für diese das Formelzeichen $w(t)$ verwendet, jedoch hat sich auch für diese Größe noch kein einheitliches Formelzeichen durchgesetzt. In Abbildung 10 sind

alle drei Größen $h(t)$, $w(t)$ und $f(t)$ für eine Funktion mit der konstanten Ausfallrate $h(t) = \lambda = 1\text{E-}5/\text{h}$, die alle 30000 Stunden getestet wird, dargestellt.

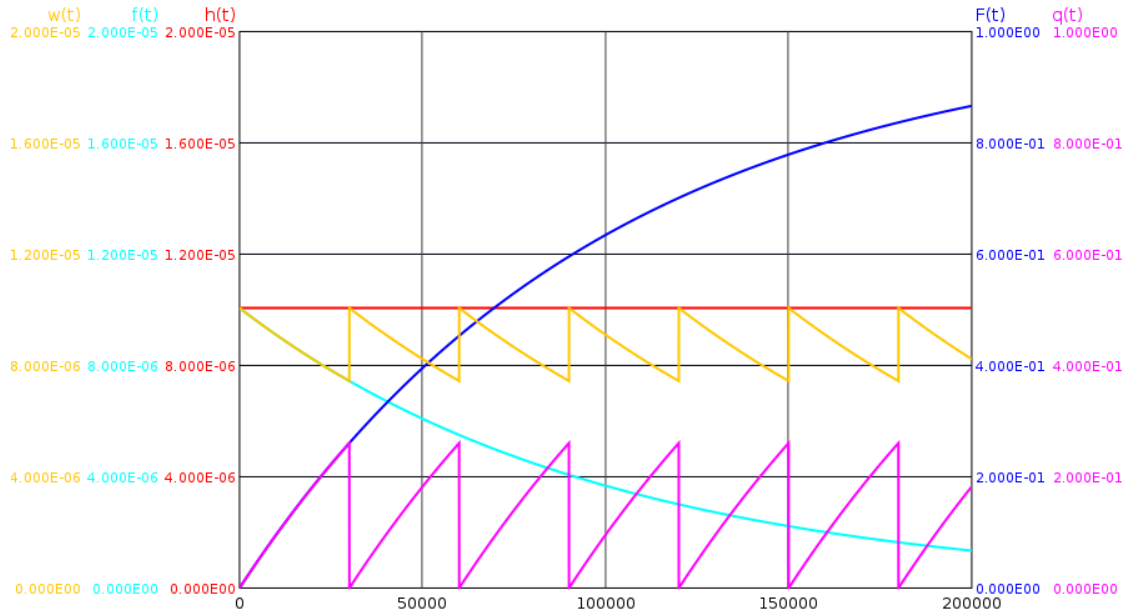


Abbildung 10: Größen bei regelmäßigen Tests

Im Gegensatz zur Dichte $f(t)$ wird $w(t)$ niemals null, da sich aufgrund der Reparatur das System ja immer (wieder) in einem Zustand befindet, in dem es (wieder) ausfallen kann. Unmittelbar nach (vollständigen) Tests, also wenn die Nichtverfügbarkeit auf null zurückgeht, steigt die Ausfallhäufigkeit wieder auf die Ausfallrate $h(t)$ an. Das Integral der Ausfallhäufigkeit $w(t)$ über die Zeit kann also beliebig groß werden. Nur bis zum ersten Ausfall bzw. dem ersten Test sind $w(t)$ und $f(t)$ identisch.

Die Ausfallrate, also die Häufigkeit des Übergangs in den Ausfallzustand unter der Bedingung, dass das System zum Zeitpunkt t ausfallfähig ist, berechnet sich nun zu

$$h(t) = \frac{w(t)}{A(t)} = \frac{w(t)}{1 - Q(t)} \quad (50)$$

wobei $A(t)$ die Verfügbarkeit bzw. $Q(t)$ die Nichtverfügbarkeit zum Zeitpunkt t bezeichnet.

5 Nichtverfügbarkeit von komplexen Funktionen

Die Nichtverfügbarkeit ist die wesentliche Größe für Sicherheitsfunktionen, die nur selten benötigt werden. Beispiele für einfache Komponenten, die solche Sicherheitsfunktionen wahrnehmen, sind Leitungsschutzschalter (soll bei Überstrom auslösen), Überdruckventile (soll bei Überdruck aufmachen) oder Deckensprinkler (soll bei zu hoher Temperatur Wasser freigeben). Ihre funktionale Architektur ist in 11 dargestellt.

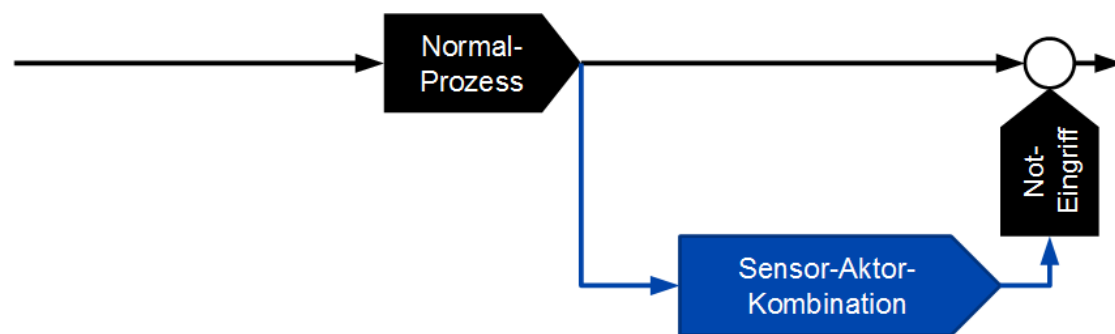


Abbildung 11: Einfaches Sicherheitssystem für seltene Anforderung

Selbstverständlich gibt es auch komplexere Systeme, die selten benötigte Sicherheitsfunktionen wahrnehmen, heutzutage meist computergesteuert. Beispiele sind Überwachungs- und Notfallsysteme in der chemischen Industrie oder in Kraftwerken, Brandmelde- und Brandbekämpfungsanlagen, Entrauchungsanlagen, Evakuierungssysteme etc. Ihre Architektur ist beispielhaft in Abbildung 12 dargestellt. Der Begriff „Prozess“ ist dabei sehr weit gefasst, das kann einfach der normale Betrieb eines Gebäudes, einer Apparatur oder einer Maschine sein.

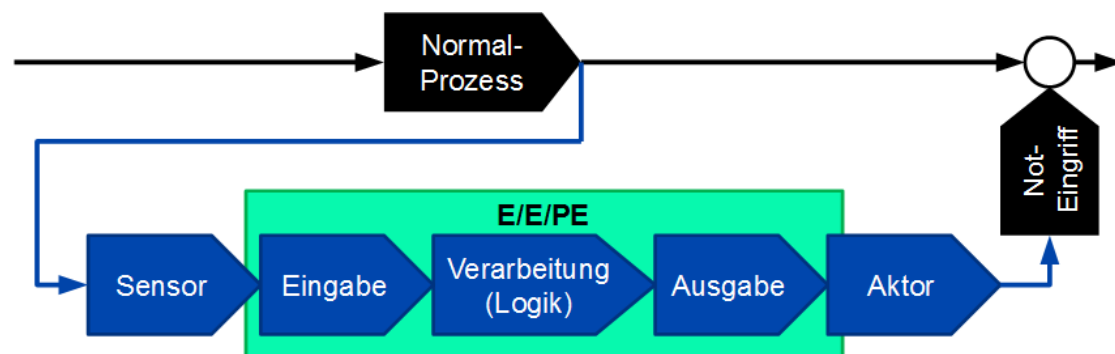


Abbildung 12: Komplexes Sicherheitssystem für seltene Anforderung

Im Normalfall bekommt man von der Existenz der Sicherheitsfunktion(en) nichts mit. Erst im Anforderungsfall (falls dieser überhaupt jemals eintritt) zeigt sich, ob die Sicherheitsfunktion tatsächlich verfügbar ist.⁴ Ein sicherheitskritischer Fehler (also einer, der die Sicherheitsfunktion im Anforderungsfall verhindert) ist nur erkennbar, wenn die

⁴Unter Umständen zeigt sich dann auch erst, ob die Sicherheitsfunktion korrekt ausgelegt ist, also beispielsweise die Aktoren richtig dimensioniert sind, aber das ist nicht Gegenstand der funktionalen Sicherheit

Komponente bzw. das System regelmäßig getestet wird ⁵ oder eben wenn sie im Anforderungsfall nicht funktioniert.

Die Nichtverfügbarkeit ist praktisch immer eine zeitabhängige Funktion $Q(t)$. Falls es keine Ereignisse mit konstanter Nichtverfügbarkeit gibt, ist die Nichtverfügbarkeit eines Systems Q_{sys} zur Zeit $t = 0$ null. Nur wenn es Ereignisse mit konstanter Nichtverfügbarkeit gibt, kann Q_{sys} schon bei $t = 0$ größer null sein. In jedem System wird es Komponenten geben, die mindestens einen Ausfallsmodus haben, der nicht sofort erkennbar ist. Die Nichtverfügbarkeit wird daher bis zum nächsten Test monoton ansteigen, und unmittelbar nach einem Test auf einen kleineren Wert (im Fall vollständiger Tests auf den Wert bei $t = 0$) abfallen. Gibt es Ausfälle, die nie erkannt werden, steigt die Nichtverfügbarkeit zumindest im Mittel bis zum Einsatzende des Systems an.

Da niemals bekannt ist, zu welchem Zeitpunkt die Sicherheitsfunktion benötigt wird, interessiert immer nur der Mittelwert der Nichtverfügbarkeit über die Lebenszeit des Prozesses oder des Sicherheitssystems:

$$\bar{Q} = \frac{1}{T_{\text{Life}}} \int_0^{T_{\text{Life}}} Q(t) dt \quad (51)$$

In [IEC 61508] wird dieser Mittelwert $\bar{Q}$ als Probability of Failure on Demand (kurz PFD) bezeichnet.

5.1 Berechnung mit Fehlerbäumen

Häufig wird das Sicherheitssystem mit Hilfe von Fehlerbäumen modelliert. Diese sind für die Modellierung solcher Systeme sehr gut geeignet, und die Nichtverfügbarkeit des Systems $\bar{Q}_{\text{sys}}$ kann sehr einfach und mathematisch exakt berechnet werden (natürlich vorausgesetzt, dass die Nichtverfügbarkeiten der Komponenten bekannt sind).

Wenngleich letztlich nur der Mittelwert der Nichtverfügbarkeit interessiert, so muss gemäß Formel (51) dennoch die zeitabhängige Funktion an ausreichend vielen Stützstellen berechnet und über diese integriert werden.

Die Basis-Ereignisse eines Fehlerbaums modellieren die Komponenten mit ihren Ausfällen sowie gegebenenfalls die Maßnahmen zur Wiederherstellung. Das Standard-Modell für ein Basis-Ereignis ist das sogenannte „wiederherstellbare Ereignis“, auch als testbares oder reparierbares Ereignis bezeichnet, siehe Anhang A.1. Dieses Modell beschreibt eine (konstante) Ausfallrate und eine (mittlere) Detektionszeit sowie gegebenenfalls auch die Reparaturzeit. Wenn der Ausfall nicht durch Diagnose oder Tests erkannt wird, also bis zum Ende der Einsatzzeit im System enthalten bleibt, muss das Ereignis mit dem Modell „nicht-wiederherstellbares Ereignis“ gemäß Anhang A.2 beschrieben werden. In diesem Fall sind auch nicht-konstante Ausfallraten möglich. Manchmal ist die Nichtverfügbarkeit im Anforderungsfall auch weder von einer Zeit seit einem letzten Test noch

⁵für bestimmte Komponenten mag dabei auch eine visuelle Inspektion ausreichen

vom Alter des Systems abhängig, oder das Ereignis beschreibt gar keinen Ausfall, sondern die Wahrscheinlichkeit des Vorhandenseins einer externen Randbedingung oder die Wahrscheinlichkeit eines Bedienfehlers. Dann ist die Nichtverfügbarkeit eine Konstante (Anhang A.3).

Als logische Verknüpfungen kommen fast ausschließlich UND und ODER zum Einsatz, darum sollen auch nur diese hier betrachtet werden.⁶

Auch wenn man heute zur Berechnung Binäre Entscheidungsdiagramme (engl. Binary Decision Diagrams, kurz BDD) verwendet, so soll dennoch die Berechnung zunächst mithilfe von Minimalschnitten (engl. Minimal Cut-Sets, MCS) erklärt werden.

Ein Minimalschnitt ist eine Kombination von Basis-Ereignissen, die zum Eintreten des Top-Ereignisses (beispielsweise dem Ausfall einer Sicherheitsfunktion) notwendig und hinreichend ist. Bei sogenannten kohärenten Fehlerbäumen – das sind Fehlerbäume, die keine negierenden Gatter wie NOT, XOR, NAND etc. enthalten – gibt es genau einen Satz von Minimalschnitten. Bei inkohärenten Fehlerbäumen spricht man von Prim-Implikanten anstelle von Minimalschnitten, und es gibt im Allgemeinen mehrere mögliche Sätze von Prim-Implikanten. Da negierende Gatter praktisch nie benötigt werden, werden sie im Folgenden nicht erwähnt.

5.1.1 Nichtverfügbarkeit einer UND-Verknüpfung

Eine UND-Verknüpfung von zwei oder mehr Basis-Ereignissen führt zu einem Minimalschnitt mit eben diesen Basis-Ereignissen. Eine UND-Verknüpfung von Zweigen eines Baums führt in der Regel auch zu längeren Minimalschnitten, die genaue Anzahl und Länge hängt von der Struktur der verknüpften Zweige ab.

Die Wahrscheinlichkeit, dass ein Minimalschnitt zu einer Zeit t erfüllt ist, also die von einem Minimalschnitt ausgehende Nichtverfügbarkeit, beträgt

$$Q_{\text{MCS}}(t) = \prod_{j=1}^m Q_j(t) \quad (52)$$

wobei m die Anzahl der Basisereignisse in diesem Minimalschnitt ist. Die Anzahl m bezeichnet man als Ordnung des Minimalschnitts.

Beispiel 5.1 *In einem Zimmer befinden sich zwei Brandmelder. Jeder hat eine Ausfallrate von $\lambda = 1\text{E}-5/\text{h}$. Die beiden Brandmelder werden etwa alle 10 000 h gleichzeitig getestet und im Fehlerfall umgehend ersetzt. Mit welcher Wahrscheinlichkeit meldet nicht mindestens einer von ihnen im Brandfall den Brand?*

Abbildung 13 zeigt den entsprechenden Fehlerbaum. Er besteht aus zwei Basis-Ereignissen vom Typ „wiederherstellbares Ereignis“, welche durch ein UND-Gatter verknüpft sind.

Es gibt nur einen Minimalschnitt, nämlich $\{BM.1 \& BM.2\}$. Da er zwei Elemente (Literale) enthält, ist es ein Minimalschnitt zweiter Ordnung. Folglich gilt mit Formeln (52)

⁶Sogenannte Mehrheitsentscheider (M-aus-N) sind nichts anderes als eine Abkürzung für ein ODER-Gatter über mehreren UND-Gattern, diese sind also eingeschlossen. Siehe hierzu Abschnitt 6.1.6.

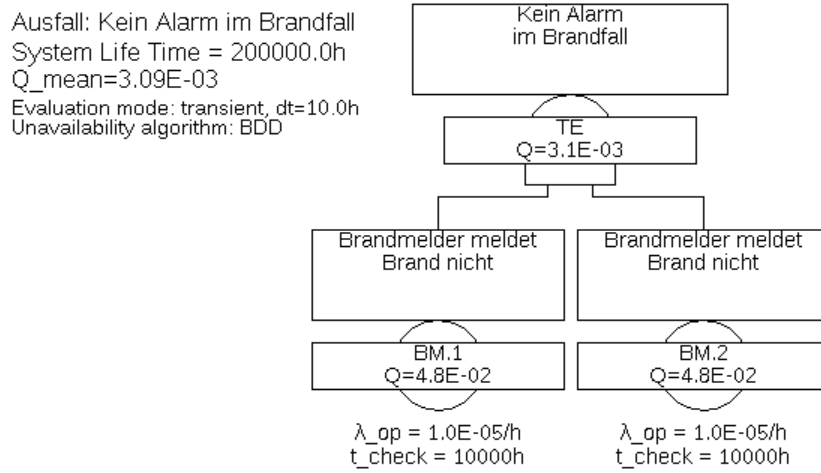


Abbildung 13: Redundante Brandmelder

für die Nichtverfügbarkeit des Minimalschnitts und (39) für die Nichtverfügbarkeiten der Brandmelder

$$Q_{\text{sys}}(t) = Q_{\text{BM.1}}(t) \cdot Q_{\text{BM.2}}(t) = Q_{\text{BM}}^2(t) = \left(1 - e^{-\lambda(t \bmod T_{\text{test}})}\right)^2 = 1 - 2e^{-\lambda(t \bmod T_{\text{test}})} + e^{-2\lambda(t \bmod T_{\text{test}})}$$

Mit den genannten Größen ergibt sich eine Periodizität mit einer Periodendauer von 10000 h, der genaue Verlauf ist in Abbildung 14 dargestellt.

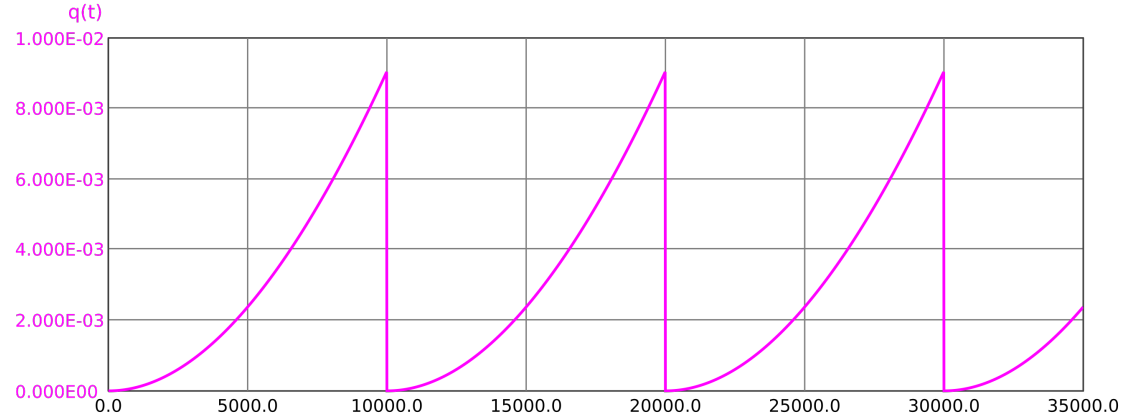


Abbildung 14: Zeitlicher Verlauf der Nichtverfügbarkeit zweier gleichartig redundanter Komponenten, die regelmäßig zu denselben Zeiten getestet werden (Ausschnitt)

Aufgrund der Periodizität genügt es, den Mittelwert über eine Periode zu berechnen:

$$\begin{aligned} \bar{Q} &= \frac{1}{T_{\text{Life}}} \int_0^{T_{\text{Life}}} Q(t) dt = \frac{1}{10000 \text{ h}} \int_0^{10000 \text{ h}} 1 - 2e^{-\lambda t} + e^{-2\lambda t} dt \\ &= \frac{1}{10000 \text{ h}} \left[t + \frac{2e^{-\lambda t}}{\lambda} - \frac{e^{-2\lambda t}}{2\lambda} \right]_0^{10000 \text{ h}} = 0,00309459... \approx 3,1\text{E-}3 \end{aligned}$$

Die System-Einsatzzeit (Lebenszeit) spielt aufgrund der regelmäßigen Tests keine Rolle.

Der Leser mag durch eigene Rechnung feststellen, dass bei Verwendung der vereinfachten Formel $Q_{\text{BM}}(t) \lesssim \lambda \cdot t$ anstelle der hier verwendeten exakten Formel $Q_{\text{BM}}(t) = 1 - \exp(-\lambda t)$ praktisch dasselbe Ergebnis herauskommt.

5.1.2 Nichtverfügbarkeit einer ODER-Verknüpfung

Eine ODER-Verknüpfung von zwei oder mehr Basis-Ereignissen führt zu entsprechend vielen Minimalschnitten. Eine ODER-Verknüpfung von Zweigen eines Baums führt in der Regel auch zu mehreren Minimalschnitten, die genaue Anzahl hängt von der Struktur der verknüpften Zweige ab.

Die Gesamt-Nichtverfügbarkeit des Systems ist näherungsweise die Summe der Nichtverfügbarkeiten der n Minimalschnitte:

$$Q_{\text{sys}}(t) \lesssim \sum_{i=1}^{n_{\text{MCS}}} Q_{\text{MCS},i}(t) = \sum_{i=1}^{n_{\text{MCS}}} \left(\prod_{j=1}^{m_{\text{Lit},i}} Q_j(t) \right) \quad (53)$$

Diese Formel ist eine Näherung, die nur gilt, wenn die Einzel-Nichtverfügbarkeiten sehr klein sind.

Eine bessere Näherung, die sich fast ebenso leicht berechnen lässt, ist die Esary-Proschan-Formel:

$$Q_{\text{sys}}(t) \lesssim 1 - \prod_{i=1}^{n_{\text{MCS}}} (1 - Q_{\text{MCS},i}(t)) \quad (54)$$

Diese Näherung kann in der Praxis gut verwendet werden, da sie immer konservativ ist (also $Q_{\text{sys}}(t)$ nie zu klein schätzt), für kleine Nichtverfügbarkeiten gegen das exakte Ergebnis tendiert, und für große Nichtverfügbarkeiten nicht größer als eins wird.⁷

Das exakte Ergebnis erhält man durch disjunkte Zerlegung der Minimalschnitte. Ein Verfahren zur disjunkten Zerlegung ist in [EN 61025] beschrieben. Dieses eignet sich jedoch nur für sehr kleine Fehlerbäume⁸.

Binäre Entscheidungsdiagramme (BDDs) können auch für sehr große Fehlerbäume mit geringem Aufwand erstellt werden, ohne dass überhaupt Minimalschnitte ermittelt werden müssen. Zudem implizieren sie bei der Berechnung schon die Disjunktion. Sie erlauben daher eine exakte Berechnung der Nichtverfügbarkeit mit deutlich geringerem Aufwand als die Näherung über Minimalschnitte. Und schließlich sind BDDs die mit Abstand schnellste Methode zum Ermitteln der Minimalschnitte. Moderne FTA-Werkzeuge nutzen daher BDDs für alle Operationen.

Beispiel 5.2 Eine automatische Brandlöschanlage besteht im Prinzip aus einem Brandmelder (BM), einer Steuerung (STRG) und einer Löschereinheit (LE). Ein Brand wird nur dann gelöscht, wenn diese drei Einheiten im Falle eines Brandes funktionieren.

⁷über Minimal-Pfade kann man auch eine untere Grenze schätzen, diese ist jedoch bei praktischen Aufgaben so weit vom tatsächlichen Wert entfernt, dass sie bedeutungslos ist

⁸und für diese ist die Überschneidung der Minimalschnitte bei korrekt ausgelegten Systemen ohnehin gering, eine disjunkte Zerlegung also unnötig

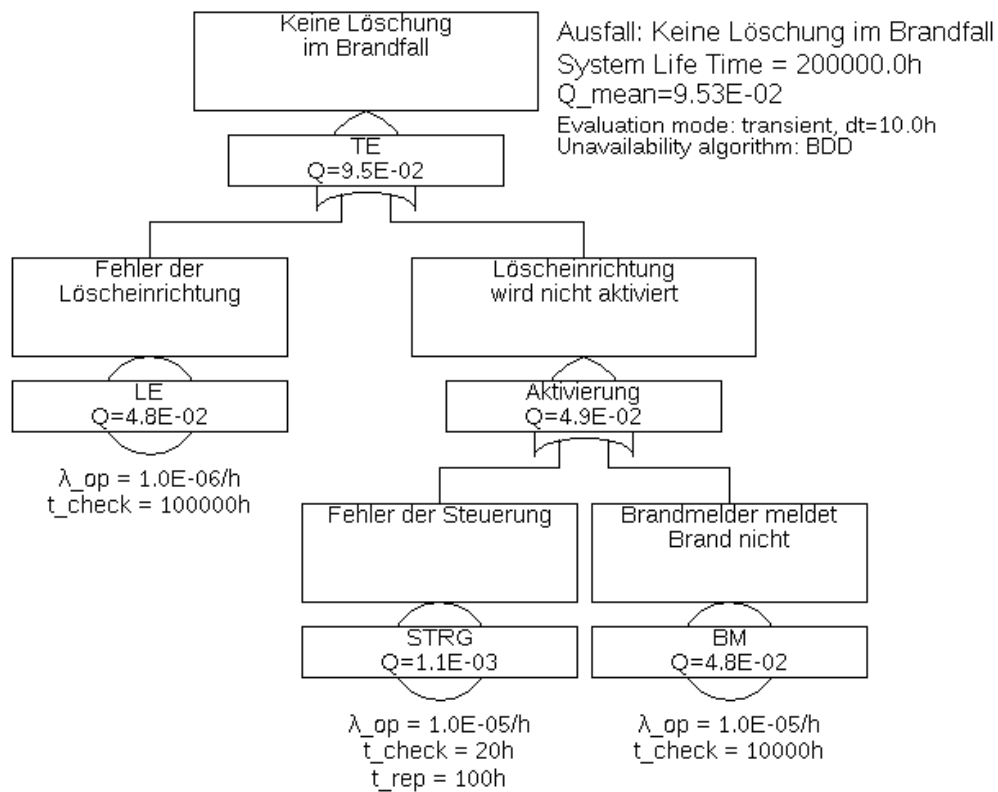


Abbildung 15: Brandlöschanlage

Dies wird durch den in Abbildung 15 dargestellten Fehlerbaum modelliert. Mathematisch könnte man alle drei Basisereignisse direkt unter das obere ODER-Gatter setzen, dies würde jedoch der FTA-Regel „Top-Down-Entwurf“ widersprechen. Diese Regel besagt, dass ein Fehlerbaum stets vom Top-Ereignis aus nach unten entwickelt werden soll und ist eine der wichtigsten Regeln überhaupt. Und wenn man überlegt, warum die Löschanlage nicht löscht, kann das unmittelbar nur daran liegen, dass sie selbst nicht funktioniert oder dass sie nicht aktiviert wird. Steuerung und Brandmelder kommen erst bei der Frage ins Spiel, warum die Löschanlage nicht aktiviert wird, also eine Ebene tiefer.

Es gibt drei Minimalschnitte, nämlich $\{LE\}$, $\{STRG\}$ und $\{BM\}$. Alle drei sind erster Ordnung. Verwendet man die Näherungsformel (53) für die System-Nichtverfügbarkeit, so erhält man

$$Q_{\text{sys}}(t) \lesssim Q_{LE}(t) + Q_{STRG}(t) + Q_{BM}(t)$$

und damit für den Mittelwert

$$\overline{Q_{\text{sys}}} \lesssim \overline{Q_{LE}} + \overline{Q_{STRG}} + \overline{Q_{BM}}$$

Mit den in Abbildung 15 erwähnten Werten und den Näherungsformeln (41) bzw. (47) erhält man schließlich

$$\begin{aligned} \overline{Q_{\text{sys}}} &\approx 0,5\lambda_{LE}T_{\text{Test},LE} + \lambda_{STRG}(0,5T_{\text{Test},STRG} + T_{\text{MRT},STRG}) + 0,5\lambda_{BM}T_{\text{Test},BM} \\ &= 0,05 + 0,0011 + 0,05 = 0,1011 \end{aligned}$$

Diese Näherungsrechnung weicht vom (hier nicht hergeleiteten) exakten Wert $Q = 0,0953...$ um nur 5% ab — eine für die Praxis längst ausreichende Genauigkeit.

Verwendet man die Abschätzung nach Esary-Proschan (54), so erhält man

$$Q_{\text{sys}}(t) \lesssim 1 - [(1 - Q_{\text{LE}}(t)) \cdot (1 - Q_{\text{STRG}}(t)) \cdot (1 - Q_{\text{BM}}(t))]$$

Mit denselben Näherungen wie zuvor für die Einzel-Nichtverfügbarkeiten ergibt sich für die mittlere System-Nichtverfügbarkeit

$$\begin{aligned} \overline{Q_{\text{sys}}} &\lesssim 1 - [(1 - 0,5\lambda_{\text{LE}}T_{\text{Test,LE}}) \\ &\quad \cdot (1 - \lambda_{\text{STRG}}(0,5T_{\text{Test,STRG}} + T_{\text{MRT,STRG}})) \cdot (1 - 0,5\lambda_{\text{BM}}T_{\text{Test,BM}})] \\ &= 1 - [(1 - 0,05) \cdot (1 - 0,0011) \cdot (1 - 0,05)] = 0,09849... \end{aligned}$$

Diese Näherung weicht vom exakten Wert $Q = 0,0953...$ um nur 3% ab.

5.1.3 Nichtverfügbarkeit von Kombinationen von UND- und ODER-Verknüpfungen

Für so einfache Systeme wie in den bisherigen Beispielen wird man kaum einen Fehlerbaum verwenden. Praktisch bestehen Fehlerbäume immer aus einer Mehrzahl von UND- und ODER-Gattern, welche oft eine Vielzahl von Basis-Ereignissen verknüpfen.

Beispiel 5.3 Abschließend sollen die beiden vorherigen Beispiele kombiniert werden. Die beiden Rauchmelder seien dabei wieder redundant, also nebeneinander montiert und so verschaltet, dass einer von beiden ausreicht, um einen Brand zu melden.

Der Fehlerbaum ist in Abbildung 16 gezeigt.

Die drei Minimalschnitte sind: $\{LE\}$, $\{STRG\}$, $\{BM.1 \ \& \ BM.2\}$

Gemäß Näherungsformel (53) gilt für die System-Nichtverfügbarkeit etwa:

$$Q_{\text{sys}}(t) \lesssim \sum_{j=1}^n Q_{\text{MCS},i}(t) = Q_{\text{LE}}(t) + Q_{\text{STRG}}(t) + Q_{\text{BM.1}}(t) \cdot Q_{\text{BM.2}}(t)$$

Für den interessierten und mit BDDs vertrauten Leser soll der Vollständigkeit halber noch das BDD angegeben werden.

Wählt man die Variablenordnung Löscheinrichtung (LE), Steuerung (STRG), Brandmelder.1 (BM.1), Brandmelder.2 (BM.2), so erhält man das in Abbildung 17 gezeigte binäre Entscheidungsdiagramm BDD.

Aus dem BDD kann direkt eine exakte Formel für die System-Nichtverfügbarkeit abgeleitet werden:

$$Q_{\text{sys}}(t) = Q_{\text{LE}}(t) + (1 - Q_{\text{LE}}(t)) \cdot [Q_{\text{STRG}}(t) + (1 - Q_{\text{STRG}}(t)) \cdot (Q_{\text{BM.1}}(t) \cdot Q_{\text{BM.2}}(t))]$$

In dieser Formel sind automatisch alle Ereignisse disjunkt. Es sei angemerkt, dass sich bei anderen Variablenordnungen andere Formeln ergeben, diese sind jedoch alle mathematisch äquivalent.

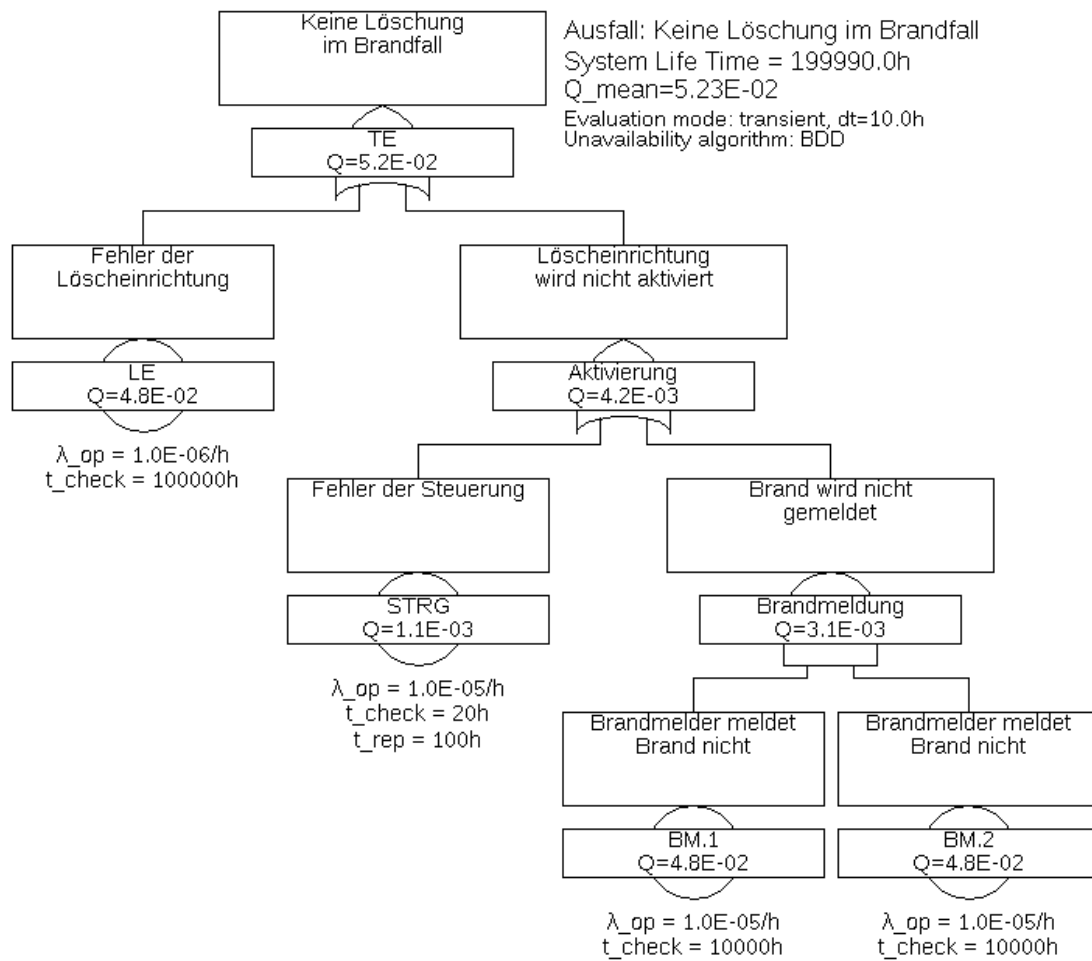


Abbildung 16: Brandlöschanlage mit redundanten Sensoren

Vergleicht man im letzten Beispiel die exakte Formel mit der Näherungsformel, so sieht man unmittelbar, dass die Näherungsformel (53) für alle Fehlerbäume, die keine negierenden Gatter enthalten, immer ein zu großes Ergebnis liefert. Für kleine Fehlerbäume ist der Unterschied bei allen korrekt ausgelegten Systemen⁹ vernachlässigbar, bei großen Fehlerbäumen mit vielen Tausend Minimalschnitten kann der Fehler jedoch selbst dann sehr groß werden. Daher können große Fehlerbäume praktisch nur mit BDDs berechnet werden, zumal schon die Ermittlung von Minimalschnitten bei großen Fehlerbäumen praktisch nur mit Hilfe von BDDs (noch besser mit ternären Entscheidungsdiagrammen) möglich ist.

5.1.4 Transiente und stationäre Berechnung, Rechnen mit Mittelwerten

Im Allgemeinen muss zur Ermittlung der mittleren Nichtverfügbarkeit eines Systems das Integral gemäß Formel (51) berechnet werden, so wie in Beispiel 5.1 gezeigt. Praktisch bedeutet das, dass der Fehlerbaum für viele Zeitpunkte berechnet werden muss, was

⁹korrekte Auslegung bedeutet, dass die Testintervalle den Ausfallraten angemessen sind, so dass alle Nichtverfügbarkeiten jederzeit sehr klein sind

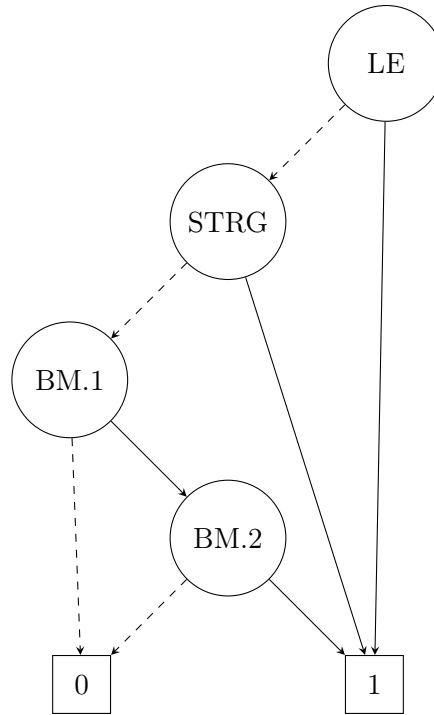


Abbildung 17: BDD für die Brandlöschanlage mit redundanten Sensoren

für große Fehlerbäume auch mit modernen Rechnern einige Zeit in Anspruch nehmen kann. Hierbei kann natürlich eine eventuell vorhandene Periodizität ausgenutzt werden, wie ebenfalls in Beispiel 5.1 geschehen. Gibt es keine Periodizität (zum Beispiel weil es mindestens ein Ereignis ohne regelmäßige Tests gibt), so wird sich kein quasi-stationärer Zustand¹⁰ einstellen. In diesem Fall muss die Berechnung immer gemäß Formel (51), also numerische Integration über die Lebenszeit, erfolgen. Da diese Berechnung auch transiente, also nicht periodische Vorgänge korrekt berücksichtigt, wird sie auch als transiente Berechnung bezeichnet, in [ASTRA TM] einfach als zeitabhängige Berechnung.

Um die Rechenzeit zu verringern, kann man auf die Idee kommen, den Fehlerbaum nur einmal mit den Mittelwerten der Nichtverfügbarkeiten der Basis-Ereignisse zu berechnen. Diese Rechnung geht von einem eingeschwungenen quasi-stationären Zustand aus, und wird daher auch als stationäre Berechnung bezeichnet.

Die Rechnung mit Mittelwerten ist jedoch auch im eingeschwungenen Zustand nicht korrekt, denn hierdurch würden Integral und Produkt in der Reihenfolge vertauscht, was mathematisch falsch ist:

$$\overline{Q_{\text{MCS}}} = \frac{1}{T_{\text{Life}}} \int_0^{T_{\text{Life}}} Q(t) dt = \frac{1}{T_{\text{Life}}} \int_0^{T_{\text{Life}}} \prod_{i=1}^n Q_i(t) dt \neq \prod_{i=1}^n \frac{1}{T_{\text{Life}}} \int_0^{T_{\text{Life}}} Q_i(t) dt = \prod_{i=1}^n \overline{Q_i}$$

Die Größe des Fehlers, der bei der Berechnung mit Mittelwerten entsteht, hängt von

¹⁰quasi-stationär bedeutet, dass die Nichtverfügbarkeit zwar um einen Mittelwert schwanken kann, der Mittelwert sich aber mit der Zeit nicht verändert

vielen Parametern ab. Im Falle von zwei gleichartigen UND-verknüpften Ereignissen wie in Beispiel 5.1, welche zur selben Zeit getestet werden, ist das berechnete Ergebnis etwa $1/3$ zu klein:

$$\overline{Q_{\text{BM.1}}} \cdot \overline{Q_{\text{BM.2}}} \approx 4,8\text{E}-2 \cdot 4,8\text{E}-2 \approx 2,3\text{E}-3 \dots \neq 3,1\text{E}-3$$

Der Fehler $1/3$ rührt aus der Integration des quadratischen Terms her, der für die in Abbildung 14 deutlich sichtbaren Parabelabschnitte verantwortlich ist. Formelmäßig wird das besonders deutlich, wenn man die Näherungsformel $Q(t) \lesssim \lambda \cdot t$ verwendet:

$$\begin{aligned} \overline{Q_{\text{korrekt}}} &= \frac{1}{T} \int_0^T Q_1(t) \cdot Q_2(t) dt \approx \frac{1}{T} \int_0^T \lambda t \cdot \lambda t dt = \frac{\lambda^2}{3T} T^3 = \frac{\lambda^2 T^2}{3} \\ \overline{Q_{\text{falsch}}} &= \frac{1}{T} \int_0^T Q_1(t) dt \cdot \frac{1}{T} \int_0^T Q_2(t) dt \approx \left(\frac{1}{T} \int_0^T \lambda t dt \right)^2 = \left(\frac{\lambda}{2T} T^2 \right)^2 = \frac{\lambda^2 T^2}{4} \end{aligned}$$

Bei höheren Potenzen, also Minimalschnitten höherer Ordnung, wird der relative Fehler noch größer, allerdings ist deren absoluter Beitrag in der Regel nur gering. Eine Berechnung mit Mittelwerten kann in der Praxis also für Überschlagsrechnungen verwendet werden, die abschließende Berechnung sollte aber immer gemäß Formel (51) erfolgen, was eine numerische Integration erforderlich macht.

Es sei angemerkt, dass eine stationäre Berechnung auch mit Maximalwerten anstatt mit Mittelwerten ausgeführt werden kann. In dem Fall ist die ermittelte Nichtverfügbarkeit immer (sehr) konservativ.

5.2 Berechnung mit Markov Modellen

Die System-Nichtverfügbarkeit kann auch mittels Markov-Modellen berechnet werden. Markov-Modelle stellen die Zustände dar, in denen sich ein System befinden kann, sowie die Übergänge (Transitionen) zwischen den Zuständen. Bei klassischen Markov-Modellen werden die Transitionen mittels Übergangsraten beschrieben. Transitionen vom Ursprungszustand weg, also insbesondere Ausfälle, werden meist mit λ abgekürzt. Transitionen in Richtung des Ursprungszustands, also Maßnahmen der Wiederherstellung, werden meist mit μ abgekürzt. Damit ist ein Markov-Modell mathematisch durch ein lineares Differenzialgleichungssystem beschrieben:

$$\dot{\vec{p}}(t) = A(t) \vec{p}(t) \quad (55)$$

Dabei ist $A(t)$ die (im Allgemeinen zeitabhängige) Transitionsmatrix und $\vec{p}$ der Vektor der Aufenthaltswahrscheinlichkeiten der Systemzustände.

Da sich das System zu jeder Zeit in genau einem Zustand befindet, muss die Summer aller Zustandswahrscheinlichkeiten stets eins sein:

$$\|\vec{p}(t)\| = \sum_{i=1}^n p_i(t) = 1 \quad (56)$$

Die Summe der Aufenthaltswahrscheinlichkeiten in den Zuständen $p_j(t) \in \vec{p}(t)$, in denen die Sicherheitsfunktion nicht gegeben ist, gibt die System-Nichtverfügbarkeit an:

$$Q(t) = \sum_{j=1}^m p_j(t) \quad (57)$$

Die mittlere Nichtverfügbarkeit ist wieder durch Formel (51) gegeben.

Die Wiederherstellungsrate μ wird in der Fachliteratur fast immer als Kehrwert der mittleren Wiederherstellungszeit definiert: ¹¹

$$\mu \stackrel{\text{def}}{=} 1/\text{MTTR} \quad (58)$$

Für Fehler, die durch regelmäßige Tests entdeckt werden, ergibt sich damit für die Wiederherstellungsrate

$$\mu = 1/\text{MTTR} = \frac{1}{0,5 \cdot T_{\text{test}}} = 2/T_{\text{test}} \quad (59)$$

Beispiel 5.4 Abbildung 18 zeigt das Markov-Modell für die Brandmeldung mittels redundanter Brandmelder, wie in Beispiel 5.1 betrachtet.

Ausfall: Kein Alarm im Brandfall
System Life Time = 200000.0h
Q_mean=2.34E-03
Evaluation mode: steady-state

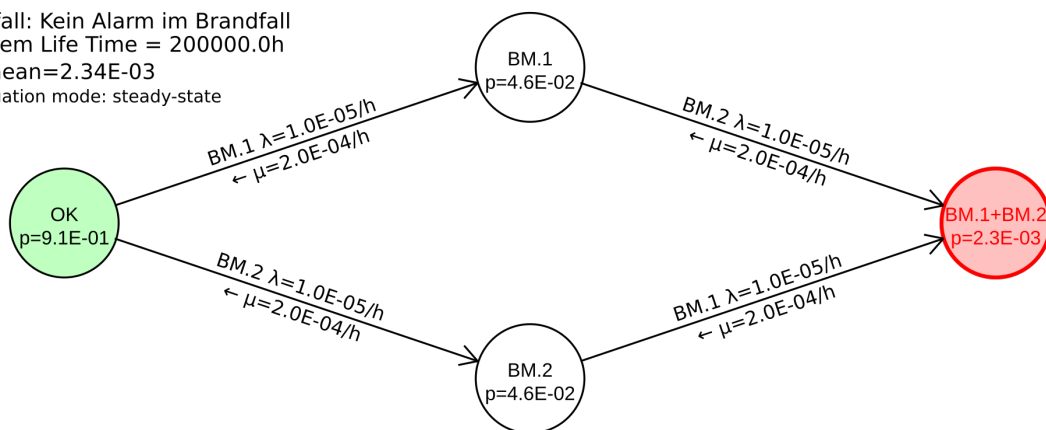


Abbildung 18: Redundante Brandmelder, stationäre Berechnung

Die Ausfallraten λ für die Brandmelder sind jeweils über den Transitionspfeilen angegeben, die Wiederherstellungsraten μ jeweils darunter und mit einem kleinen Pfeil für die gegensätzliche Richtung versehen. ¹²

Mit dem Zustandsvektor

$$\vec{p} = \begin{pmatrix} \text{OK} \\ \text{BM.1} \\ \text{BM.2} \\ \text{BM.1 + BM.2} \end{pmatrix}$$

¹¹Dass es sich hierbei tatsächlich um eine Definition und nicht um eine sachlich begründbare Formel handelt, wird in Beispiel 5.4 mit Beispiel 5.5 sichtbar.

¹²Häufig werden separate Linien für die Wiederherstellung dargestellt, die Darstellung mit nur einer Linie erscheint jedoch übersichtlicher.

gilt für das lineare Differenzialgleichungssystem

$$\begin{pmatrix} -2\lambda & \mu & \mu & 0 \\ \lambda & -\mu - \lambda & 0 & \mu \\ \lambda & 0 & -\mu - \lambda & \mu \\ 0 & \lambda & \lambda & -2\mu \end{pmatrix} \vec{p}(t) = \dot{\vec{p}}(t) \quad (60)$$

5.2.1 Stationäre Berechnung

Wenn jeder Ausfall detektierbar ist, gibt es aus jedem Zustand auch eine Transition heraus. Folglich werden sich die Zustände nach beliebig langer Zeit im Gleichgewicht befinden, die zeitliche Ableitung des Zustandsvektors also zu null werden. Sind alle Detektions- und Reparaturzeiten relativ kurz im Verhältnis zur Lebenszeit des Systems, wird das Gleichgewicht nach relativ kurzer Zeit praktisch eingenommen sein.

Die Aufenthaltswahrscheinlichkeiten in diesem stationären Systemzustand kann man leicht berechnen, indem man $\dot{\vec{p}}(t) = 0$ setzt und dann eine beliebige Gleichung durch die Summe der Zustandswahrscheinlichkeiten ersetzt, welche immer eins sein muss.

Da der stationäre Zustand ewig währt, hat der Einschwingvorgang keinen signifikanten Einfluss auf das Integral in Formel (51), der Mittelwert der Nichtverfügbarkeit ist also etwa gleich der Nichtverfügbarkeit im stationären Zustand:

$$\overline{Q_{\text{sys}}} \approx Q_{\text{stat}} \quad (61)$$

Beispiel 5.5 Ersetzt man in Gleichungssystem (60) die vierte Zeile durch die Summenzeile, so ist die stationäre Lösung durch das folgende lineare Gleichungssystem beschrieben:

$$\begin{pmatrix} -2\lambda & \mu & \mu & 0 \\ \lambda & -\mu - \lambda & 0 & \mu \\ \lambda & 0 & -\mu - \lambda & \mu \\ 1 & 1 & 1 & 1 \end{pmatrix} \vec{p}_{\text{stat}} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}$$

Für den Zustandsvektor im eingeschwungenen Zustand ergibt sich

$$\vec{p}_{\text{stat}} = \begin{pmatrix} \text{OK} \\ \text{BM.1} \\ \text{BM.2} \\ \text{BM.1} + \text{BM.2} \end{pmatrix} = \begin{pmatrix} \frac{\mu^2}{\mu^2 + 2\lambda\mu + \lambda^2} \\ \frac{\lambda\mu}{\mu^2 + 2\lambda\mu + \lambda^2} \\ \frac{\lambda\mu}{\mu^2 + 2\lambda\mu + \lambda^2} \\ \frac{\lambda^2}{\mu^2 + 2\lambda\mu + \lambda^2} \end{pmatrix}$$

Die Nichtverfügbarkeit ist die Aufenthaltswahrscheinlichkeit des Zustands BM.1+BM.2, also $\overline{Q_{\text{stat}}} = \frac{\lambda^2}{\mu^2 + 2\lambda\mu + \lambda^2}$.

Mit den für Beispiel 5.1 verwendeten Zahlenwerten $\lambda = 1,0\text{E-}5/\text{h}$ und $T_{\text{test}} = 10\,000\text{ h}$ ergibt sich $\mu = 2/(10\,000\text{ h}) = 2,0\text{E-}4/\text{h}$ und damit

$$\overline{Q_{\text{stat}}} = \frac{(1,0\text{E-}5/\text{h})^2}{(2,0\text{E-}4/\text{h})^2 + 2 \cdot 1,0\text{E-}5/\text{h} \cdot 2,0\text{E-}4/\text{h} + (1,0\text{E-}5/\text{h})^2} \approx 0,0023$$

Die mittlere Nichtverfügbarkeit wurde in Beispiel 5.1 exakt zu $\overline{Q_{\text{sys}}} = 0,003\,094\dots$ berechnet, das über die stationäre Auswertung des Markov-Modells ermittelte Ergebnis ist also deutlich zu optimistisch. Dies liegt zum Einen daran, dass Formel (58) nur für kontinuierliche Wartung und Reparatur gilt, und Formel (59) immer etwas optimistisch ist, zum Anderen aber auch an der Struktur des Markov-Modells, welches die Realität offensichtlich nicht richtig widerspiegelt — siehe hierzu das nächste Beispiel.

5.2.2 Transiente Berechnung

In vielen praktischen Anwendungen wird der eingeschwungene Zustand nicht einmal ansatzweise erreicht, da die Einsatzdauer zu kurz ist. Falls es Ausfälle gibt, die nicht detektiert oder repariert werden können, gibt es sogar absorbierende Zustände, der stationäre Zustand ist also durch die Kumulation der Aufenthaltswahrscheinlichkeit in einem oder mehreren Ausfallzuständen gegeben, so dass $Q(t \rightarrow \infty) = 1$ gilt¹³. In diesem Fall ist die stationäre Lösung des Differenzialgleichungssystems uninteressant, statt dessen muss die mittlere Nichtverfügbarkeit mittels Formel (51) während des Übergangs vom Ursprungszustand bis zum Ende der Einsatzzeit des Systems berechnet werden. Dies erfordert eine numerische Integration des Differenzialgleichungssystems.

Die numerische Integration eröffnet Möglichkeiten der Modellierung, die über klassische Markov-Modelle hinausgehen. Insbesondere ist es möglich, zeitlich veränderliche Übergangsraten zu verwenden, und sogar Übergänge zu bestimmten Zeitpunkten zu berücksichtigen. Letzteres wiederum ermöglicht die realitätsnahe Berücksichtigung von regelmäßigen Tests.

Beispiel 5.6 Abbildung 19 zeigt ein Markov-Modell für die redundanten Brandmelder, bei welchem berücksichtigt wird, dass die Tests und somit die Wiederherstellung nicht kontinuierlich, sondern zu bestimmten Zeitpunkten erfolgt. Außerdem ist berücksichtigt, dass bei Defekt beider Brandmelder (Zustand „BM.1+BM.2“) beide Defekte zur selben Zeit erkannt werden und auch die Reparatur zur selben Zeit stattfindet, also das System in den Ursprungszustand überführt wird.

Anmerkung: Am Transitionspeil von „OK“ zu „BM.1+BM.2“ ist keine Ausfallrate angegeben, diese ist daher null. Nur die regelmäßige Wiederherstellung alle 10000 h ist hier relevant, diese steht unter dem Transitionspeil und ist mit einem kleinen Pfeil nach links angedeutet. Umgekehrt steht unter den Transitionspeilen von „BM.1“ und „BM.2“ zu „BM.1+BM.2“ keine Wiederherstellung, diese ist also null.

Das Ergebnis dieser Modellierung und der Berechnung mit einer Integrationsschrittweite

¹³absorbierende Zustände sind immer Fehlerzustände, ansonsten ist das Modell nicht korrekt

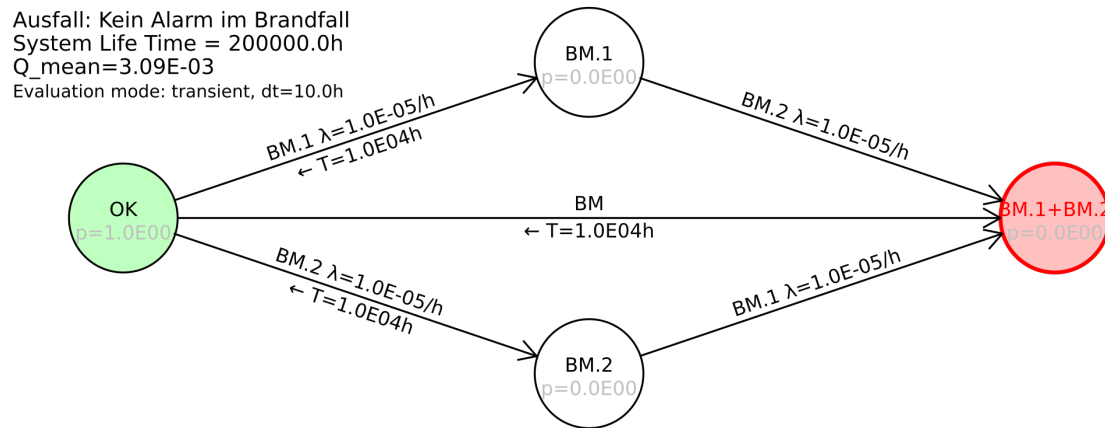


Abbildung 19: Redundante Brandmelder mit Wiederherstellung zu diskreten Zeitpunkten

von 10 Stunden ist $Q = 3,09\text{E-}3$ und stimmt nun praktisch mit dem exakten Wert überein.

Klassische Markov-Modelle mit konstanten Übergangsraten sind für die Berechnung der Nichtverfügbarkeit nur bedingt geeignet, die Ergebnisse sind meist zu optimistisch. Erweiterte Markov-Modelle mit zeitdiskreten Übergängen ermöglichen die realitätsnahe Modellierung und Berechnung und sind daher wesentlich besser geeignet.

Es muss erwähnt werden, dass die Berechnung zeitdiskreter Übergänge eine hinreichend kleine Integrationsschrittweite voraussetzt. Für kleine Testintervalle oder gar kontinuierliche Diagnose müssen konstante Übergangsraten verwendet werden.

Beispiel 5.7 Abbildung 20 zeigt das Markov-Modell für die in Beispiel 5.3 eingeführte Brandlöschanlage mit redundanten Brandmeldern. Die Wiederherstellung der Steuerung STRG wird als kontinuierlicher Übergang behandelt, da die Integrationsschrittweite mit 10 h nur unwesentlich kleiner ist als das Testintervall (20 h). Das Ergebnis stimmt mit dem des Fehlerbaums überein.

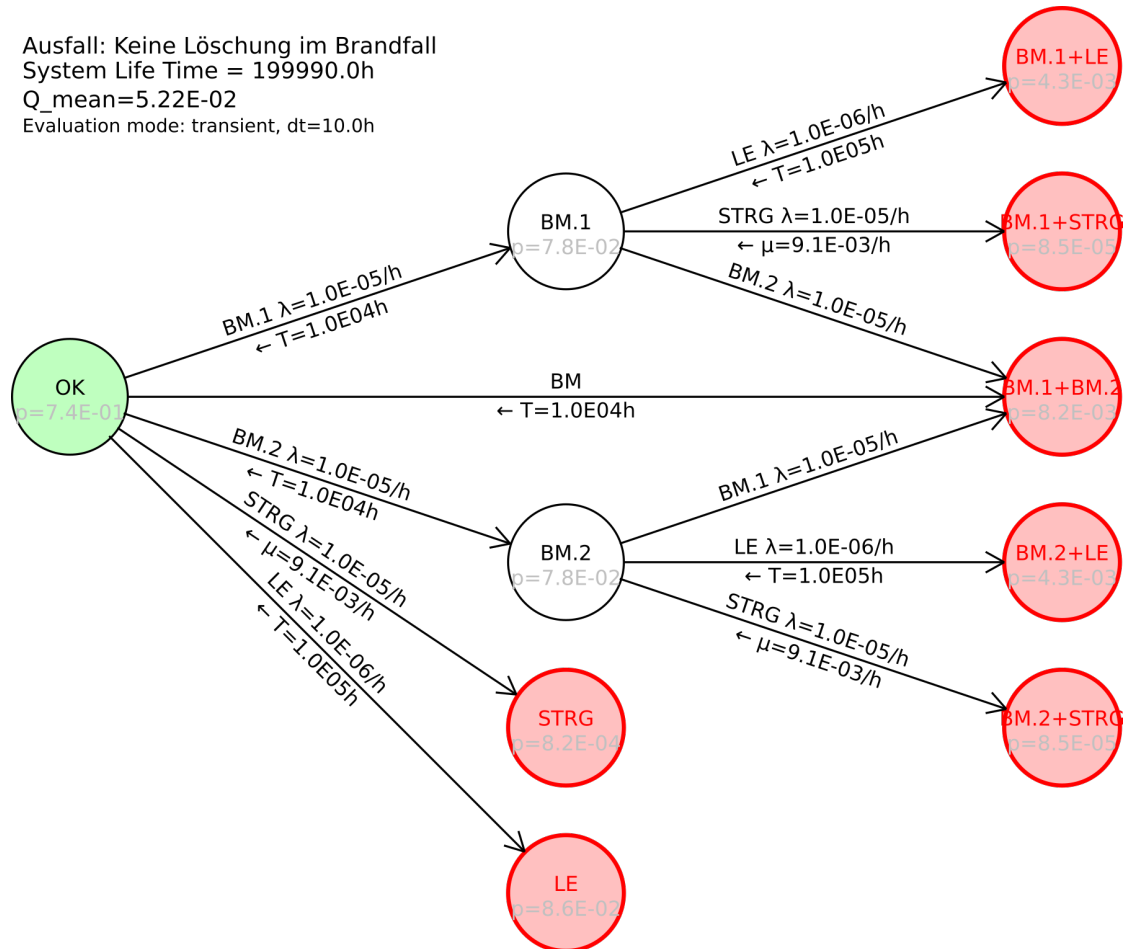


Abbildung 20: Brandlöschanlage

6 Ausfallrate von komplexen Funktionen

Die Ausfallrate ist die wesentliche Größe für Sicherheitsfunktionen, die kontinuierlich oder zumindest häufig (quasi-kontinuierlich) benötigt werden. Beispiele für Systeme, die kontinuierlich benötigte Sicherheitsfunktionen implementieren, sind Flugantriebe (sollen immer den geforderten Schub liefern und vor allem nicht fälschlich stehenbleiben), Lageregelungen (sollen immer die vorgegebene Lage sicherstellen), Antriebsregelungen (sollen immer geforderte Drehmomente, Geschwindigkeiten oder Positionen gewährleisten oder nicht zur Unzeit anlaufen), Eisenbahn-Signalanlagen (sollen nie einen zu freigiebigen Signalbegriff anzeigen bzw. einen zu freigiebigen Fahrbefehl geben) aber auch Airbag-Steuerungen (sollen niemals fälschlich den Airbag auslösen), ABS-Steuerungen (sollen den Bremsdruck nie zu stark reduzieren), Zug-Türsteuerungen (sollen die Tür nie zur falschen Zeit öffnen). Sobald die Sicherheitsfunktion versagt, ist unmittelbar eine gefährliche Situation gegeben. Das Wort „unmittelbar“ heißt nicht, dass zwingend ein Schaden entstehen muss, sondern nur, dass bei üblichen externen Bedingungen ein Schaden nicht unwahrscheinlich ist.¹⁴

Beispiele für Systeme, die quasi-kontinuierlich benötigte Sicherheitsfunktionen implementieren, gehören Fahrwerke von Flugzeugen (müssen nur bei der Landung funktionieren, aber die Landung ist unausweichlich), Bremsen von sämtlichen Fahrzeugen (müssen nur bei Bremsanforderung funktionieren, aber die Anforderung kommt fast sicher – nur im Ausnahmefall wird ein Ausrollen möglich sein).

Die prinzipielle Struktur solcher Sicherheitsfunktionen ist in Abbildung 21 dargestellt.

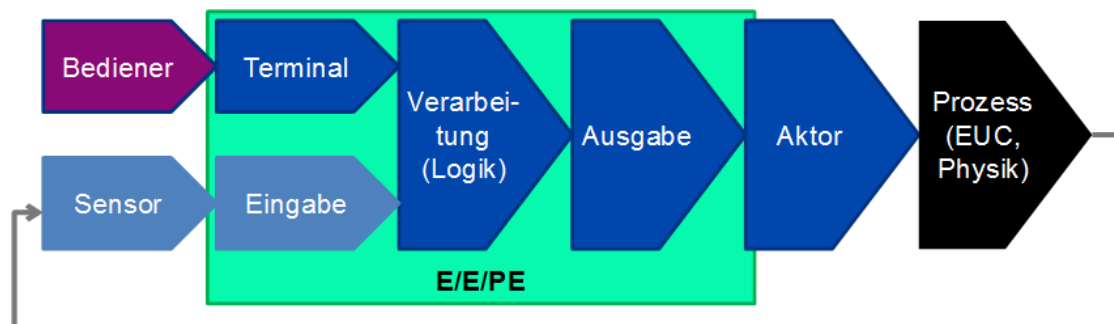


Abbildung 21: *Typische Architektur von Sicherheitsfunktionen mit kontinuierlicher Anforderung*

Charakteristisch ist, dass ein Ausfall der Sicherheitsfunktion unmittelbar Einfluss auf das Verhalten des physikalischen Prozesses hat (welcher etwa durch die Bewegungsgleichungen des Fahrzeugs, des Flugzeugs, der Maschine oder die Reaktionsdynamik der Chemikalien und Apparate gegeben ist) und somit zur Gefährdung führt – egal ob der Schaden sofort oder erst nach einer absehbaren Zeit eintritt.

Für kontinuierliche Sicherheitsfunktionen ist der Begriff der Prozessfehlertoleranzzeit

¹⁴Im Gegensatz zu den Sicherheitsfunktionen mit seltener Anforderungen, die überhaupt erst bei einer unüblichen externen Bedingung (Anforderung) benötigt werden.

(PFTZ, auch Prozesssicherheitszeit genannt) existenziell: Das ist die Zeit, für die die Sicherheitsfunktion verletzt werden darf, ohne dass sich hieraus eine Gefährdung ergibt. Im Fall einer hochdynamischen Antriebsregelung oder einer Lageregelung eines Kampffjets sind dies maximal wenige Millisekunden, im Fall einer Bremse mögen es wenige Sekunden sein, im Fall der Brennstoffzufuhr eines Großkraftwerks vielleicht einige Minuten. Eventuell vorhandene Diagnosemaßnahmen müssen in dieser Zeit den Fehler erkennen und eine adäquate Reaktion initiieren, zum Beispiel eine Sicherheitsabschaltung veranlassen oder die defekte Steuerung isolieren und einen redundanten Steuerpfad aktivieren. Eine reine Fehlermeldung an den Bediener ist bei kontinuierlichen oder quasi-kontinuierlichen Sicherheitsfunktionen in der Regel nicht ausreichend, da der Bediener die Funktion meist nicht wiederherstellen kann, bevor der Schaden eintritt.

Die Ausfallrate, welche man für ein System, welches über eine Lebenszeit T hinweg betrieben wird, praktisch feststellen wird, ist unmittelbar gegeben durch die Anzahl der beobachtbaren Ausfälle in der Lebenszeit dividiert durch selbige, siehe Formel (62), identisch zu Formel (29):

$$\overline{h_{\text{sys}}}(T) = \frac{N_{\text{sys}}(T)}{T} = \frac{1}{\text{MTTF}(T)} \quad (62)$$

Die mittlere Zeit zwischen zwei Ausfällen MTTF¹⁵ kann bei bekannter Ausfalldichtenfunktion $f(t)$ mit Formel (28) berechnet werden:

$$\text{MTTF}(T) = \frac{\int_0^T t \cdot f(t) dt + T}{F(T)} - T \quad (63)$$

Die Ausfalldichtenfunktion $f(t)$ kann für beliebige Zeitpunkte t mit Fehlerbäumen oder durch transiente Auswertung von Markov Modellen berechnet werden.

In [IEC 61508] wird dieser Mittelwert $\bar{h}$ als Probability of Failure per Hour (kurz PFH) bezeichnet.¹⁶

Die obigen Formeln sind dabei für beliebige Ausfallraten- bzw. Ausfalldichtefunktionen gültig, und auch unabhängig davon, ob das System während der Lebenszeit T wahrscheinlich nie oder wahrscheinlich mehrfach ausfallen wird, und auch unabhängig davon, ob das System oder Teile davon regelmäßig inspiziert und repariert werden.

Auf oberster Ebene stellt $\overline{h_{\text{sys}}}$ die Gefährdungsrate dar, oft mit HR (für engl. hazard rate) abgekürzt (vgl. [EN 50126]). Wenn es sich bei dem betrachteten System nur um

¹⁵Man mag hier anstelle von MTTF die Bezeichnung MTBF (Mean Time Between Failures) erwartet haben. Das wäre begrifflich tatsächlich korrekter, aber leider wird die Bezeichnung MTBF (überflüssigerweise) in der Literatur anders definiert und steht daher hier nicht zur Verfügung. Da andererseits die aus Abschnitt 3 bekannten Überlegungen und Formeln auch für ein System aus vielen Komponenten und mit Tests und Reparatur gelten, besteht auch kein triftiger Grund für eine andere Bezeichnung.

¹⁶Anmerkung: Einige Formeln im informativen Anhang B in [IEC 61508-6] sind hierzu nicht konsistent, die Bedeutung der PFH als synonyme Begriff zu $\bar{h}$ wie hier definiert geht aus den übrigen Teilen der Norm jedoch klar hervor und ist die einzig sinnvolle Definition.

eine Teilfunktion einer Sicherheitsfunktion handelt, wird $\overline{h_{\text{sys}}}$ entsprechend als Funktional Failure Rate (FFR) bezeichnet.

$\overline{h_{\text{sys}}}$ ist das relevante Maß für die Sicherheit für alle Sicherheitsfunktionen, bei deren Versagen es ohne weitere Bedingungen zu einem Schaden kommen kann (also Sicherheitsfunktionen mit kontinuierlicher oder zumindest häufiger Anforderung).

Anmerkung: Viele Steuerungen nehmen sowohl kontinuierliche Sicherheitsfunktionen wahr, als auch selten benötigte. In diesem Fall muss für die Steuerung sowohl die Nichtverfügbarkeit im Anforderungsfall $\overline{Q}$ als auch die Häufigkeit eines falschen Kommandos an die Aktorik $\overline{h}$ bestimmt werden.¹⁷

6.1 Berechnung mit Fehlerbäumen

Die Berechnung der Ausfallrate $\overline{h}$ ist wesentlich schwieriger als die Berechnung der Nichtverfügbarkeit $\overline{Q}$, sowohl bezüglich der Aufstellung eines korrekten Fehlerbaums als auch bezüglich der eigentlichen mathematischen Berechnung. Dennoch ist für die meisten Sicherheitsfunktionen die Fehlerbaumanalyse eine geeignete Methode zur Bestimmung der Ausfallrate, und oft die einzige praktikable Methode der Modellierung. Es gibt jedoch leider keinerlei Standardisierung hierfür¹⁸, obwohl die wesentlichen mathematischen Grundlagen schon in [NUREG] erwähnt sind. Daher ist es unerlässlich, dass sich der Analyst intensiv mit den Eigenschaften und Eigenheiten des verwendeten Werkzeugs vertraut macht, und im Zweifelsfall anhand von einfachen Tests von der korrekten Funktion des Werkzeugs überzeugt. Die folgenden Beispiele können Basis solcher Tests sein.

6.1.1 System ohne Redundanzen

Im einfachsten Fall wird die Sicherheitsfunktion von einer Anzahl n Komponenten realisiert, welche alle für die Sicherheitsfunktion zwingend erforderlich sind. Fällt eine Komponente aus, fällt die Sicherheitsfunktion aus. Der Fehlerbaum besteht nur aus ODER-Gattern.

Beispiel 6.1 *Der in Abbildung 22 gezeigte Fehlerbaum beschreibt solch ein System, welches aus den Komponenten Sensorik, Logik und Aktorik besteht. Dabei sei angenommen, dass das System zwingend laufen muss, es also keine Möglichkeit einer Not-Abschaltung im Falle von erkannten Fehlern gibt. Die ist häufig der Fall, zum Beispiel kann ein Flugzeug nicht einfach in einen sicheren Zustand gebracht werden, wenn die Geschwindigkeits-sensorik als fehlerhaft erkannt wird, denn diese wird für den Weiterflug bis zur Landung zwingend benötigt.*

Versagt eine dieser n Komponenten auf gefährliche Weise¹⁹, ist die Sicherheitsfunktion

¹⁷In der Regel sind hierfür zwei unterschiedliche Fehlerbäume oder Markov-Modelle nötig, da sich insbesondere die Diagnose bezüglich $\overline{Q}$ von der für $\overline{h}$ unterscheidet, und daher andere Basisereignisse (auf Basis einer anderen FMEDA) nötig sind und oft auch andere Gatter.

¹⁸[EN 61025] macht keinerlei Angaben, wie Fehlerbäume zur Berechnung von Ausfallraten verwendet werden können – weder bezüglich der Modellierung noch bezüglich der Berechnung).

¹⁹bezüglich gefährlicher und ungefährlicher Ausfälle siehe spätere Beispiele

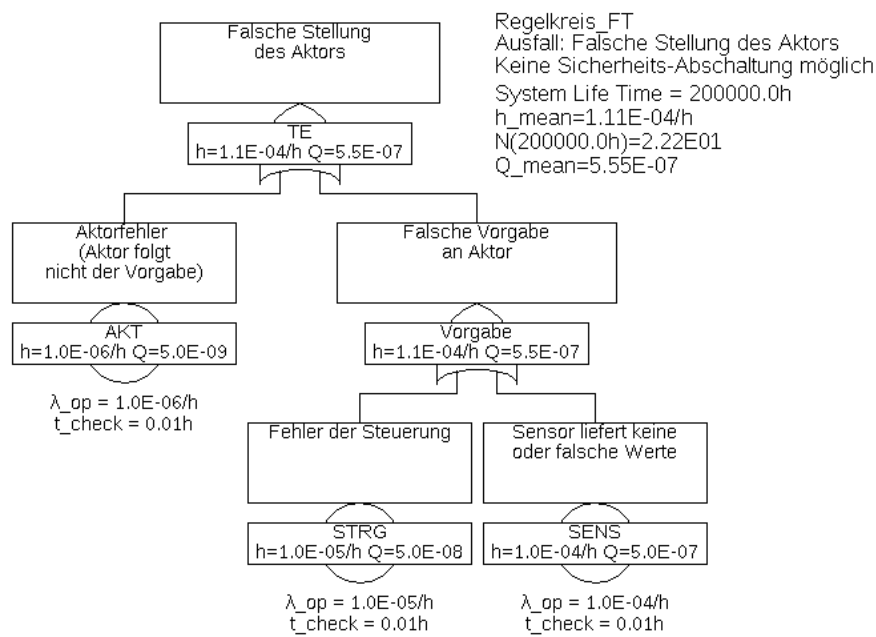


Abbildung 22: Fehlerbaum einer einfachen Sicherheitsfunktion mit kontinuierlicher Anforderung

nicht mehr gewährleistet. Die Minimalschnitte sind offensichtlich: $\{SENS\}$, $\{STRG\}$, $\{AKT\}$.

Für die Gesamtausfallrate gilt die schon aus Abschnitt 3 bekannte Formel

$$h(t) = \sum_{i=1}^n h_i(t) \quad (64)$$

Da jeder Fehler unmittelbar zum Ausfall führt, spielen weder System-Lebenszeit noch Fehler-Detektions- oder Reparaturzeiten eine Rolle, sondern ausschließlich die Ausfallraten der Komponenten.

Nimmt man für alle Komponenten konstante Ausfallraten an, so ergibt sich mit obiger Formel auch gleich die mittlere Ausfallrate, mit den im Fehlerbaum angegebenen Werten also $h(t) = \text{const} = \bar{h} = 1,0\text{E-}6/\text{h} + 1,0\text{E-}5/\text{h} + 1,0\text{E-}4/\text{h} = 1,11\text{E-}4/\text{h}$.

Dieses Beispiel war zweifellos trivial, und kaum jemand würde auf die Idee kommen, für solch ein System überhaupt einen Fehlerbaum (oder ein Markov-Modell) aufzustellen.

6.1.2 System mit Redundanzen

Wirklich interessant und als Modell nahezu unverzichtbar werden Fehlerbäume erst für Systeme mit Redundanzen (Mehrkanaligkeit). Im Fall von Redundanzen führt nicht jeder Einzelausfall zum Versagen der Sicherheitsfunktion, im Fehlerbaum wird also mindestens ein UND-Gatter enthalten sein, und es wird mindestens einen Minimalschnitt geben, der mehr als ein Basis-Ereignis enthält, also eine Ordnung größer eins hat.

Beispiel 6.2 In Beispiel 6.1 ist offensichtlich die Sensorik eine Schwachstelle. Daher sollen nun zwei Sensoren in einer redundanten Anordnung verwendet werden, also so, dass beide Sensoren dieselbe physikalische Größe messen. Wie schon in Beispiel 6.1 sei angenommen, dass das System zwingend laufen muss, es also keine Möglichkeit einer Not-Abschaltung im Falle von erkannten Fehlern gibt.

Wenn ein Sensor gar keine Werte oder offensichtlich falsche Werte liefert, kann nun der Messwert des anderen Sensors verwendet werden. Wenn jedoch nicht klar ist, welcher von beiden Sensoren defekt ist, kann die Steuerung auch weiterhin keine korrekte Stellgröße berechnen. Dasselbe gilt auch in dem Fall, dass ein Sensor bekanntermaßen defekt ist, und nun noch der zweite ausfällt, bevor der erste repariert wurde.

Die Sensoren müssen nun also bezüglich ihrer Fehlermodi unterschieden werden: Es gibt Defekte der Sensoren, die von der Steuerung erkannt werden können (SENS_ED, wie beispielsweise Drahtbruch), und solche, die nicht von der Steuerung erkannt werden können (SENS_NED). Der hierzu gehörige Fehlerbaum ist in Abbildung 23 gezeigt.

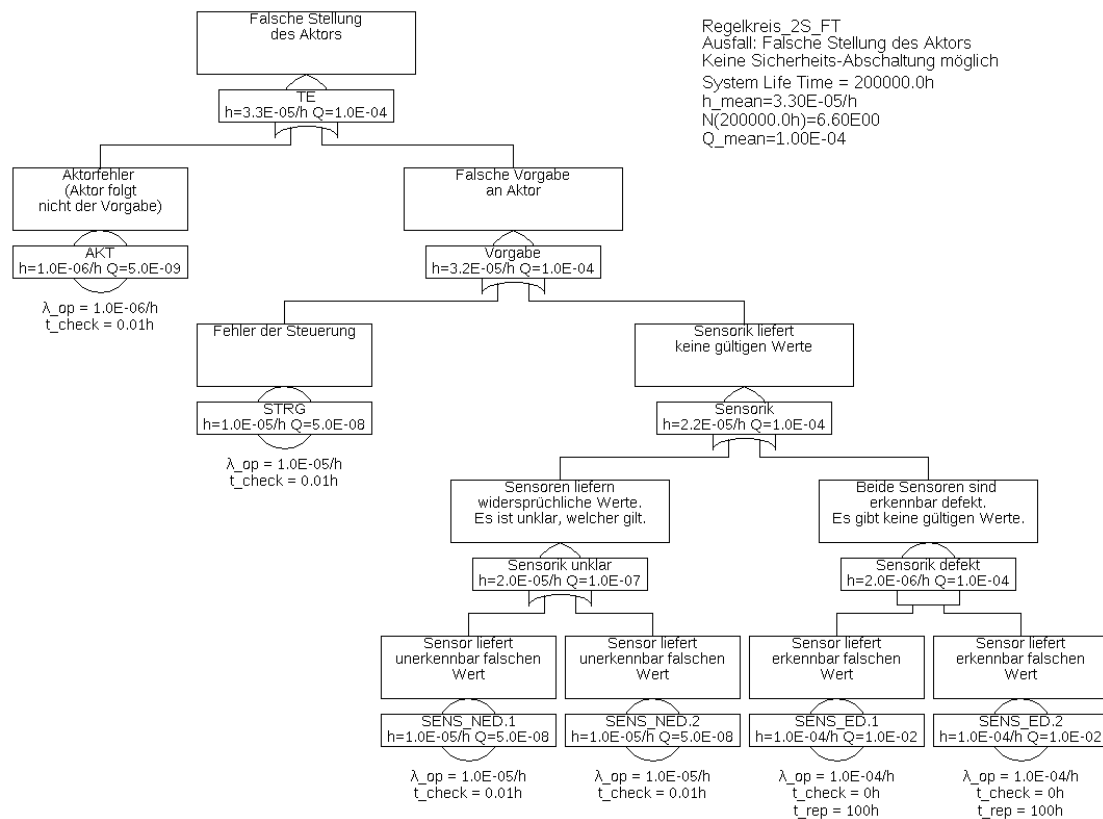


Abbildung 23: Fehlerbaum einer Sicherheitsfunktion mit kontinuierlicher Anforderung und redundanten Sensoren

Die Minimalschnitte sind:

- {AKT}
- {STRG}

- $\{SENS_NED.1\}$
- $\{SENS_NED.2\}$
- $\{SENS_ED.1 \text{ \& } SENS_ED.2\}$

Die Frage ist nun, wie die Eintrittsrate $h_{MCS}(t)$ des Minimalschnitts $\{SENS_ED.1 \text{ \& } SENS_ED.2\}$ berechnet werden kann. Die Aussage dieses Minimalschnitts ist folgende:

1. Sensor 1 ist bekanntermaßen defekt (also nicht verfügbar, $Q_{SENS_ED.1}$), und nun fällt auch noch Sensor 2 aus (mit Ausfallrate $\lambda_{SENS_ED.2}$)
ODER
2. Sensor 2 ist bekanntermaßen defekt (also nicht verfügbar, $Q_{SENS_ED.2}$), und nun fällt auch noch Sensor 1 aus (mit Ausfallrate $\lambda_{SENS_ED.1}$).

Für den Minimalschnitt gilt daher ²⁰

$$h_{MCS}(t) \lesssim \lambda_{SENS_ED.1} \cdot Q_{SENS_ED.2}(t) + \lambda_{SENS_ED.2} \cdot Q_{SENS_ED.1}(t)$$

Bei den Ereignissen $SENS_ED.x$ wird nun auch die Nichtverfügbarkeit benötigt. Diese hängt gemäß Formel (48) im Allgemeinen von der Zeit zur Detektion und der Reparaturzeit ab. Da in diesen Ereignissen nur die Ausfälle betrachtet werden, die sofort erkennbar sind, wird die Detektionszeit zu null modelliert. Die Reparaturzeit ist die Zeit, für die der andere Sensor noch durchhalten muss, sei es bis ein sicherer Zustand erreicht ist (z. B. die Landung des Flugzeugs erfolgt ist), oder bis eine Reparatur bei laufendem Betrieb erfolgt ist. Sie ist hier mit 100 h angenommen.

Mit Formel (47) gilt

$$\overline{Q} \approx \lambda \cdot MRT = 1E-4/h \cdot 100 \text{ h} = 0,01$$

und somit für den Minimalschnitt

$$h_{MCS}(t) = \overline{h_{MCS}} = 1E-4/h \cdot 0,01 + 1E-4/h \cdot 0,01 = 2E-6/h$$

Da auch die Ausfallraten der anderen Minimalschnitte konstant sind, beträgt die Gesamt-Ausfallrate des Systems somit

$$h_{sys} = 1E-6/h + 1E-5/h + 2 \cdot 1E-5/h + 2E-6/h = 3,3E-5/h$$

Die im Beispiel verwendete Formel für die Eintrittsrate eines Minimalschnitts lässt sich auf Minimalschnitte beliebiger Ordnung $m = n_{Lit}$ erweitern:

$$\begin{aligned} h_{MCS}(t) &\lesssim h_1(t) \cdot Q_2(t) \cdot Q_3(t) \cdot \dots \cdot Q_m(t) \\ &\quad + h_2(t) \cdot Q_1(t) \cdot Q_3(t) \cdot \dots \cdot Q_m(t) \\ &\quad + \dots \\ &\quad + h_m(t) \cdot Q_1(t) \cdot Q_2(t) \cdot \dots \cdot Q_{m-1}(t) \\ &= \sum_{j=1}^m \left(h_j(t) \cdot \prod_{k=1, k \neq j}^m Q_k(t) \right) \end{aligned} \tag{65}$$

²⁰bezüglich der Exaktheit siehe Kommentar zu Formel (65)

Formel (65) ist nur für $Q_i \rightarrow 0$ korrekt. Die exakte Formel für zwei Ereignisse mit beliebig großen Nichtverfügbarkeiten wird in Beispiel 6.7 hergeleitet. Der Fehler fällt jedoch erst für große Nichtverfügbarkeiten ins Gewicht, ist also für korrekt ausgelegte Systeme (also wenn die Detektions- und Reparaturzeiten wesentlich kleiner sind als die MTTF) irrelevant. Die Formel ist zudem immer konservativ, so dass selbst bei nicht korrekt ausgelegten Systemen die Ausfallrate nicht zu klein geschätzt wird.

Für die System-Ausfallrate (oder allgemeiner: die Eintrittsrate des Top-Events) gilt

$$\begin{aligned} h_{\text{sys}}(t) &\lesssim h_{\text{MCS1}}(t) + h_{\text{MCS2}}(t) + \dots + h_{\text{MCSn}}(t) \\ &= \sum_{i=1}^{n_{\text{MCS}}} \left(\sum_{j=1}^{n_{\text{Lit}, \text{MCS}_i}} \left(h_j(t) \cdot \prod_{k=1, k \neq j}^{n_{\text{Lit}, \text{MCS}_i}} q_{i,k}(t) \right) \right) \end{aligned} \quad (66)$$

Diese Formel gilt exakt nur, wenn alle Minimalschnitte nur aus einem Ereignis bestehen. Andernfalls kann es sein, dass sich die Minimalschnitte gegenseitig überlappen, so dass das Ergebnis etwas zu groß wird. Dies lässt sich wie bei der Berechnung der System-Nichtverfügbarkeit durch Disjunktion der Minimalschnitte berücksichtigen. Diese Operation ist jedoch sehr aufwändig und stößt selbst bei Verwendung von BDDs bei großen Fehlerbäumen an die Leistungsgrenzen moderner PCs.

6.1.3 Einkanalig Fail-Safe

Im letzten Beispiel wurde angenommen, dass der Prozess nicht einfach abgeschaltet und in einen sicheren Zustand gebracht werden kann, wenn ein Fehler erkannt wird. Bei vielen Prozessen ist dies durchaus möglich, beispielsweise kann ein Zug immer noch sicher zum Stillstand gebracht werden, wenn die Weg- oder Geschwindigkeitsmessung ausfällt. Dabei muss nicht einmal bekannt sein, welche Komponente genau ausgefallen ist, sondern man kann den sicheren Zustand auch im Fall von Inkonsistenzen jeglicher Art anfordern. Dies soll im nächsten Beispiel verdeutlicht werden.

Gelegentlich spricht man von einer Fail-Safe-Architektur, wenn die Steuerung in der Lage ist, den Prozess im Fall erkannter Fehler in einen sicheren Zustand zu bringen. Der Begriff ist allerdings äußerst unscharf, denn nirgends ist definiert, welche Fehlermodi oder welcher Anteil der gefährlichen Fehlermodi erkannt werden müssen, um ein System „fehlersicher“ nennen zu dürfen.²¹

Beispiel 6.3 *In diesem Beispiel wird angenommen, dass die Steuerung in dem Fall, dass sie einen Fehler erkennt, den Aktor so ansteuert, dass der Prozess in einen sicheren Zustand geht (beispielsweise Abschalten der Energiezufuhr oder Anlegen der Bremsen).*

In diesem Fall sind die erkennbaren Ausfälle der Sensoren SENS_ED nicht mehr gefährlich und können somit ignoriert werden. Und auch der Zustand, dass ein Sensor falsche Werte liefert, jedoch nicht klar ist welcher, ist unkritisch, da im Fall einer Diskrepanz die Steuerung den Aktor ebenfalls so ansteuern wird, dass der Prozess einen sicheren Zustand erreicht.

²¹Eine vollständige Erkennung und Behandlung aller kritischen Fehler gilt gemeinhin als unmöglich.

Gefährlich ist nur der Fall, dass beide Sensoren gleichzeitig unerkennbar falsche Werte liefern, also die Ausfälle *SENS_NED.1* und *SENS_NED.2* gleichzeitig vorliegen, und die falschen Werte noch dazu so ähnlich sind, dass dies für die Steuerung nicht als Fehler erkennbar ist. Die beiden Ausfälle müssten also nicht nur im Detail ähnlich sein, sondern auch noch so schnell hintereinander passieren, dass dies für die Steuerung nicht sichtbar wäre, also typischerweise innerhalb der Prozessfehlertoleranzzeit (PFTZ). Man mag verleitet sein anzunehmen, dass so ein Fall praktisch nie eintreten wird. Tatsächlich wird dieser Fall wohl auch nicht durch unabhängige zufällige Ereignisse eintreten, die Erfahrung zeigt aber, dass es immer wieder zu gleichzeitigen Ausfällen redundanter Komponenten aufgrund nicht berücksichtigter äußerer Umstände kommt. Flug AF447 oder die Katastrophe von Fukushima sind bekannte Beispiele hierfür. Um die Realität möglichst korrekt zu modellieren, sollte man einen Faktor für den Anteil von gemeinsamen Ausfällen aufgrund nicht näher bekannter oder berücksichtigter äußerer Umstände annehmen. In [IEC 61508] wird dieser Common-Cause-Factor genannt und mit β bezeichnet. In [IEC 61508] sind auch Tabellen enthalten, die bei der Schätzung dieses Faktors hilfreich sein können. Abbildung 24 zeigt den so erstellten Fehlerbaum. Hier wurde $\beta = 2\%$ zwischen den Ereignissen *SENS_NED.1* und *SENS_NED.2* angenommen und diese daher in *SENS_NED_CC* umbenannt.

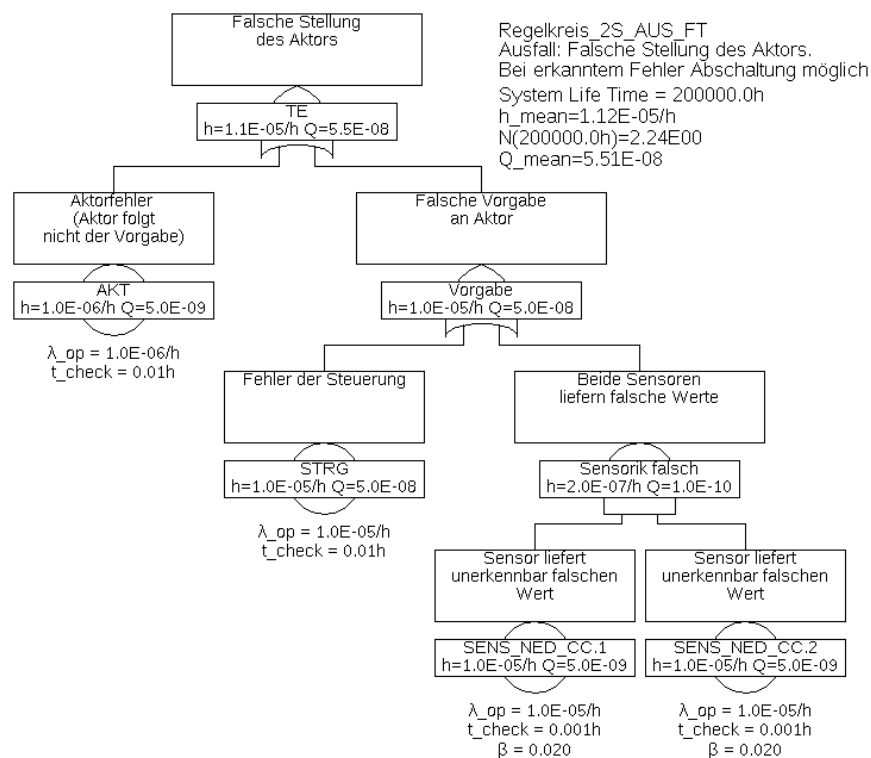


Abbildung 24: Fehlerbaum einer Sicherheitsfunktion mit kontinuierlicher Anforderung und redundanten Sensoren und leicht erreichbaren sicheren Zustand

Die Minimalschnitte und deren Teil-Eintrittsraten sind in Tabelle 1 gelistet.

Darin bezeichnet *SENS_NED_CC.COM* das Common-Cause-Ereignis, dass beide Sen-

Tabelle 1: Minimalschnitte für Beispiel 6.3

Minimalschnitt	Eintrittsrate $\bar{h}$
STRG	1,0E-05/h
AKT	1,0E-06/h
SENS_NED_CC.COM	2,0E-07/h
SENS_NED_CC.1 & SENS_NED_CC.2	9,604E-14/h

soren gleichzeitig (aufgrund eines gleichzeitig wirkenden äußeren Ereignisses) unerkennbar ausfallen. Seine Eintrittsrate beträgt $\beta \cdot \lambda_{\text{SENS_NED_CC}} = 0,02 \cdot 1,0\text{E}-5/\text{h} = 2,0\text{E}-7/\text{h}$.

Die Ausfallrate der Sensorik (Gatter „Sensorik falsch“) ändert sich damit von $2,2\text{E}-5/\text{h}$ auf $2,0\text{E}-7/\text{h}$. Die Sensorik hat in diesem Fall also keinen großen Anteil an der Gesamt-Ausfallrate mehr.

Allgemein gilt: Bei der Modellierung können Ausfälle, die unmittelbar zur sicheren Seite gehen, oder bei deren Eintritt eine immer zur Verfügung stehende Maßnahme mit höchster Wahrscheinlichkeit einen sicheren Zustand herbeiführt, weggelassen werden. Es muss jedoch unbedingt geprüft werden, ob diese Maßnahmen auch tatsächlich vorhanden und geeignet sind, in allen Fällen einen sicheren Zustand (des Prozesses!) zu erreichen. Um diese Prüfung durchführen zu können, müssen alle angenommenen Maßnahmen und sicheren Zustände zwingend erwähnt und insbesondere bei generischen Komponenten in die für den Kunden bestimmte Dokumentation des Produkts übernommen werden.

Bei Ereignissen unterhalb von UND-Gattern ist eine korrekte Modellierung der Nichtverfügbarkeit nötig. Diese basiert auf Fehlerdetektions- und Wiederherstellungszeiten. Zudem kann es notwendig sein, gleichzeitige Ausfälle redundanter Komponenten zu berücksichtigen, beispielsweise durch Common-Cause-Faktoren β .

Zur korrekten Modellierung von Diagnose und Inspektion ist die Kenntniss des physikalischen oder technischen Prozesses, in den die Sicherheitsfunktion eingebettet ist, notwendig. Ist diese (noch) nicht bekannt, müssen alle getroffenen Annahmen als Bedingungen bezüglich der Gültigkeit des Fehlerbaums dokumentiert und ggf. an Kunden weitergegeben werden.

6.1.4 Modellierung von regelmäßigen Tests und Diagnose

In Beispiel 6.3 wurde angenommen, dass sich der Prozess leicht in einen sicheren Zustand bringen lässt. Die Sicherheit ist maßgeblich durch die Ausfallrate der Steuerung bestimmt. Es liegt nun nahe, eine Diagnoseeinheit zu ergänzen, welche die Aktivität der Steuerung überwacht. Wenn die Diagnoseeinheit einen Fehler erkennt, bringt sie den Prozess (z. B. die Maschine oder die verfahrenstechnische Anlage) typischerweise über einen zusätzlichen, einfach aufgebauten binären Not-Aktor (beispielsweise ein Relais oder ein Abschaltventil) in einen sicheren Zustand (Stillstand, Leerlauf). Die Grundarchitektur solcher Steuerungen oder Regelungen ist in Abbildung 25 dargestellt. Sie wird oft als

einkanalig mit Diagnose, kurz 1oo1D (für engl. 1-out-of-1) bezeichnet.

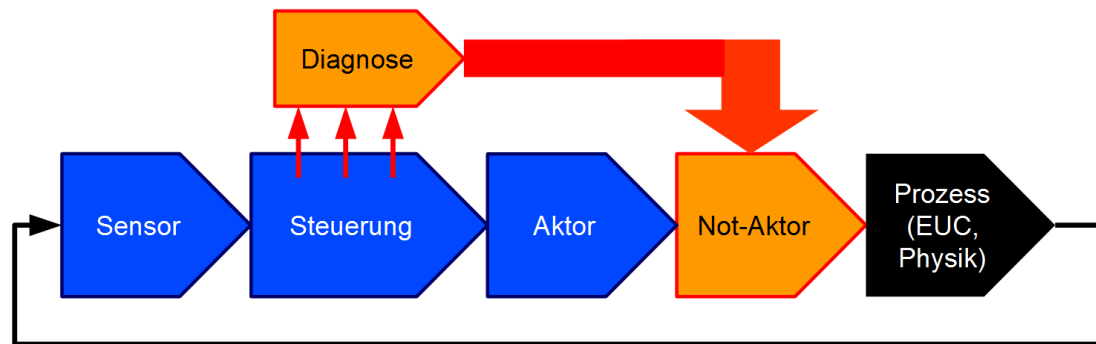


Abbildung 25: Regelkreis mit Diagnose und Notabschaltung

Die systematische Qualität der Diagnoseeinheit, also der Anteil der durch die Diagnose rechtzeitig erkennbaren Ausfälle an der Gesamtheit der (gefährlichen) Ausfallmodi des diagnostizierten Bauteils, wird als Diagnose-Deckungsgrad (engl. Diagnostic Coverage, DC) bezeichnet. Angenommen die Diagnoseeinheit überwacht die Versorgungsspannungen, den Prozessortakt und die regelmäßige Abarbeitung der kritischen Software-Tasks (Watchdog-Funktion), nicht jedoch die Logik der Recheneinheiten, die Speicher, oder die Ein-/Ausgabeeinheiten (A/D- und D/A-Wandler etc.) der Steuerung, so erhält man gemäß Tabellen in einschlägigen Normen vielleicht einen Diagnose-Deckungsgrad von 70%.

Es gibt nun mindestens drei Möglichkeiten, die Diagnose zu modellieren:

1. Man verringert die Ausfallrate der diagnostizierten Komponente (hier STRG) entsprechend des Diagnose-Deckungsgrads. Die Diagnose taucht also im Fehlerbaum nicht explizit auf. Der Vorteil ist offensichtlich ein einfacher Fehlerbaum. Nachteil ist, dass die Diagnose nicht ausdrücklich als sicherheitsrelevante Komponente erwähnt wird, man also stillschweigend voraussetzt, dass die Diagnose immer funktionieren wird. Tatsächlich aber unterliegt auch die Diagnose und insbesondere der Abschaltpfad zufälligen Ausfällen und muss daher in den meisten Anwendungen regelmäßig getestet werden.
2. Man teilt die Ausfälle der diagnostizierten Komponente gemäß Diagnose-Deckungsgrad auf zwei Basisereignisse auf, eines für die von der Diagnose nicht detektierbaren Ausfälle (AUSFALL_UNDET), eines für die detektierbaren (AUSFALL_DET).

Der Ausfall der Diagnose selbst wird ebenfalls als Basisereignis (DIAG) mit einer bestimmten Ausfallrate und Testintervall modelliert. Das Ereignis für die erkennbaren Ausfälle wird mit der Diagnose mittels UND-Gatter verbunden, vgl. Abbildung 26 links.

3. Wie zuvor, aber das Basisereignis des Diagnoseausfalls wird als Bedingung (engl. Condition) gekennzeichnet (DIAG_COND), und mittels INHIBIT-Gatter mit dem zu diagnostizierenden Ausfall (AUSFALL_DET) verbunden, vgl. Abbildung 26 rechts.

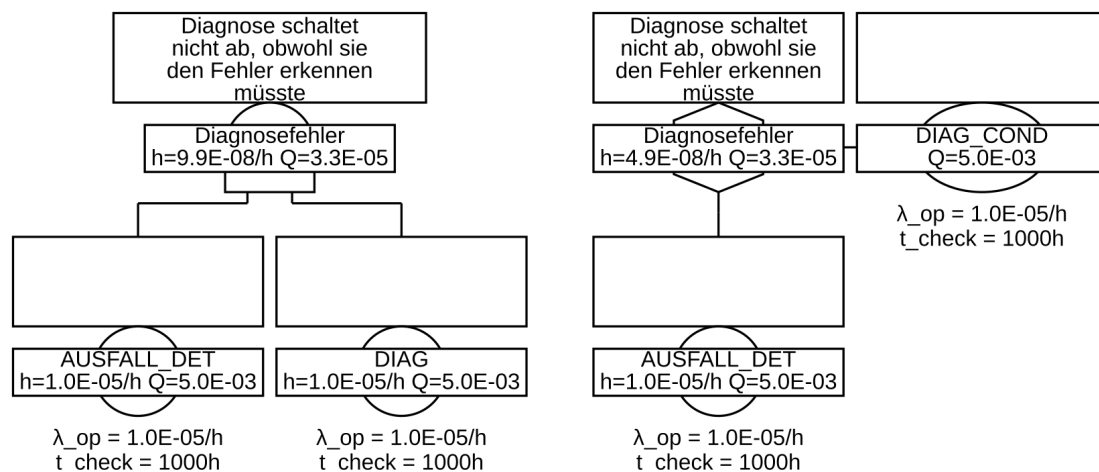


Abbildung 26: Und-Gatter und Inhibit-Gatter mit Bedingung

Der Unterschied der Varianten 2 und 3 ist folgender: Beim UND-Gatter gilt Formel (65). Bei der Verknüpfung von AUSFALL_DET und DIAG ergibt sich also

$$h_{\text{UND}} = h_{\text{AUSFALL_DET}} \cdot Q_{\text{DIAG}} + h_{\text{DIAG}} \cdot Q_{\text{AUSFALL_DET}}$$

Das spiegelt aber nicht die Realität wider, denn die Ausfallrate der Diagnose h_{DIAG} hat keinerlei Einfluss auf die Eintrittsraten des Systemausfalls, sondern nur deren Nichtverfügbarkeit Q_{DIAG} . Solange also Q_{DIAG} konstant ist, sollte auch die Ausfallrate der Verknüpfung gleich bleiben. Dies erreicht man, indem man ein Ereignis als Bedingung kennzeichnet, und mittels INHIBIT an die durch Eintrittsraten h und Nichtverfügbarkeiten Q beschriebenen „normalen“ Teile des Fehlerbaums anbindet. Damit wird die Eintrittsraten der Bedingung ignoriert (als wäre sie null), und somit wird der zweite Summand in vorheriger Formel null – genau das, was hier gewünscht ist:

$$h_{\text{INHIBIT}} = h_{\text{AUSFALL_DET}} \cdot Q_{\text{DIAG}} + 0 \cdot Q_{\text{AUSFALL_DET}}$$

Anmerkung: Die Unterscheidung von UND und INHIBIT mag in unterschiedlichen Werkzeugen unterschiedlich streng gehandhabt werden, da es auch hier keine Standardisierung gibt.

Anmerkung: Oft beschreibt man mit einem Bedingungs-Ereignis auch die Wahrscheinlichkeit, dass bestimmte Randbedingungen vorliegen. In dem Fall wird diese Wahrscheinlichkeit direkt als Konstante eingegeben, siehe Anhang A.3.

Anmerkung: Bedingungen können miteinander verknüpft werden (z. B. mittels UND-, oder ODER-Gattern), bevor sie als Gesamt-Bedingung an ein INHIBIT-Gatter angeschlossen werden.

Beispiel 6.4 Die Abbildungen 27 und 28 stellen den Fehlerbaum für eine Architektur dar, in der die Steuerung diagnostiziert wird und im Falle eines erkennbaren Ausfalls der Prozess abgeschaltet wird.

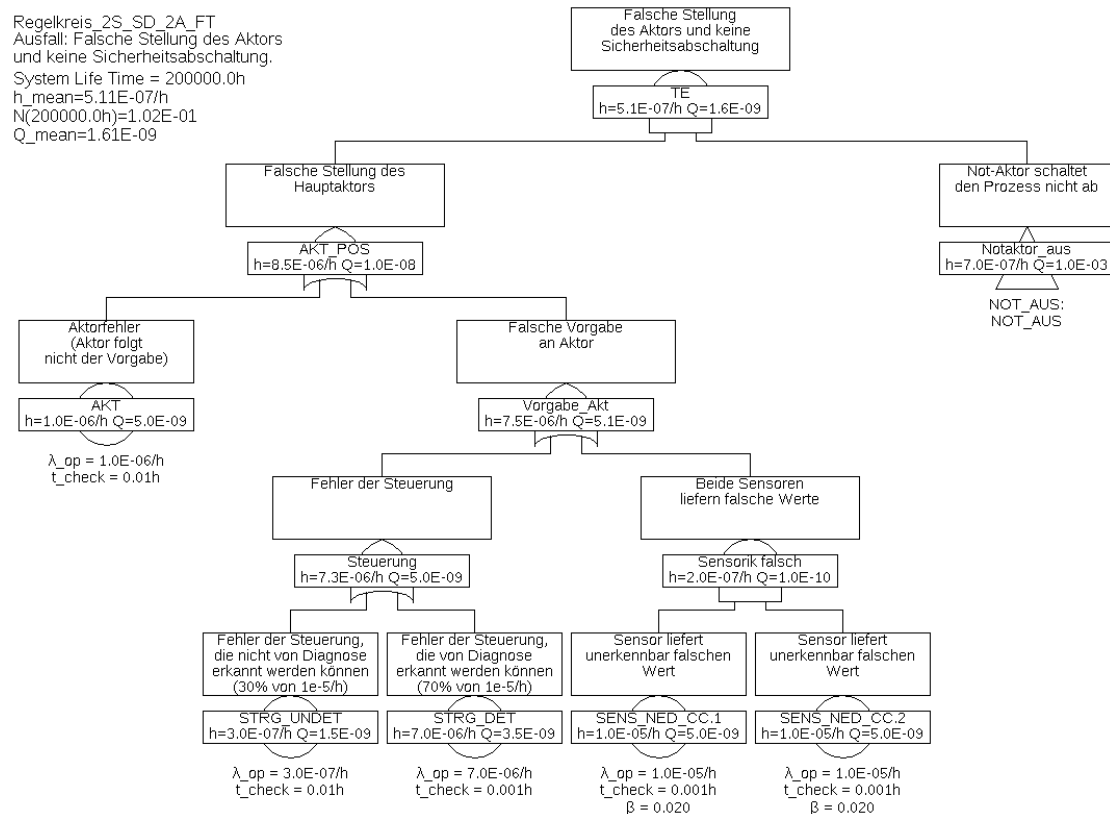


Abbildung 27: Regelkreis mit Diagnose und Notabschaltung

Der Fehlerbaum wurde aus Platzgründen in zwei Teil-Fehlerbäume aufgeteilt. Der Teil-Baum für das Gatter „Notaktor_ aus“ wurde mittels eines Transfer-Gatters mit dem übergeordneten Fehlerbaum verbunden. Das Transfer-Gatter selbst hat keinerlei Auswirkung auf die Berechnung, es ist nur ein Verweis auf einen Zweig an anderer Stelle.

Dabei wird angenommen, dass die Abschaltung über den Notaktor auch dann erfolgt, wenn die Steuerung erkennt, dass der Haupt-Aktor defekt ist. Dies ist in vielen Anwendungen leicht und rechtzeitig daran zu erkennen, dass die Regelgröße zunehmend vom Sollwert abweicht, oder die Steuerung liest die Stellgröße über einen zusätzlichen Sensor zurück. Einen Aktor-Fehler kann die Steuerung der Diagnose einfach dadurch melden, dass sie sich in einen von der Diagnose erkennbaren Fehlerzustand versetzt (z. B. einen Watchdog nicht mehr zurücksetzt).

Die Minimalschnitte sind in Tabelle 2 gelistet.

Sobald es Minimalschnitte mit mehr als einem Ereignis gibt, sind die Nichtverfügbarkeit der darin enthaltenen Ereignisse und somit über die Fehlerdetektions- und Reparaturzeiten wesentlich. Diese müssen daher korrekt gewählt und begründet werden, wie in Tabelle 3 dargestellt. Die Prozessfehlertoleranzzeit wurde mit 0,01 h angenommen.

Die Eintrittsrates gefährlicher Ausfälle des Systems beträgt nun also nur noch $5,1 \cdot 10^{-7} \text{h}$ und wird maßgeblich von den unerkannten Fehlern der Steuerung bestimmt. Weitere Ver-

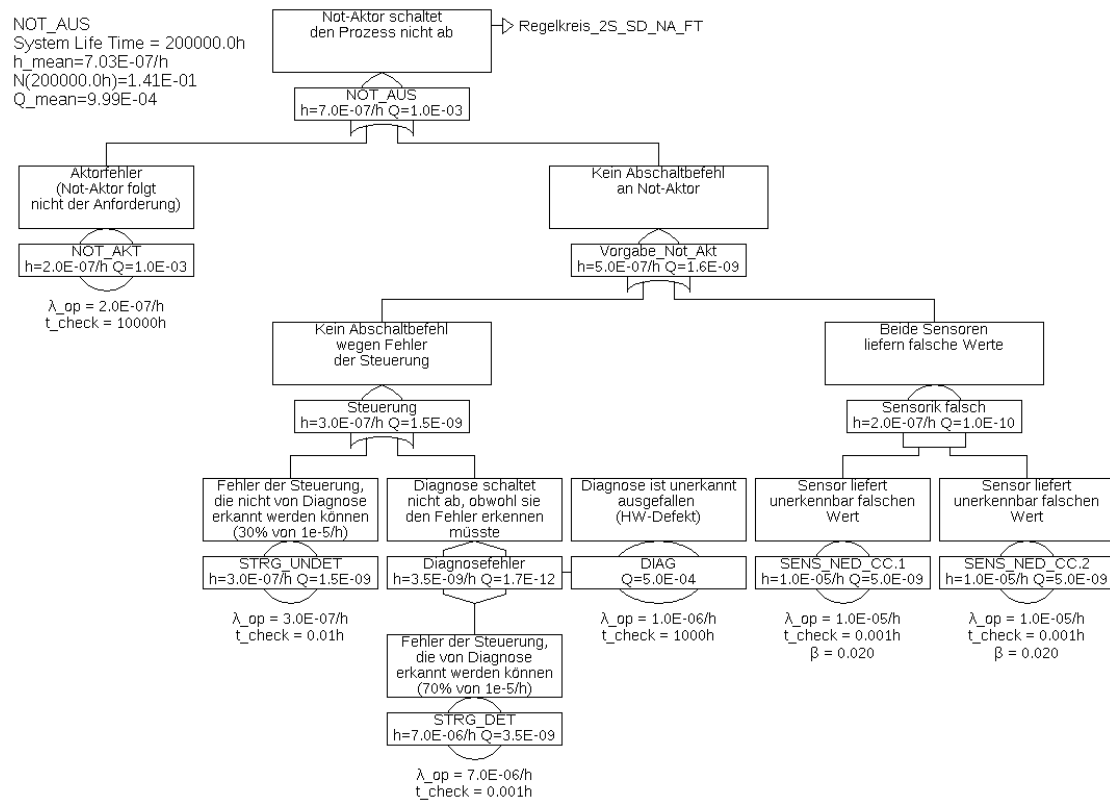


Abbildung 28: Regelkreis mit Diagnose und Notabschaltung, Teilbaum für Notabschaltung

Tabelle 2: Minimalschnitte für Beispiel 6.4

Minimalschnitt	Eintrittsrates $\bar{h}$
STRG_UNDET	3,0E-07/h
SENS_NED_CC.COM	2,0E-07/h
STRG_DET & NOT_AKT	6,988E-09/h
STRG_DET & DIAG	3,495E-09/h
AKT & NOT_AKT	9,983E-10/h
SENS_NED_CC.1 & SENS_NED_CC.2	9,604E-14/h

besserungen könnten durch Einsatz spezieller Sicherheitsprozessoren (beispielsweise Dual-Core Prozessoren im Lock-Step) und spezieller Speicher (ECC) erreicht werden.

6.1.5 Zweikanalig Fail-Safe

Das in Beispiel 6.4 betrachtete System wird gelegentlich als fehlersicheres einkanaliges System mit Diagnose (1oo1D) bezeichnet.

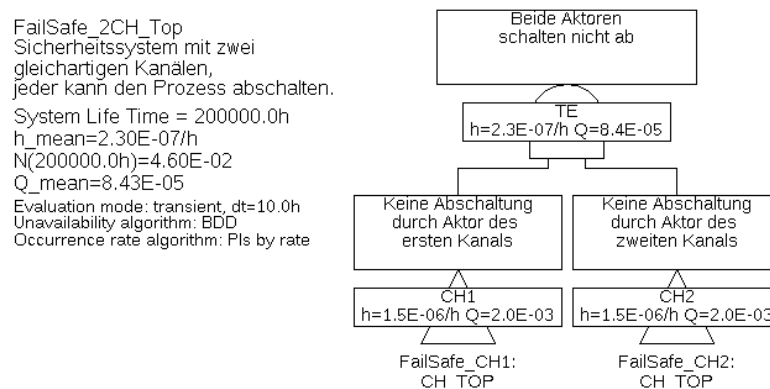
Steuerungen für hohe Sicherheitsanforderungen baut man oft aus zwei gleichartigen Einzelsteuerungen auf, wovon jede mit einer eigenen Diagnose versehen ist. Ist jede der Steuerungen in der Lage, den Prozess im Fall erkannter Fehler in einen sicheren Zustand zu bringen, wird dies manchmal als zweikanaliges fehlersicheres System, oder 1-aus-2

Tabelle 3: Basisereignisse für Beispiel 6.4

Ereignis	Fehleroffenbarung	Fehleroffenbarungszeit
STRG_UNDET	Systemausfall, Einzelfehler	unrelevant
STRG_DET	Diagnose	0,001 h
NOT_AKT	jährlicher Test	ca. 10000 h
DIAG	Einschalt-Selbsttest nach monatlicher Wartung/Reinigung	ca. 1000 h
AKT	Steuerung (unerwartetes Prozessverhalten)	0,01 h
SENS_NED_CC	a) Differenz der Sensoren	unmittelbar (0,001 h)
SENS_NED_CC	b) Gleichzeitiger Ausfall beider Sensoren: über Common-Cause β berücksichtigt, Einzelfehler	unrelevant

System mit Diagnose bezeichnet, kurz 1oo2D. Man sollte diese Bezeichnungen aber nur mit Vorsicht verwenden, da sie nicht harmonisiert sind, und in der Literatur durchaus unterschiedlich und sogar widersprüchlich verwendet werden.²²

Beispiel 6.5 Der Fehlerbaum eines zweikanaligen Steuerungssystems, bestehend aus der von den vorherigen Beispielen schon bekannten Sensorik, welche gemeinsam von zwei Steuerungen mit jeweils eigener Diagnose und eigenem Aktor verwendet wird, ist in Abbildung 29 mit dem unterlagerten Baum gemäß Abbildung 30 dargestellt. Da die Kanäle gleich aufgebaut sind, ist nur der Teil-Fehlerbaum des ersten Kanals dargestellt.

**Abbildung 29:** Homogen-redundantes Steuerungssystem mit Sicherheitsabschaltung im Fehlerfall

Jede Steuerung kann bei erkannten Fehlern den Prozess über ihren Aktor in einen sicheren Zustand versetzen, auch die Diagnoseeinheit dieses Kanals nutzt den Aktor hierfür.

²²In [IEC 61508] wird diese Architektur mit 1oo2 statt 1oo2D bezeichnet, da gemäß dieser Norm immer eine Diagnose vorhanden sein muss. In Teil 6 der Norm wird mit 1oo2D eine Art Fail-Operational Architektur bezeichnet, was sehr unüblich ist und daher oft missverstanden wird, zumal der Begleittext dies nicht klar erläutert.

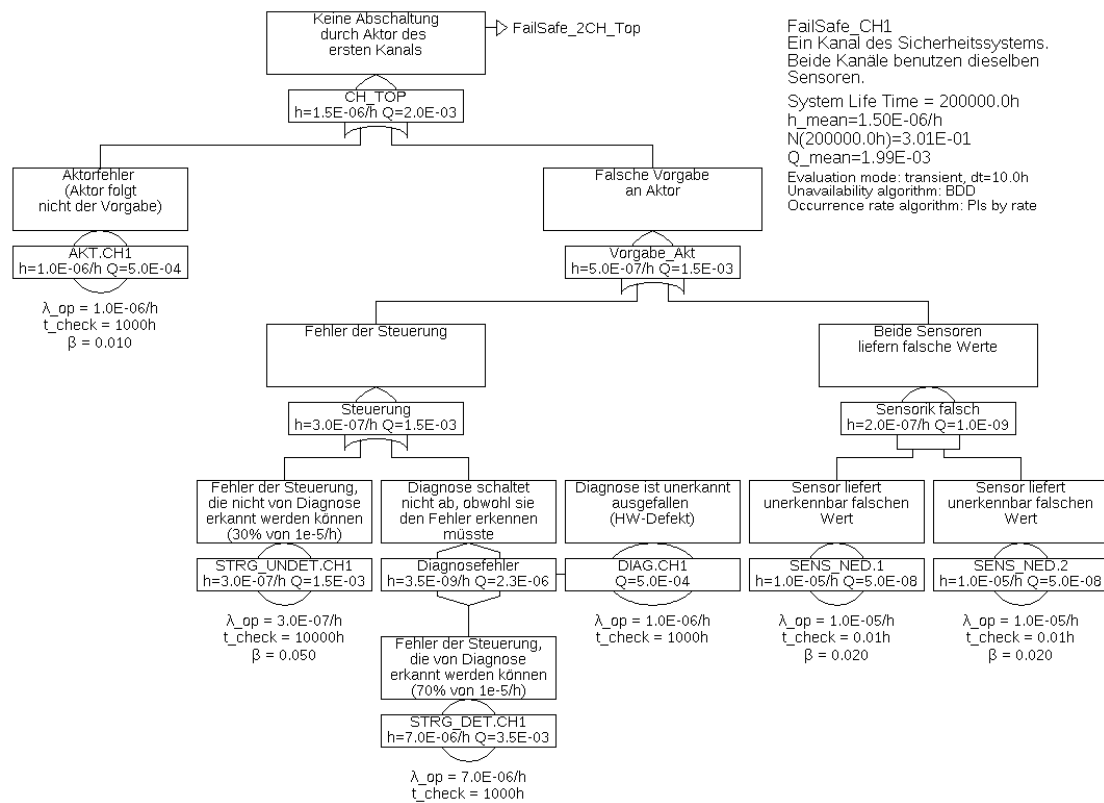


Abbildung 30: Homogen-redundantes Steuerungssystem mit Sicherheitsabschaltung im Fehlerfall

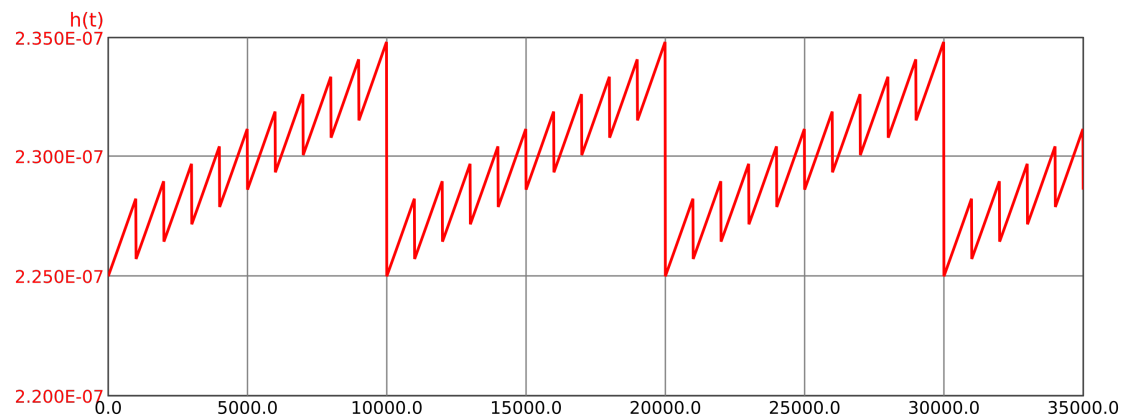
Zumindest für undetektierbare Ausfälle von Sensorik und Elektronik, sowie Ausfälle von Aktoren kann man Ausfälle aufgrund gemeinsamer Ursachen meist nicht ausschließen. Daher wurden hier nicht nur (wie bisher) die unerkennbaren Ausfälle der Sensoren, sondern auch die unerkennbaren Ausfälle der Steuerungen und die Ausfälle der Aktoren mit Common-Cause-Faktoren versehen.

Für das System können die in Tabelle 4 gelisteten Minimalschnitte ermittelt werden. Aus Tabelle 4 geht deutlich hervor, dass die Ausfälle aufgrund gemeinsamer Ursache die wesentlichen Ereignisse sind. Das ist nicht nur in diesem Beispiel so, sondern entspricht der praktischen Erfahrung. Aus diesem Grund muss bei mehrkanaligen Systemen immer eine Common-Cause-Analyse durchgeführt werden, und eine Vielzahl von Maßnahmen ergriffen werden, um die Rate von Ausfällen aufgrund gemeinsamer Ursache (mathematisch beschrieben durch den β -Faktor) möglichst klein zu halten.

Tabelle 4: Minimalschnitte für Beispiel 6.5

Minimalschnitt	Eintrittsrate $\bar{h}$
SENS_NED.COM	2,0E-07/h
STRG_UNDET.COM	1,5E-08/h
AKT.COM	1,0E-08/h
AKT.CH1 & STRG_UNDET.CH2	1,548E-09/h
AKT.CH2 & STRG_UNDET.CH1	1,548E-09/h
AKT.CH1 & AKT.CH2	9,699E-10/h
STRG_UNDET.CH1 & STRG_UNDET.CH2	8,106E-10/h
STRG_DET.CH2 & STRG_UNDET.CH1 & DIAG.CH2	5,745E-12/h
STRG_DET.CH1 & STRG_UNDET.CH2 & DIAG.CH1	5,745E-12/h
AKT.CH1 & STRG_DET.CH2 & DIAG.CH2	4,542E-12/h
AKT.CH2 & STRG_DET.CH1 & DIAG.CH1	4,542E-12/h
SENS_NED.1 & SENS_NED.2	9,604E-13/h
STRG_DET.CH1 & DIAG.CH1 & STRG_DET.CH2 & DIAG.CH2	2,393E-14/h

Der Verlauf der Ausfallrate über der Zeit ist in Abbildung 31 dargestellt. Auf den er-

**Abbildung 31:** Zeitlicher Verlauf der Ausfallrate eines homogen-redundanten Systems mit regelmäßigen Tests

sten Blick mag verwundern, dass die Ausfallrate des Systems nicht konstant ist, obwohl doch die Ausfallraten aller Komponenten konstant sind. Dies liegt an der zeitvarianten Nichtverfügbarkeit jedes Kanals, welche gemäß Formel (65) in die Ausfallrate eingeht. Aufgrund der hohen Common-Cause-Anteile (siehe Minimalschnitte in Tabelle 4), welche direkt – ohne Multiplikation mit einer Nichtverfügbarkeit – in die System-Ausfallrate eingehen, ist die Zeitabhängigkeit allerdings nur gering (man beachte die Skala für $h(t)$ in Abbildung 31).

6.1.6 Fail-Operational Systeme

Wie im einleitenden Beispiel 6.1 erwähnt, gibt es eine Vielzahl von Prozessen, die sich nicht einfach in einen sicheren Zustand bringen lassen.

Falls die Sicherheitsanforderungen nicht sehr hoch sind, genügt oft eine zweikanalige Architektur, wobei im Falle eines erkannten Fehlers jeder Kanal sich selbst abschalten kann oder sich gegenüber einer Auswahlhaltung als defekt erklären kann. Eine gute Eigen-diagnose jedes einzelnen Kanals ist hierfür unerlässlich, denn nur wenn klar ist, welcher Kanal defekt ist, kann auf den intakten Kanal umgeschaltet werden. Die Umschaltung selbst muss durch eine nachgeschaltete Auswahlhaltung erfolgen.

Für höhere Sicherheitsanforderungen ist der Diagnose-Deckungsgrad der Kanal-internen Diagnose der einzelnen Kanäle oft nicht ausreichend, es kommt also zu einer zu großen Rate von widersprüchlichen Ausgaben der einzelnen Kanäle, ohne dass ein Kanal anzeigt, dass er defekt ist. Die Auswahlhaltung hätte in diesem Fall eine 50% Chance, den richtigen auszuwählen. Die Wahrscheinlichkeit, den richtigen Kanal abzuschalten, lässt sich dadurch wesentlich verbessern, dass die Ausgaben von drei Kanälen verglichen werden. Unter der Bedingung, dass systematische Fehler und Ausfälle aufgrund gemeinsamer Ursache hinreichend selten sind, werden in der Regel mindestens zwei Kanäle die korrekte Ausgangsgröße ermitteln. Die Auswahlhaltung wird daher so aufgebaut, dass sie die Ergebnisse der beiden Kanäle verwendet, deren Ergebnisse identisch sind oder zumindest am nächsten beieinander liegen. Eine solche Architektur wird als 2-aus-3-Architektur (2oo3 oder auch 2oo3D, falls jeder Kanal eine eigene Diagnose hat) bezeichnet.²³

Unabhängig von der Anzahl der Kanäle darf die Auswahlhaltung nur eine minimale Ausfallrate haben, um nicht das für die Gesamt-Sicherheit maßgebliche Element zu sein. Manchmal erfolgt die Auswahl auch über die Mechanik, beispielsweise in dem jeder von drei Kanälen einen Aktor treibt, welcher von zwei anderen mechanisch überstimmt werden kann. Bei entsprechender „Intelligenz“ der Auswahlhaltung kann ein 2oo3-System sogar mit zwei ausgefallenen Kanälen noch korrekt arbeiten, wenn die Ausfälle von der Kanal-Diagnose erkannt und an die Auswahllogik gemeldet werden.

Beispiel 6.6 *Dieses Beispiel verwendet die Komponenten des Beispiels 6.1 wieder. Allerdings werden nun drei Sensoren und drei Steuerungen (Rechner) eingesetzt, und zwar so, dass jede der Steuerungen die Werte aller drei Sensoren bekommt. Beim Ausfall eines Sensors können also alle drei Steuerungen weiterarbeiten, im Gegensatz zu einer Architektur, bei der jede Steuerung nur auf einen Sensor Zugriff hätte. Zudem kann so jede Steuerung Fehler der Sensorik erkennen. Die von den drei Steuerungen berechneten Vorgaben für den einzigen Aktor werden von einer Auswahl-Logik verglichen und gemäß Mehrheitsentscheidung (2-von-3) ausgewählt.*

Der Fehlerbaum ist in Abbildung 32 dargestellt, mit den Unterbäumen in 33 und 34.

Die Gatter „Steuerung“ in Abbildung 33 und „Sensorik“ in Abbildung 34 sind Kombinations-Gatter, sie werden weiter unten erklärt.

Neben den Ausfallraten aller Komponenten sind die Nichtverfügbarkeiten der Steuerungen und der Sensoren relevant. Hier wurde angenommen, dass sich alle Ausfälle der Steuerungen sowie alle unabhängigen Ausfälle der Sensoren sofort durch Diskrepanzen

²³Im Gegensatz zu 1oo2, 2oo2, 1oo3, 3oo3 ist bei 2oo3 eine Verwechslung nicht möglich, da jede Sichtweise dasselbe Ergebnis liefert.

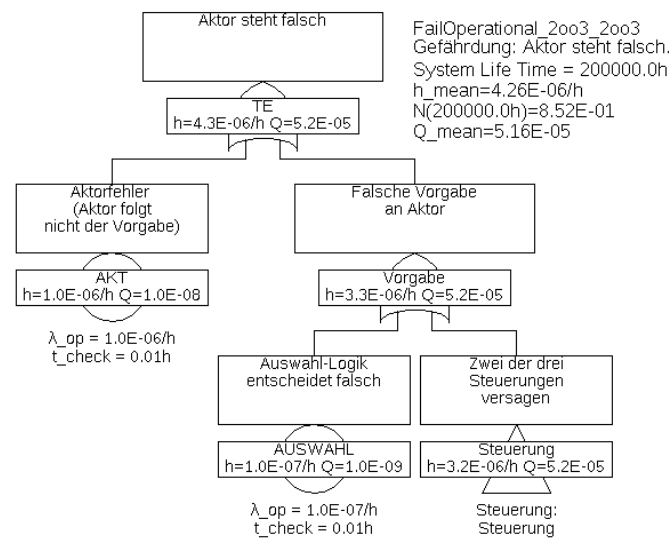


Abbildung 32: Homogen-redundantes Steuerungssystem mit 3 Kanälen ohne Sicherheitsabschaltung

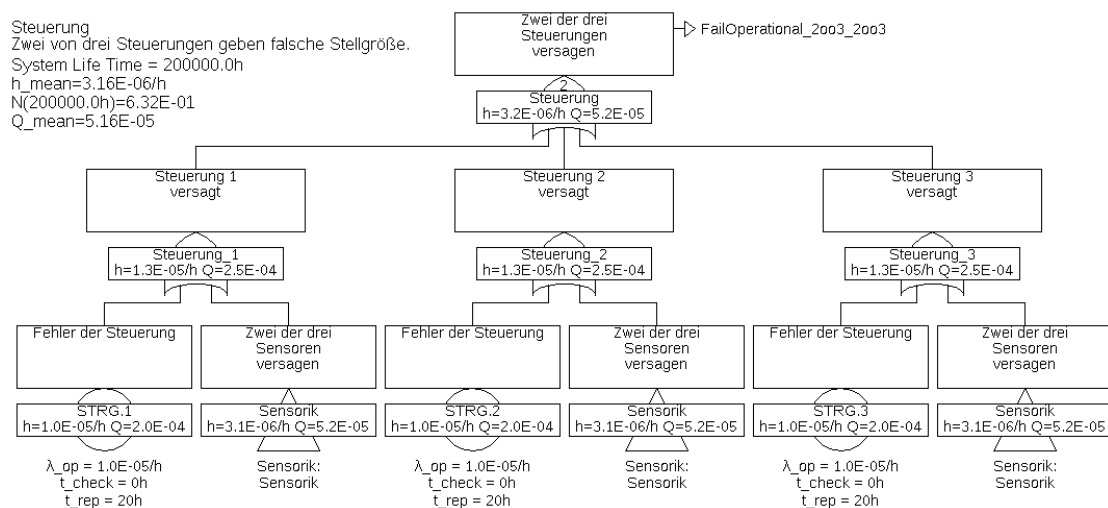


Abbildung 33: Redundante Steuerungen (2003) mit Nutzung derselben Sensorik

in der Auswahl-Einheit offenbaren, die Detektionszeit t_{check} also null ist. Die Zeit bis zur Reparatur, also die Zeit, die die verbliebenen Kanäle durchhalten müssen, wurde mit 20 h angenommen (man denke etwa an ein Langstreckenflugzeug, welches erst nach der Landung repariert werden kann, oder an ein Schiff, bei dem die Reparatur eines Aggregats zwar auf See erfolgt, aber eine Weile dauert).

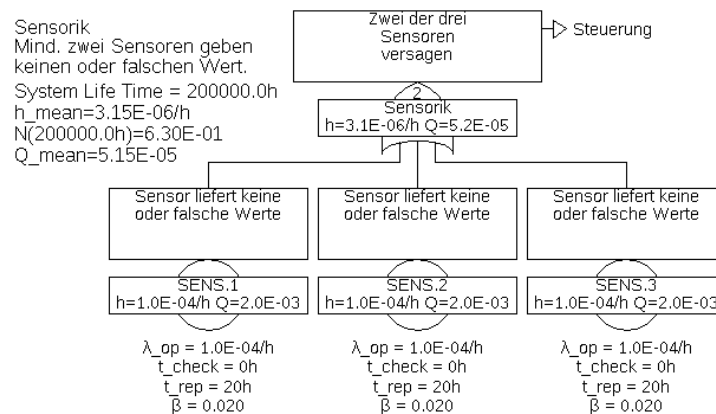


Abbildung 34: Redundante Sensorik (2oo3) zur gemeinsamen Nutzung

Die Minimalschnitte sind in Tabelle 5 gelistet.

Tabelle 5: Minimalschnitte für Beispiel 6.6

Minimalschnitt	Eintrittsrate $\bar{h}$
SENS.COM	2,0E-06/h
AKT	1,0E-06/h
SENS.1 & SENS.2	3,842E-07/h
SENS.1 & SENS.3	3,842E-07/h
SENS.2 & SENS.3	3,842E-07/h
AUSWAHL	1,0E-07/h
STRG.1 & STRG.2	4,0E-09/h
STRG.1 & STRG.3	4,0E-09/h
STRG.2 & STRG.3	4,0E-09/h

Die Gesamt-Ausfallrate ist im Vergleich zum Beispiel 6.1 von $1,11E-4/h$ auf $4,26E-6/h$ gesunken, und wird nun hauptsächlich durch die Ausfälle der Sensoren aufgrund gemeinsamer Ursache (β wurde mit 2% angenommen) bestimmt.

Im vorherigen Beispiel wurden sogenannte KOMBINATIONEN-Gatter (engl. Combination-Gate), auch Mehrheitsentscheider (engl. Voting-Gate) genannt, für die Gatter „Steuerung“ und „Sensorik“ verwendet. Sie sind nur eine Abkürzung für die entsprechende Kombination aus UND- und ODER-Gattern. Die Zahl im Gatter (oft mit m abgekürzt, hier also $m = 2$) gibt an, wieviele der n Eingänge mindestens versagen müssen (also erfüllt sein müssen), damit das durch das Gatter beschriebene Ereignis eintritt. $m = 1$ bedeutet also ein reines ODER-Gatter, $m = n$ bedeutet ein reines UND-Gatter, $1 < m < n$ bedeutet eine Kombination eines ODER- mit mehreren UND-Gattern, wie in Abbildung 35 beispielhaft für ein 2-von-3-Gatter gezeigt.

Zur Berechnung werden die Kombinations-Gatter in die entsprechende Kombination von UND- und ODER-Gattern umgewandelt. Es gibt daher keine besonderen Formeln oder Berechnungsmethoden für diese Gatter.

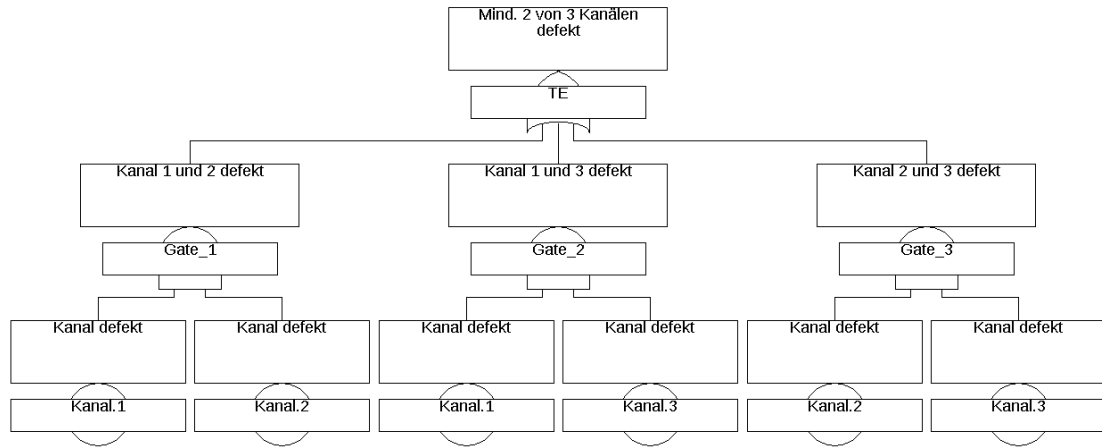


Abbildung 35: Entsprechung eines 2-aus-3-Gatters mit diskreten ODER- und UND-Gattern

6.1.7 Transiente und stationäre Berechnung

Wie die mittlere System-Nichtverfügbarkeit kann auch die mittlere System-Ausfallrate sowohl über eine stationäre Berechnung als auch eine transiente Berechnung ermittelt werden, siehe hierzu Abschnitt 5.1.4.

Einziger Unterschied ist, dass der in Abschnitt 5.1.4 beschriebene mathematische Fehler beim Rechnen mit Mittelwerten für die Ausfallrate erst bei Minimalschnitten dritter Ordnung (also Minimalschnitten mit drei Basis-Ereignissen) auftritt, und nicht schon bei Minimalschnitten zweiter Ordnung wie bei der Nichtverfügbarkeit. Das liegt daran, dass gemäß Formel (65) bei einem Minimalschnitt zweiter Ordnung noch keine Multiplikation von Nichtverfügbarkeiten auftritt. Wenn also klar ist, dass sich das System (abgesehen von einer im Vergleich zur Einsatzzeit vernachlässigbaren Einschwingphase) in einem quasi-stationären Zustand befinden wird, kann die System-Ausfallrate meist in guter Näherung mit Mittelwerten berechnet werden.

6.2 Berechnung mit Markov-Modellen

Bezüglich der Modellierung gibt es keine Unterschiede zu Kapitel 5.2.

Wie für die Berechnung der Nichtverfügbarkeit muss zunächst entweder der stationäre Zustand berechnet werden, oder das lineare Differenzialgleichungssystem muss über die Lebenszeit integriert werden.

Die Eintrittshäufigkeit $w_i(t)$ eines Zustands i ist die Summe der m Transitionsraten $h_{i,j}$, die zu diesem Zustand führen, jeweils multipliziert mit der Aufenthaltswahrscheinlichkeit im Ursprungszustand der jeweiligen Kante $p_{\text{ursprung}_{i,j}}$:

$$w_i(t) = \sum_{j=1}^m h_{\text{ein}_{i,j}}(t) \cdot p_{\text{ursprung}_{i,j}}(t) \quad (67)$$

Die System-Ausfallhäufigkeit ergibt sich aus der Summe der Zustands-Eintrittshäufigkeiten

für alle n Ausfallzustände

$$w_{\text{sys}}(t) = \sum_{i=1}^n w_i(t) = \sum_{i=1}^n \sum_{j=1}^m h_{\text{ein}_{i,j}}(t) \cdot p_{\text{ursprung}_{i,j}}(t) \quad (68)$$

Die Ausfallhäufigkeit $w(t)$ ist nur dann identisch zur gesuchten Ausfallrate $h(t)$, wenn die Aufenthaltswahrscheinlichkeit in allen Ausfallzuständen null ist, da in der Praxis im Fall einer (quasi-)kontinuierlich benötigten Sicherheitsfunktion ein Systemausfall praktisch sofort erkannt wird (nämlich durch das Eintreten eines Unfalls), was zur unmittelbaren Beendigung des Betriebs führt. Der Betrieb wird erst nach der Reparatur oder dem Ersatz des Systems durch ein anderes oder neues wieder aufgenommen, die Zeit bis dahin darf zur Berechnung der Ausfallrate nicht berücksichtigt werden (sonst würde sie zu optimistisch, wie man sich leicht mit Extrem-Beispielen vor Augen führen kann).

Möchte man also ein Markov-Modell zur Berechnung der Ausfallrate $h_{\text{sys}}(t)$ bzw. $\overline{h_{\text{sys}}}$ einsetzen, muss man entweder von allen Ausfallzuständen eine Transition mit einer sehr hohen Rate μ zurück in einen Nicht-Ausfallzustand modellieren (in der Regel zurück in den Ausgangszustand), oder man muss die Ausfallhäufigkeit $w(t)$ durch die Wahrscheinlichkeit, nicht in einem Ausfallzustand zu sein, dividieren:

$$h_{\text{sys}}(t) = \frac{w_{\text{sys}}(t)}{1 - \sum_{i=1}^n p_i(t)} \quad (69)$$

Falls die Wahrscheinlichkeiten der Ausfallzustände null sind, ergibt sich $h(t) = w(t)$. Formel (69) kann man und sollte man sicherheitshalber immer anwenden, auch wenn – wie zuvor erwähnt – bereits Transitionen mit großer Wiederherstellungsrate μ im Modell eingefügt wurden.

Beispiel 6.7 Die genannten Formeln sollen nun auf ein einfaches Markov-Modell angewandt werden. Hierfür wird ein einfaches diversitär-zweikanaliges System angenommen, bestehend aus den unterschiedlich aufgebauten (diversitären) Kanälen A und B . Die Kanäle A und B haben daher unterschiedliche Ausfallraten λ_A und λ_B . Beide Kanäle A und B werden regelmäßig getestet, jedoch in unterschiedlichen Intervallen T_A und T_B . Daraus ergeben sich unterschiedliche $\mu_A = 2/T_A$ und $\mu_B = 2/T_B$.

Wenn beide Kanäle versagen, versagt die kontinuierlich benötigte Sicherheitsfunktion, beendet also ihren Betrieb durch einen Unfall. Die Reparatur oder das Ersetzen des Systems nach einem Unfall führt zurück in den Ursprungszustand OK. Die Aufenthaltswahrscheinlichkeit im Zustand $A \& B$ während des Betriebs ist null, was durch ein im Vergleich zu den Ausfallraten sehr großes μ_{rep} modelliert wird.

Das zugehörige Markov-Modell ist in Abbildung 36 gezeigt.

Die Transitionsmatrix lautet:

$$T = \begin{pmatrix} -\lambda_A - \lambda_B & \mu_A & \mu_B & \mu_{\text{rep}} \\ \lambda_A & -\mu_A - \lambda_B & 0 & 0 \\ \lambda_B & 0 & -\mu_B - \lambda_A & 0 \\ 0 & \lambda_B & \lambda_A & -\mu_{\text{rep}} \end{pmatrix}$$

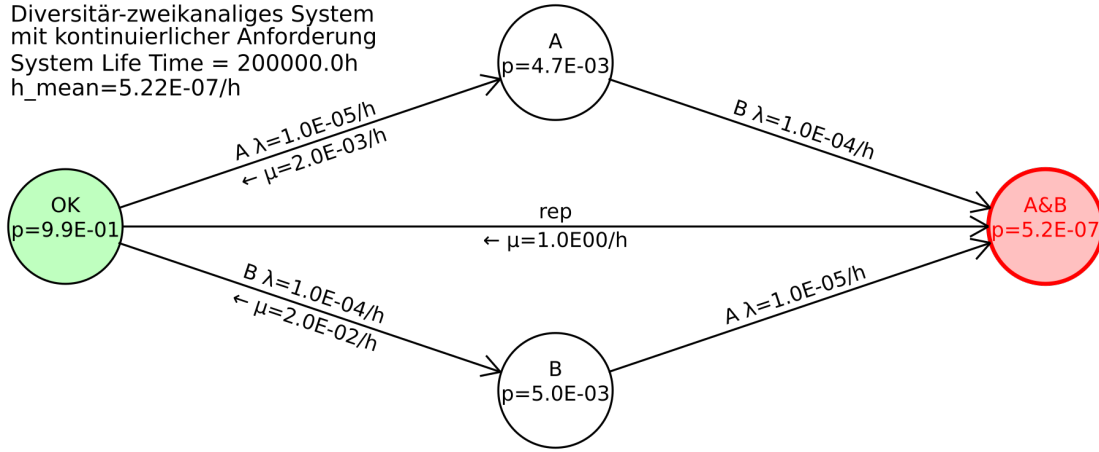


Abbildung 36: Diversität-redundantes System für kontinuierliche Anforderung

Für die stationäre Berechnung wird eine der Zeilen zu 1 gesetzt (hier wurde die letzte genommen):

$$\begin{pmatrix} -\lambda_B - \lambda_A & \mu_A & \mu_B & \mu_{\text{rep}} \\ \lambda_A & -\mu_A - \lambda_B & 0 & 0 \\ \lambda_B & 0 & -\mu_B - \lambda_A & 0 \\ 1 & 1 & 1 & 1 \end{pmatrix} \cdot \vec{p}(t) = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}$$

Die System-Eintrittsrates ergibt sich gemäß Formeln (68) und (69) zu

$$h_{\text{sys}} = \frac{p_A \cdot \lambda_B + p_B \cdot \lambda_A}{1 - p_{A\&B}}$$

Die in dieser Formel benötigten Zustandswahrscheinlichkeiten p_A , p_B und $p_{A\&B}$ werden durch Lösen des Gleichungssystems ermittelt. Damit ergibt sich h_{sys} zu

$$\begin{aligned} h_{\text{sys}} &= \frac{\lambda_A \lambda_B^2 + \lambda_A^2 \lambda_B + \lambda_A \lambda_B \mu_A + \lambda_A \lambda_B \mu_B}{\lambda_A^2 + \lambda_B^2 + \lambda_A \lambda_B + (\lambda_A + \lambda_B) \mu_A + (\lambda_A + \lambda_B) \mu_B + \mu_A \mu_B} \\ &= \frac{\lambda_A \lambda_B (\lambda_A + \lambda_B + \mu_A + \mu_B)}{\lambda_A (\lambda_A + \mu_A + \mu_B) + \lambda_B (\lambda_B + \mu_A + \mu_B) + \lambda_A \lambda_B + \mu_A \mu_B} \end{aligned}$$

Es sei bemerkt, dass die Rate μ_{rep} im Ergebnis nicht mehr auftaucht, da sie durch Formel (69) herausgekürzt wird. Dies entspricht genau der Erwartung, dass der Zahlenwert dieser Rate irrelevant sein muss. Diese Formel gilt nun für beliebige Testintervalle, im Gegensatz zu Formel (65), welche nur für ausreichend kleine Testintervalle gilt (andernfalls wird Formel (65) etwas zu groß).

Für geeignete Testintervalle $T_{\text{Test},A} \ll 1/\lambda_A$ und $T_{\text{Test},B} \ll 1/\lambda_B$ kann man λ_A gegenüber μ_A und λ_B gegenüber μ_B vernachlässigen und erst recht $\lambda_A \lambda_B$ gegenüber $\mu_A \mu_B$. Damit

gilt näherungsweise:

$$\begin{aligned} h_{\text{sys}} &\lesssim \frac{\lambda_A \lambda_B (\mu_A + \mu_B)}{\lambda_A (\mu_A + \mu_B) + \lambda_B (\mu_A + \mu_B) + \mu_A \mu_B} \\ &= \frac{\lambda_A \lambda_B \left(\frac{1}{\mu_A} + \frac{1}{\mu_B} \right)}{\lambda_A \left(\frac{1}{\mu_A} + \frac{1}{\mu_B} \right) + \lambda_B \left(\frac{1}{\mu_A} + \frac{1}{\mu_B} \right) + 1} \end{aligned}$$

Unter derselben Bedingung (geeignete Testintervalle) gilt auch $\lambda_A/\mu_A \ll 1$ und $\lambda_B/\mu_B \ll 1$ und folglich die Näherung

$$h_{\text{sys}} \lesssim \frac{\lambda_A \lambda_B \left(\frac{1}{\mu_A} + \frac{1}{\mu_B} \right)}{\frac{\lambda_A}{\mu_B} + \frac{\lambda_B}{\mu_A} + 1}$$

Ersetzt man nun die Reparaturraten durch die Testintervalle durch $\mu_i \approx 2/T_{\text{Test},i}$, so erhält man

$$h_{\text{sys}} \approx \frac{\lambda_A \lambda_B (T_{\text{Test},A} + T_{\text{Test},B})}{\lambda_A T_{\text{Test},B} + \lambda_B T_{\text{Test},A} + 2}$$

Unter der zusätzlichen Bedingung, dass die Testintervalle auch für die Ausfallrate des jeweils anderen Kanals ausreichend kurz wären, also $\lambda_A T_{\text{Test},B} \ll 1$ und $\lambda_B T_{\text{Test},A} \ll 1$ können die Produkte im Nenner vernachlässigt werden:

$$h_{\text{sys}} \approx \lambda_A \lambda_B \frac{T_{\text{Test},A} + T_{\text{Test},B}}{2}$$

Dies ist die bereits für die Berechnung der Ausfallrate eines Minimalschnitts bekannte Formel (65), angewandt auf den einzigen Minimalschnitt dieses Markov-Modells $\{A \& B\}$:

$$h_{\text{sys}} = h_{\text{MCS}} \lesssim \lambda_A Q_B + \lambda_B Q_A \lesssim \lambda_A \lambda_B \frac{T_{\text{Test},B}}{2} + \lambda_B \lambda_A \frac{T_{\text{Test},A}}{2} = \lambda_A \lambda_B \frac{T_{\text{Test},A} + T_{\text{Test},B}}{2}$$

6.3 Erwartungswert der Ausfälle

Für reparierbare bzw. ersetzbare Systeme ist neben der Ausfallrate noch der Erwartungswert der Ausfälle $N(t)$ über die Lebenszeit T interessant, welcher sich unmittelbar aus Formel (62) ergibt:

$$N_{\text{sys}}(T) = \overline{h_{\text{sys}}}(T) \cdot T = \text{PFH} \cdot T \quad (70)$$

7 Unzuverlässigkeit von komplexen Funktionen

Die Unzuverlässigkeit ist die wesentliche Größe für Sicherheitsfunktionen, die über eine ganz bestimmte Zeit hinweg funktionieren sollen, und die während dieser Zeit nicht oder nur in begrenztem Umfang gewartet und instandgesetzt werden können.

Nur wenige Systeme und deren Sicherheitsfunktionen fallen in diese Kategorie, etwa bemannte Raumfahrtmissionen. Hier ist nicht die Häufigkeit von Ausfällen, also Unfälle pro Stunde oder Unfälle pro Kilometer gefragt, sondern die Wahrscheinlichkeit, dass die Mission erfolgreich ist, also alle RaumfahrerInnen wieder gesund zur Erde zurückkehren (daher auch der gelegentlich verwendete Begriff „Mission Time“ anstelle von System-Lebenszeit oder Einsatzzeit). Die Mission ist fest vorgegeben, ein einfaches Erreichen eines sicheren Zustands ist nicht möglich (auch wenn es für gewisse Notfälle Abbruch-Szenarien geben mag). Abnutzungsbedingte Ausfälle während der Mission sind nicht anzunehmen. Alle Komponenten werden vor jeder Mission intensiv überprüft, so dass sie wie neu sind. Ab dem Start können die meisten sicherheitsrelevanten Komponenten nicht repariert werden. Das heißt nicht, dass es keine Diagnose geben kann, diese dient dann allerdings nur dazu, eine Komponente abzuschalten, um weiteren Schaden zu verhindern (sofern dies möglich ist) oder um auf eine Ersatz-Komponente umzuschalten (sofern vorhanden).

Da nur wenige Sicherheitsingenieure jemals die Unzuverlässigkeit solcher Systeme berechnen müssen, ist dieses Kapitel recht kurz gehalten.

Selbstverständlich ist die Unzuverlässigkeit nicht nur im Rahmen der Sicherheit relevant (dort wie gesagt eher selten). Auch die Frage „Wie wahrscheinlich ist, dass ein neues Auto innerhalb der ersten 5 Jahre unplanmäßig in die Werkstatt muss?“ ist eine Frage nach der Unzuverlässigkeit. Sie lässt sich allerdings leicht ohne Fehlerbäume oder Markov-Modelle beantworten, da es üblicherweise keine Redundanzen gibt, und daher direkt die Formeln (24) und (8) angewandt werden können, also

$$F(T) = 1 - e^{-\int_0^T \sum_{i=1}^n h_i(t) dt} = 1 - e^{-\sum_{i=1}^n \int_0^T h_i(t) dt} = 1 - \prod_{i=1}^n e^{-\int_0^T h_i(t) dt} \quad (71)$$

Diese Formel kann selbst im Falle komplizierter zeitabhängiger Ausfallraten $h_i(t)$ leicht mithilfe einer Tabellenkalkulation berechnet werden.

7.1 Berechnung mit Fehlerbäumen

Fehlerbäume können gut zur Berechnung der Unzuverlässigkeit komplexer Systeme verwendet werden. Dabei gibt es zwei Möglichkeiten der Berechnung:

1. Direkt über die Unzuverlässigkeiten $F(T)$ der Basis-Ereignisse, mit denselben mathematischen Methoden für die Nichtverfügbarkeit in Abschnitt 5 gezeigt. Bei Berechnung über Minimalschnitte wird also zunächst die Unzuverlässigkeit jedes Minimalschnitts berechnet gemäß

$$F_{\text{MCS}}(T) = \prod_{i=1}^m F_i(T) \quad (72)$$

und anschließend die System-Unzuverlässigkeit mittels der Formel von Esary-Proschan:

$$F_{\text{sys}}(T) \lesssim 1 - \prod_{j=1}^n (1 - F_{\text{MCS},i}(T)) \quad (73)$$

Noch schneller und genauer ist die Verwendung von BDDs.

Diese Methode ist sehr schnell, funktioniert aber nur, wenn keine Verknüpfungen mit Nichtverfügbarkeiten oder sonstigen Wahrscheinlichkeiten, die keine Unzuverlässigkeiten sind (wie etwa die Wahrscheinlichkeit, dass eine externe Randbedingung erfüllt ist), stattfinden. Insbesondere darf der Fehlerbaum also keine Bedingungen oder INHIBIT-Gatter enthalten. Sobald ein Basis-Ereignis eine Nichtverfügbarkeit oder sonstige Bedingung beschreibt, wird diese Methode zu optimistische Ergebnisse liefern ²⁴.

2. Über die im vorherigen Abschnitt 6 gezeigte Berechnung der Ausfallrate $h(t)$ (transiente Berechnung) und Anwenden der Formel (8).

$$F_{\text{sys}}(T) = 1 - e^{-\int_0^T h_{\text{sys}}(t) dt}$$

Die Methode ist wesentlich langsamer, da die Berechnung der Ausfallrate $h(t)$ schon aufwändiger ist, und diese auch noch für viele Zeitpunkte erfolgen muss. Sie liefert jedoch auch dann korrekte Ergebnisse, wenn der Fehlerbaum Basis-Ereignisse enthält, die Nichtverfügbarkeiten oder sonstige Bedingungen beschreiben. Dies sollte zwar bei Sicherheitsfunktionen, für die wirklich die Unzuverlässigkeit relevant ist, die Ausnahme sein, kann aber gelegentlich vorkommen.

Beispiel 7.1 Für den in Abbildung 37 gezeigten Fehlerbaum soll die Unzuverlässigkeit zum Zeitpunkt $T = 200\,000\text{ h}$ mit beiden vorgestellten Methoden berechnet werden. Die Unzuverlässigkeit der Basis-Ereignisse A und B wird hier durch Weibull-Verteilungen beschrieben, wobei für A eine zunehmende Ausfallrate (Weibull-Exponent > 1) und für B eine abnehmende Ausfallrate zutrefte (Weibull-Exponent < 1).

Methode 1:

Die Unzuverlässigkeit des Basis-Ereignisses A ergibt sich zu:

$$F_A(T) = 1 - e^{-(\lambda \cdot T)^k} = 1 - e^{-(6,0\text{E}-6/\text{h} \cdot 200\,000\text{ h})^{4,0}} \approx 0,874$$

Für Basis-Ereignis B ergibt sich

$$F_B(T) = 1 - e^{-(4,0\text{E}-6/\text{h} \cdot 200\,000\text{ h})^{0,4}} \approx 0,599$$

und für Basis-Ereignis C gilt

$$F_C(T) = 1 - e^{-1,0\text{E}-5/\text{h} \cdot 200\,000\text{ h}} \approx 0,865$$

²⁴Ein gutes Tool wird den Benutzer entsprechend warnen oder automatisch auf einen anderen Algorithmus wechseln.

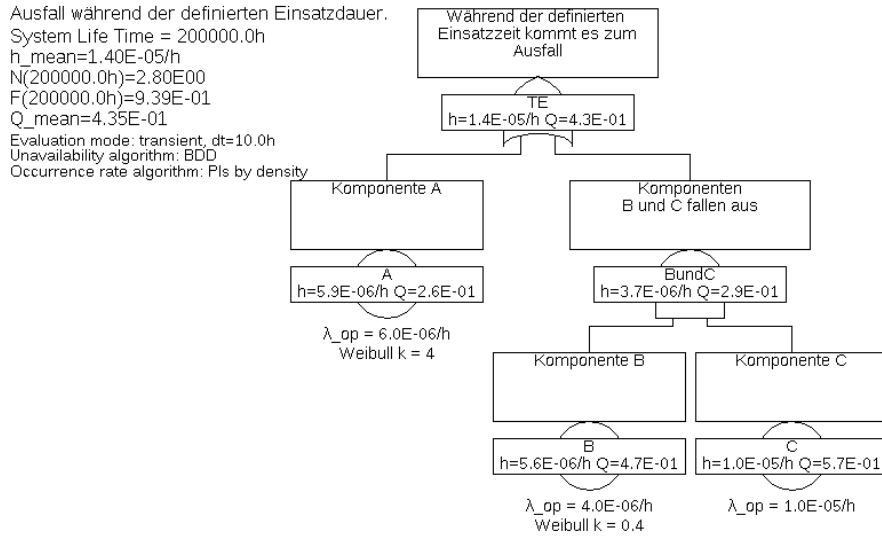


Abbildung 37: Berechnung der Unzuverlässigkeit mittels Fehlerbaum

Die Minimalschnitte sind $\{A\}$ und $\{B \cup C\}$, für die System-Unzuverlässigkeit ergibt sich gemäß Formel (73)

$$F_{\text{sys}}(T) \approx 1 - (1 - F_A(T)) \cdot (1 - F_B(T) \cdot F_C(T)) \approx 1 - (1 - 0,874) \cdot (1 - 0,599 \cdot 0,865) = 0,939$$

Moderne Werkzeuge werden BDDs verwenden, damit ergibt sich die System-Unzuverlässigkeit zu

$$F_{\text{sys}}(T) = F_A(T) + (1 - F_A(T)) \cdot (F_B(T) \cdot F_C(T)) \approx 0,874 + (1 - 0,874) \cdot (0,599 \cdot 0,865) = 0,939$$

Dabei wurde die Variablen-Reihenfolge A, B, C gewählt, jede andere Reihenfolge liefert dasselbe Ergebnis.

Methode 2:

Die Ausfallraten der Basis-Ereignisse A und B sind mit Formel (16) gegeben, die Nichtverfügbarkeiten $Q(t)$ durch Formel (18). Die System-Ausfallrate $h(t)$ kann nun mit Formel (65) berechnet werden. Die Unzuverlässigkeit wird damit zu $F_{\text{sys}}(T) \approx 0,964$ berechnet. Dies ist etwas zu groß, da die Nichtverfügbarkeiten hier sehr groß sind (das System also praktisch nicht verwendbar ist). Wendet man stattdessen die Formel (78) aus Anhang B.1 an, so ergibt sich das korrekte Ergebnis $F_{\text{sys}}(T) \approx 0,939$. Der Zeitverlauf von Ausfalldichte $f(t)$, Ausfallrate $h(t)$ und Unzuverlässigkeit $F(t)$ ist in Abbildung 38 gezeigt.

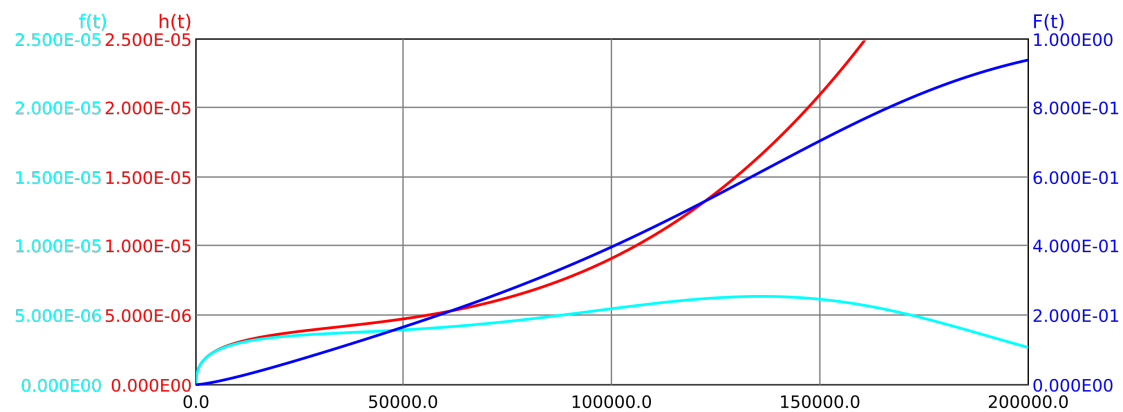


Abbildung 38: Zeitlicher Verlauf von Ausfallrate, Ausfalldichte und Unzuverlässigkeit für Beispiel 7.1

7.2 Berechnung mit Markov Modellen

Die System-Unzuverlässigkeit kann auch mittels Markov-Modellen berechnet werden. Die Berechnung muss nun grundsätzlich durch Integration des Differenzialgleichungssystems erfolgen (transiente Berechnung), da ja nie ein stationärer Zustand angenommen wird. Die System-Unzuverlässigkeit zu einem bestimmten Zeitpunkt ist durch die Summe der Aufenthaltswahrscheinlichkeiten in den Ausfallzuständen zu diesem Zeitpunkt gegeben.

Wenn wie in Beispiel 7.1 Ausfallverteilungen mit nicht-konstanten Ausfallraten verwendet werden, sind die Transitionsraten zeitabhängig. Das macht die Integration erheblich aufwändiger, da nun die zur Integration von steifen Differenzialgleichungssystemen benötigte Jakobi-Matrix bei jedem Zeitschritt mindestens einmal invertiert werden muss.

Beispiel 7.2 *Abbildung 39 zeigt das Markov-Modell, welches die Struktur des Fehlerbaums aus Beispiel 7.1 abbildet.*

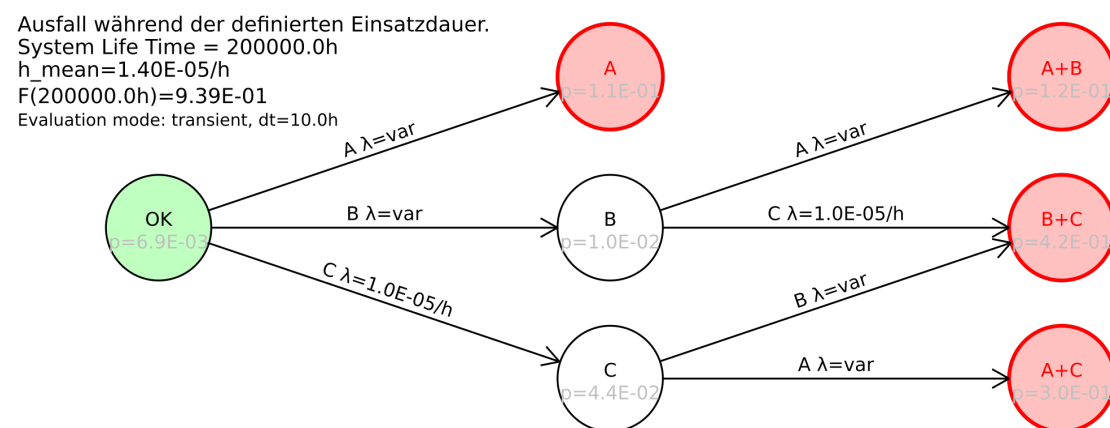


Abbildung 39: Markov Modell mit Berechnung der Unzuverlässigkeit

Die Ergebnisse sind praktisch identisch zu denen, die in Beispiel 7.1 mit Fehlerbäumen berechnet wurden.

A Modelle für Basis-Ereignisse

Die Basis-Ereignisse von Fehlerbäumen oder die Kanten (Transitionen) von Markov-Modellen modellieren Ereignisse oder Bedingungen (Zustände). Hierfür werden unterschiedliche Modelle benötigt, je nachdem ob es sich um einmalige oder mehrmalige zufällige Ereignisse handelt, um regelmäßige Ereignisse, oder um Bedingungen (Zustände). Nachfolgend werden die wichtigsten erwähnt, es gibt aber auch andere.

A.1 Modell Wiederherstellbar

Mit diesem Modell, das auch als „Dormant Failure Model“ (deutsch: schlafende-Fehler-Modell) oder „Reparierbar“ oder „Testbar“ bezeichnet wird ²⁵, können beispielsweise die folgenden Ausfallszenarien modelliert werden:

- Ausfälle, die direkt zum Systemausfall führen (und damit unmittelbar erkannt werden)
- Ausfälle, die nur in Gegenwart oder bei nachfolgendem Auftreten anderer Ereignisse zum Systemausfall führen, die aber durch regelmäßige Tests erkannt werden können

Dieses Modell ist das Standardmodell, mit diesem können fast alle Ausfälle von technischen Komponenten wie Elektrik/Elektronik, Hydraulik, Pneumatik etc. abgebildet werden, von einer einfachen Leitung bis hin zu ganzen Steuerungen.

Wenn weder Detektionszeit noch Reparaturzeit vernachlässigbar klein sind, ist die Nichtverfügbarkeit $Q(t)$ durch die in Kapitel 4 erwähnte Formel (48) sehr gut angenähert:

$$Q(t) = 1 - \frac{e^{-\frac{\lambda \cdot (t \bmod T_{\text{test}})}{\lambda \cdot \text{MRT} + 1}}}{\lambda \cdot \text{MRT} + 1}$$

mit dem Testintervall T_{test} und der mittleren Reparaturzeit pro Ausfall MRT (einschließlich Zeit zur Ersatzbeschaffung).

Die mittlere Nichtverfügbarkeit $\bar{Q}$ ist gegeben durch die in Kapitel 4 hergeleitete Formel (44)

$$\bar{Q} = \frac{e^{-\lambda \cdot T_{\text{test}}} - 1}{\lambda \cdot T_{\text{test}} + \lambda \cdot \text{MRT} \cdot (1 - e^{-\lambda \cdot T_{\text{test}}})} + 1$$

sowie die dort erwähnte Formel (45) bei vernachlässigbarer Detektionszeit $T_{\text{test}} \rightarrow 0$

$$\bar{Q} = \frac{\lambda \cdot \text{MRT}}{\lambda \cdot \text{MRT} + 1}$$

bzw. Formel (46) bei vernachlässigbarer Reparaturzeit $\text{MRT} \rightarrow 0$:

$$\bar{Q} = \frac{e^{-\lambda \cdot T_{\text{test}}} - 1}{\lambda \cdot T_{\text{test}}} + 1$$

²⁵Manche Werkzeuge bieten separate Modell für Testbar und Reparierbar an. Die klare Trennung vereinfacht das Verständnis, jedoch werden manchmal sowohl eine Fehlerdetektionszeit größer null als auch eine Reparaturzeit größer null benötigt.

Das Modell kann oft auch verwendet werden, wenn es kein definiertes Test- oder Inspektionsintervall gibt, jedoch ein Ausfall sich während des Betriebs in einer unkritischen Situation offenbart. Falls es keine Tests gibt und sich ein Ausfall nur offenbart, wenn noch ein anderes Ereignis eintritt, muss das Testintervall auf die nominale Einsatzdauer gesetzt werden, oder das Modell „Nicht-Wiederherstellbar“ verwendet werden.

Wenn die Ausfallrate nicht konstant ist, muss eine mittlere Ausfallrate über Formel (26) in Verbindung mit (25) berechnet werden:

$$\lambda_{\text{eff}} = \frac{1}{\text{MTTF}} = \frac{1}{\int_0^{\infty} t \cdot f(t) dt} = \frac{1}{\int_0^{\infty} t \cdot h(t) \cdot e^{-\int_0^t \sum_{i=1}^n h_i(\tau) d\tau} dt} \quad (74)$$

Das Modell kann auch als Beschreibung von Bedingungen verwendet werden.

A.2 Modell Nicht-Wiederherstellbar

Mit diesem Modell können beispielsweise die folgenden Ausfallszenarien modelliert werden:

- Ausfälle, die direkt zum Systemausfall führen
- Ausfälle, die nur zum Systemausfall führen, wenn auch schon andere Ereignisse eingetreten sind oder noch eintreten, und die auch nur dann erkannt werden

Nur bei diesem Ereignismodell macht es Sinn, bei einer transienten Berechnung mit zeitvarianten Ausfallraten (also beispielsweise Weibull-Verteilungen) zu rechnen.

Die Unzuverlässigkeit berechnet sich gemäß der bekannten Formel (8) zu

$$F(T) = 1 - e^{-\int_0^T h(t) dt} \quad (75)$$

Die mittlere Ausfallrate bezüglich $F(T)$ ergibt sich damit zu

$$\begin{aligned} F(T) &= 1 - e^{-\int_0^T h(t) dt} = 1 - e^{-\overline{h(T)} \cdot T} \\ \Rightarrow \overline{h(T)} &= \frac{1}{T} \int_0^T h(t) dt \end{aligned} \quad (76)$$

Wenn also eine Unzuverlässigkeit $F(T)$ für ein nicht-reparierbares Element mit einer prinzipiell zeitabhängigen Ausfallrate berechnet werden soll, muss entweder die Unzuverlässigkeit über die Integration mittels Formel (75) berechnet werden, oder es muss eine konstante Ausfallrate verwendet werden, welche gemäß Formel (76) berechnet wurde. Es darf nicht die gemäß Formel (29) berechnete mittlere effektive Ausfallrate λ_{eff} verwendet werden!

Die mittlere Nichtverfügbarkeit über die geplante Einsatzzeit $\overline{Q(T)}$ ist gegeben durch

$$\overline{Q(T)} = \frac{\int_0^T F(t) dt}{T} \quad (77)$$

Die maximale Nichtverfügbarkeit ergibt sich am Ende der geplanten Einsatzzeit.

Das Modell darf offensichtlich nur verwendet werden, wenn sicher ist, dass die Komponente zum Zeitpunkt $t = 0$ nicht defekt ist. Es verbietet sich also, irgendwelche kurzen (geplanten) Einsatzzeiten wie etwa einen einzelnen Flug oder eine einzelne Autofahrt anzunehmen, da vor diesen nicht sichergestellt ist, dass alle Komponenten wie neu sind (im Gegensatz zu einer Raumfahrt-Mission).

Das Modell kann auch als Beschreibung von Bedingungen verwendet werden.

A.3 Modell Konstant

Insbesondere für konstante Nichtverfügbarkeiten $Q = \text{const}$, aber auch für die (konstante) Wahrscheinlichkeit, dass eine externe Randbedingung erfüllt ist, wird das Modell „Konstante“ verwendet.

Das Modell wird (fast) nur als Beschreibung von Bedingungen verwendet.

A.4 Allgemeine Empfehlungen zu Basismodellen

Es ist sinnvoll, mehrere Ausfallmodi in einem Basis-Ereignis zusammenzufassen. Wichtig ist, dass alle Ausfallmodi, die von diesem einen Basis-Ereignis modelliert werden sollen, sich bezüglich der Modellierung der Wiederherstellung nicht unterscheiden, also insbesondere die gleiche Fehleroffenbarungszeit und ggf. Reparaturzeit besitzen.

Bei komplexen Komponenten (wie beispielsweise einer Elektronik-Karte oder einem ganzen Steuerungssystem) ist es weder notwendig noch sinnvoll, jede einzelne der zig-tausend Ausfallarten einzeln als Basis-Ereignis aufzunehmen. Typischerweise ist die Anzahl der an den Schnittstellen beobachtbaren Ausfallarten ohnehin recht überschaubar (typisch: binäres Signal High statt Low oder umgekehrt, analoges Signal zu hoch/zu niedrig, Buskommunikation ausgefallen/Lebenszeichen ungültig, Busvariable unerkennbar falsch).

So ist es in der Regel sinnvoll, eine komplexe Baugruppe mit zwei bis vier Basis-Ereignissen zu beschreiben:

1. ein Basis-Ereignis für Ausfälle, die sofort erkannt werden,
2. ein Basis-Ereignis für Ausfälle, die im Rahmen von täglichen Tests erkannt werden,
3. ein Basis-Ereignis für Ausfälle, die im Rahmen von regelmäßigen Inspektionen erkannt werden,
4. ein Basis-Ereignis für Ausfälle, die nie bzw. nur im Anforderungsfall erkannt werden.

Die Zuordnung aller möglichen Ausfälle zu einer dieser Kategorien erfolgt typischerweise in einer FMEA. Die Ausfallrate für jedes dieser maximal vier Basis-Ereignisse ist die Summe der Einzel-Ausfallraten aller Fehlermodi, die in der FMEA der jeweiligen Kategorie zugeordnet wurden.

B Ergänzungen zur Berechnung der System-Ausfallrate mit Fehlerbäumen

B.1 Verbesserung bei sehr hohen Nichtverfügbarkeiten

Formel (66) aus Abschnitt 6 liefert für große Nichtverfügbarkeiten etwas zu konservative Ergebnisse. Da große Nichtverfügbarkeiten immer ein Zeichen für nicht korrekt ausgelegte Systeme sind, ist dies in der Praxis nicht relevant. Dennoch soll hier noch eine Formel erwähnt werden, die weniger konservative Ergebnisse liefert.

Verwendet man statt der Ausfallraten jeweils die bezüglich der Wiederherstellung bedingten Ausfallhäufigkeiten $w(t)$, so kann man zunächst die System-Ausfallhäufigkeit $w_{\text{sys}}(t)$ mittels folgender Formel berechnen:

$$\begin{aligned}
 h_{\text{sys}}(t) &= \sum_{i=1}^{n_{\text{MCS}}} \frac{w_{\text{MCS}_i}(t)}{1 - Q_{\text{MCS}_i}(t)} \\
 &= \sum_{i=1}^{n_{\text{MCS}}} \frac{\sum_{j=1}^{n_{\text{Lit}, \text{MCS}_i}} \left(w_j(t) \cdot \prod_{k=1, k \neq j}^{n_{\text{Lit}, \text{MCS}_i}} q_{i,k}(t) \right)}{1 - \prod_{k=1}^{n_{\text{Lit}, \text{MCS}_i}} q_{i,k}(t)}
 \end{aligned} \tag{78}$$

C Weitere Verteilungsfunktionen

C.1 Normal-Verteilung

Die Normal-Verteilung hat die Dichtefunktion

$$f(t) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(t-\mu)^2}{2\sigma^2}} \quad (79)$$

mit dem Mittelwert $\mu = \text{MTTF}$ und der Standardabweichung σ . Dabei ist zu beachten, dass die Funktion bereits bei $t = -\infty$ beginnt bzw. dieser Anteil nicht zu 0 gesetzt werden kann.

Die Verteilungsfunktion ist folglich gegeben durch

$$F(t) = \int_{-\infty}^t f(t)dt = \frac{1}{\sqrt{2\pi\sigma^2}} \int_{-\infty}^t e^{-\frac{(t-\mu)^2}{2\sigma^2}} dt = 0,5 \left(1 + \text{erf} \left(\frac{t-\mu}{\sqrt{2\sigma^2}} \right) \right) \quad (80)$$

wobei $\text{erf}(x)$ die sogenannte Fehlerfunktion ist. Für dieses Integral existiert keine geschlossene Darstellung, es muss daher immer numerisch bestimmt werden. Entsprechend gibt es auch für die Ausfallrate $h(t)$ keine geschlossene Darstellung.

Die Abbildung 40 zeigt eine Normal-Verteilung mit Mittelwert $\mu = 1\text{E}6 \text{ h}$ und Standardabweichung $\sigma = 1\text{E}5 \text{ h}$.

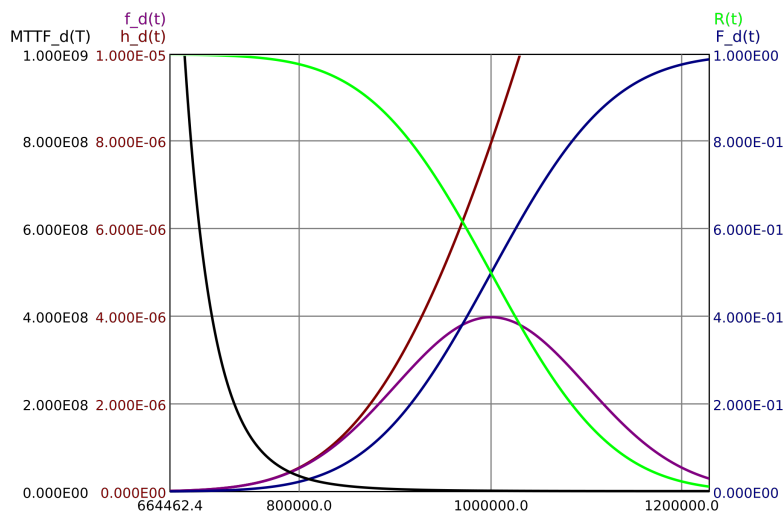


Abbildung 40: Normal-Verteilung mit Mittelwert $\mu = 1\text{E}6 \text{ h}$ und Standardabweichung $\sigma = 1\text{E}5 \text{ h}$

C.2 Gleichverteilung

Die Gleichverteilung zeichnet sich durch eine konstante Ausfalldichte $f(t) = \text{const}$ innerhalb eines Intervalls $t_1 \dots t_2$ aus. Außerhalb dieses Intervalls ist sie 0. Sie wird daher

auch Rechteckverteilung genannt. Sie kommt in der Natur und Technik nicht vor, eignet sich jedoch für Gedankenexperimente oder zur Plausibilisierung von Formeln.

$$f(t) = \frac{1}{t_2 - t_1} \quad \text{für } t_1 \leq t < t_2 \quad , \text{sonst } 0 \quad (81)$$

$$F(t) = \begin{cases} 0 & \text{für } t < t_1 \\ \frac{t - t_1}{t_2 - t_1} & \text{für } t_1 \leq t < t_2 \\ 1 & \text{für } t \geq t_2 \end{cases} \quad (82)$$

$$R(t) = \begin{cases} 1 & \text{für } t < t_1 \\ \frac{t_2 - t}{t_2 - t_1} & \text{für } t_1 \leq t < t_2 \\ 0 & \text{für } t \geq t_2 \end{cases} \quad (83)$$

$$h(t) = \frac{\frac{1}{t_2 - t_1}}{\frac{t_2 - t}{t_2 - t_1}} = \frac{1}{t_2 - t_1} \cdot \frac{t_2 - t_1}{t_2 - t} = \frac{1}{t_2 - t} \quad \text{für } t_1 \leq t < t_2 \quad , \text{sonst } 0 \quad (84)$$

$$\text{MTTF} = \int_{t_1}^{t_2} t \cdot \frac{1}{t_2 - t_1} dt = \frac{1}{t_2 - t_1} \left[\frac{t^2}{2} \right]_{t_1}^{t_2} = \frac{1}{t_2 - t_1} \frac{t_2^2 - t_1^2}{2} = \frac{t_1 + t_2}{2} \quad (85)$$

Eine Gleichverteilung ist in Abbildung 41 dargestellt. Man sieht, dass die Ausfallrate $h(t)$ für $t \rightarrow t_2$ gegen unendlich geht, also bei t_2 eine Polstelle hat.

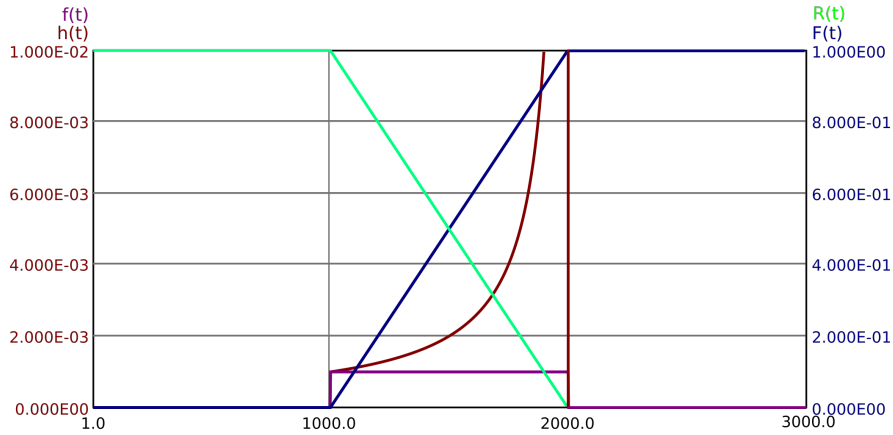


Abbildung 41: Gleichverteilung zwischen t_1 und t_2

C.3 Delta-Verteilung

Die Delta-Verteilung ist die mathematische Beschreibung des Determinismus: Nur zum Zeitpunkt T nimmt die Dichte $f(t)$ einen Wert ungleich 0 an. Die Dichte muss also die Eigenschaft eines Dirac-Stoßes der Höhe 1 zum Zeitpunkt T haben:

$$f(t) = \delta(t - T) \quad (86)$$

Dabei ist $\delta(t)$ die Dirac-Funktion mit der Eigenschaft $\int_{-\infty}^{+\infty} \delta(t) dt = 1$. Zum Zeitpunkt T springt also die Unzuverlässigkeit von 0 auf 1:

$$F(t) = \int_0^{\infty} f(t) dt = \int_0^{\infty} \delta(t - T) dt = \sigma(t - T) \quad (87)$$

Dabei bezeichnet $\sigma(t - T)$ die Einheitsprung-Funktion zur Zeit T . Für die Zuverlässigkeit ergibt sich unmittelbar:

$$R(t) = 1 - \sigma(t - T) \quad (88)$$

Die MTTF ist offensichtlich T , was sich auch rechnerisch ergibt:

$$\text{MTTF} = \int_0^{\infty} t \cdot \delta(t - T) dt = T \quad (89)$$

Die Delta-Verteilung ist ein Spezialfall zahlreicher Verteilungen, zum Beispiel der Gleichverteilung (nämlich für $t_1 \rightarrow t_2$) oder der Normalverteilung (für Streuung $\sigma \rightarrow 0$).

D Importanzen

Importanzen geben den Einfluss eines jeden Basis-Ereignisses auf eine System-Kenngröße an. In der Literatur finden sich eine ganze Reihe von Importanzen, welche oft unterschiedlich definiert sind, und fast immer ohne Nennung der System-Kenngröße, für die sie definiert wurden. So wird auch bezüglich der Importanzen oft nur von „Ausfallwahrscheinlichkeiten“ gesprochen.

Importanzen für die System-Ausfallrate h_{sys} (in [IEC 61508] PFH genannt) sucht man in der Literatur praktisch vergeblich. Das ist insofern nachvollziehbar, als dass Importanzen fast immer im Zusammenhang mit Fehlerbäumen definiert werden, und die Berechnung der System-Ausfallrate mit Fehlerbäumen ebenfalls nur selten (z. B. in [NUREG]) behandelt wird. Einige Importanzen lassen sich direkt auf die Ausfallrate übertragen, einige sinngemäß, und einige Importanzen können für die Ausfallrate gar nicht sinnvoll definiert werden.

Wenngleich Importanzen meist für die Verwendung mit Fehlerbäumen definiert wurden, lassen sich doch einige auch auf andere Modelle, wie etwa Markov-Modelle anwenden.

D.1 Allgemeine Hinweise

In den Abschnitten 5 bis 7 wurde dargestellt, dass oft eine transiente (zeitabhängige) Berechnung nötig ist, um korrekte Werte zu erhalten. Bei Importanzen kann man hierauf oft verzichten, da zum Einen viele Importanzen per Definition schon relative Größen darstellen, sich Ungenauigkeiten in Zähler und Nenner also aufheben, und zum Anderen der Zweck der Importanzen nur der ist, Basis-Ereignisse oder Minimalschnitte zu priorisieren, wofür es ebenfalls nur auf Verhältnisse oder Größenordnungen und nicht auf bestimmte Zahlenwerte ankommt.

Um die Formeln kurz und einprägsam zu halten, wird im Folgenden auf die Erwähnung der Abhängigkeit von der Systemlebenszeit T oder der Mittelwertbildung verzichtet: Statt $F(T)$ wird kurz F geschrieben, statt $\bar{Q}(T)$ wird kurz Q geschrieben, und statt $\bar{h}(T)$ wird kurz h .

D.2 Partielle Ableitung (PD) und Birnbaum-Importanz (BI)

Unmittelbar naheliegend als Maß für die Wichtigkeit einzelner Basis-Ereignisse ist die partielle Ableitung (partial derivative, PD) des System-Werts Q , F oder h .

Die partiellen Ableitungen der System-Nichtverfügbarkeit Q und der System-Zuverlässigkeit F werden auch Birnbaum-Importanz genannt.²⁶

²⁶Es ist keine Quelle bekannt, die eine partielle Ableitung der System-Ausfallrate als „Birnbaum-Importanz“ bezeichnet.

D.2.1 Partielle Ableitung für die System-Nichtverfügbarkeit

Die Ableitung der System-Nichtverfügbarkeit Q_{Sys} nach der Nichtverfügbarkeit jedes Basis-Ereignisses Q_x ist gegeben durch:

$$I_{Q,x}^{\text{PD}} = \frac{\partial Q_{\text{Sys}}}{\partial Q_x} = \frac{Q_{\text{Sys}}(\mathbf{Q} + \partial Q_x) - Q_{\text{Sys}}(\mathbf{Q})}{\partial Q_x} \quad (90)$$

Dabei bezeichnet $\mathbf{Q}$ den Vektor der (Mittelwerte der) Nichtverfügbarkeiten aller Basis-Ereignisse.

Mit der Näherungsformel (53) für die System-Nichtverfügbarkeit für Fehlerbäume berechnet sich die partielle Ableitung zu

$$I_{Q,x}^{\text{PD}} \approx \frac{\partial \sum_{i=1}^{n_{\text{MCS}}} \left(\prod_{j=1}^{m_{\text{Lit},i}} Q_j(t) \right)}{\partial Q_x} = \sum_{i=1}^{n_{\text{MCS}}} \begin{cases} 0 & \text{wenn Basis-Ereignis } x \notin \text{MCS}_i \\ \prod_{j=1, j \neq x}^{m_{\text{Lit},i}} Q_j & \text{wenn Basis-Ereignis } x \in \text{MCS}_i \end{cases} \quad (91)$$

Bei Verwendung von BDDs kann die partielle Ableitung $\frac{\partial Q_{\text{Sys}}}{\partial Q_x}$ leicht exakt ermittelt werden. Verschiebt man jedes Basis-Ereignis der Reihe nach an die Spitze des BDDs, wie in Abbildung 42 dargestellt, ergibt sich

$$I_{Q,x}^{\text{PD}} = \frac{\partial((1 - Q_x) \cdot \text{BDD}_0 + Q_x \cdot \text{BDD}_1)}{\partial Q_x} = \text{BDD}_1 - \text{BDD}_0 \quad (92)$$

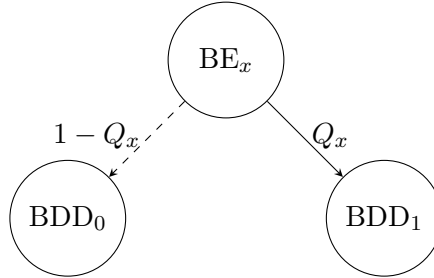


Abbildung 42: Zur Berechnung der partiellen Ableitung mit BDDs

Dabei ist BDD_0 der Low-Zweig für Basis-Ereignis x also die System-Nichtverfügbarkeit im Fall, dass Basis-Ereignis x nicht ausgefallen ist, und BDD_1 der High-Zweig also die System-Nichtverfügbarkeit im Fall, dass Basis-Ereignis x ausgefallen ist. Damit kann man also auch schreiben

$$I_{Q,x}^{\text{PD}} = \text{BDD}_{x,1} - \text{BDD}_{x,0} = Q_{\text{Sys}}(Q_x := 1) - Q_{\text{Sys}}(Q_x := 0) \quad (93)$$

Dabei meint $Q_{\text{Sys}}(Q_x := 1)$ die System-Nichtverfügbarkeit, die sich ergibt, wenn man die Nichtverfügbarkeit von Basis-Ereignis x zu 1 setzt, und die Nichtverfügbarkeit aller anderen Basis-Ereignisse auf ihren ursprünglichen Werten belässt.

Da $\text{BDD}_{x,0}$ die Wahrscheinlichkeit angibt, mit der das System nicht verfügbar ist, auch wenn Komponente x in Ordnung ist, und $\text{BDD}_{x,1}$ die Wahrscheinlichkeit, dass das System

dann nicht verfügbar ist, wenn auch noch Komponente x ausfällt, ist die Differenz die Wahrscheinlichkeit, dass das System sich in einem Zustand befindet, in dem Komponente x kritisch ist, also der Ausfall von Komponente x zum Systemausfall führen würde.

D.2.2 Partielle Ableitung für die System-Unzuverlässigkeit

Auch für die System-Unzuverlässigkeit kann die partielle Ableitung angegeben werden:

$$I_{F,x}^{\text{PD}} = \frac{\partial F_{\text{Sys}}}{\partial F_x} = \frac{F_{\text{Sys}}(\mathbf{F} + \partial F_x) - F_{\text{Sys}}(\mathbf{F})}{\partial F_x} \quad (94)$$

Dabei bezeichnet $\mathbf{F}$ den Vektor der Unzuverlässigkeiten aller Basis-Ereignisse zu einem bestimmten Zeitpunkt (in der Regel zum System-Lebensende).

Beispiel D.1 Der Fehlerbaum eines Systems sei BE1 UND BE2. Damit gilt:

$$F_{\text{sys}}(T) = F_{\text{BE1}}(T) \cdot F_{\text{BE2}}(T)$$

Die Ableitung nach $F_{\text{BE1}}(T)$ ist $F_{\text{BE2}}(T)$ und umgekehrt.

Wie in Abschnitt 7 erläutert, kann ein Fehlerbaum zur Berechnung der System-Unzuverlässigkeit auch Bedingungen enthalten, also Basis-Ereignisse, die durch ihre Nichtverfügbarkeit Q beschrieben sind. Für diese Basis-Ereignisse kann man ersatzweise die partielle Ableitung $I_{F,x}^{\text{PD}} = \frac{\partial F_{\text{Sys}}}{\partial Q_x}$ bilden, allerdings gelten die oben genannten Formeln nicht oder nur noch näherungsweise.

D.2.3 Partielle Ableitung für die System-Ausfallrate

Für die System-Ausfallrate h macht eine partielle Ableitung nur nach der Eintrittsrate h_x eines Basis-Ereignisses $\frac{\partial h_{\text{Sys}}}{\partial h_x}$ wenig Sinn, da die System-Ausfallrate h gemäß Formel (66) auch von der Nichtverfügbarkeit jedes Basis-Ereignisses abhängt:

$$h_{\text{sys}}(t) \lesssim \sum_{i=1}^{n_{\text{MCS}}} \left(\sum_{j=1}^{n_{\text{Lit},\text{MCS}_i}} \left(h_j(t) \cdot \prod_{k=1, k \neq j}^{n_{\text{Lit},\text{MCS}_i}} q_{i,k}(t) \right) \right) \quad (95)$$

Man könnte natürlich zwei Ableitungen $I_{h,h,x}^{\text{PD}} = \frac{\partial h_{\text{Sys}}}{\partial h_x}$ und $I_{h,Q,x}^{\text{PD}} = \frac{\partial h_{\text{Sys}}}{\partial Q_x}$ bilden. Allerdings hängt Q_x bei den meisten Basis-Ereignissen wiederum von der Ausfallrate desselben ab:

$$h_{\text{Sys}} = \text{fkt}(h_x, Q_x = \text{fkt}(h_x)) \quad (96)$$

Für regelmäßig getestete und reparierte Komponenten etwa gilt für die mittlere Nichtverfügbarkeit $\bar{Q} \approx \lambda \cdot (T_{\text{Test}}/2 + \text{MRT}) = h \cdot (T_{\text{Test}}/2 + \text{MRT})$.

Sinnvoller ist es daher, die Importanz $I_{h,x}^{\text{PD}}$ als Ableitung nach der (mittleren) Ausfallrate des Basis-Ereignisses λ_i zu definieren:

$$I_{h,x}^{\text{PD}} = \frac{\partial h_{\text{Sys}}}{\partial \lambda_x} = \frac{h_{\text{Sys}}(\boldsymbol{\lambda} + \partial \lambda_x) - h_{\text{Sys}}(\boldsymbol{\lambda})}{\partial \lambda_x} \approx \frac{\partial \left(\sum_{i=1}^{n_{\text{MCS}}} h_{\text{MCS},i} \right)}{\partial \lambda_x} = \sum_{i=1}^{n_{\text{MCS}}} \frac{\partial h_{\text{MCS},i}}{\partial \lambda_x} \quad (97)$$

Mit Formel (65) für die Ausfallrate $h_{\text{MCS},i}$ jedes Minimalschnitts

$$\begin{aligned}
 h_{\text{MCS}} &\lesssim h_1 \cdot Q_2 \cdot Q_3 \cdot \dots \cdot Q_m \\
 &\quad + h_2 \cdot Q_1 \cdot Q_3 \cdot \dots \cdot Q_m \\
 &\quad + \dots \\
 &\quad + h_m \cdot Q_1 \cdot Q_2 \cdot \dots \cdot Q_{m-1}
 \end{aligned} \tag{98}$$

ergibt sich

$$\begin{aligned}
 \frac{\partial h_{\text{MCS},i}}{\partial \lambda_x} &\approx \frac{\partial(h_1 \cdot Q_2 \cdot Q_3 \cdot \dots \cdot Q_m)}{\partial \lambda_x} \\
 &\quad + \frac{\partial(h_2 \cdot Q_1 \cdot Q_3 \cdot \dots \cdot Q_m)}{\partial \lambda_x} \\
 &\quad + \dots \\
 &\quad + \frac{\partial(h_m \cdot Q_1 \cdot Q_2 \cdot \dots \cdot Q_{m-1})}{\partial \lambda_x} \\
 &= \sum_{j=1}^m \frac{\partial \left(h_j \cdot \prod_{k=1, k \neq j}^m Q_k \right)}{\partial \lambda_x}
 \end{aligned} \tag{99}$$

Wenn Basis-Ereignis x in MCS_i nicht enthalten ist, ist diese Ableitung null. Andernfalls ist der Summand mit $j = x$ gleich $\prod_{k=1, k \neq j}^m Q_k$ (wobei die Nichtverfügbarkeiten dieses Produkts alle von Basis-Ereignis x unabhängig sind), und alle Summanden mit $j \neq x$ sind gleich $h_j \frac{\partial Q_x}{\partial \lambda_x} \prod_{k=1, k \neq j, k \neq x}^m Q_k$.

Somit gilt

$$I_{h,x}^{\text{PD}} \approx \sum_{i=1}^{n_{\text{MCS}}} \begin{cases} 0 & \text{wenn BE } x \notin \text{MCS}_i \\ \prod_{k=1, k \neq x}^{m_{\text{Lit},i}} Q_k + \frac{\partial Q_x}{\partial \lambda_x} \cdot \sum_{j=1, j \neq x}^{m_{\text{Lit},i}} \left(h_j \cdot \prod_{k=1, k \neq j, k \neq x}^{m_{\text{Lit},i}} Q_k \right) & \text{wenn BE } x \in \text{MCS}_i \end{cases} \tag{100}$$

Beispiel D.2 Ein System bestehe aus zwei unterschiedlichen Komponenten mit konstanten Ausfallraten λ_1 und λ_2 , welche in unterschiedlichen Intervallen $T_{\text{Test},i}$ regelmäßig getestet und ggf. umgehend repariert werden. Das System falle dann gefährlich aus, wenn eine der Komponenten ausgefallen ist, und in diesem Zustand noch die zweite Komponente ausfällt. Der Fehlerbaum ist damit BE1 UND BE2. Es gibt also nur einen Minimalschnitt, nämlich $\{\text{BE1}, \text{BE2}\}$. Damit gilt:

$$\overline{h_{\text{sys}}} \lesssim \lambda_1 \cdot \overline{Q_2} + \lambda_2 \cdot \overline{Q_1}$$

Für die mittlere Nichtverfügbarkeit jeder Komponente gilt $\overline{Q_x} \approx \lambda_x \cdot T_{\text{Test},x}/2$ und somit für deren Ableitung nach λ_x : $\frac{\partial Q_x}{\partial \lambda_x} \approx T_{\text{Test},x}/2$.

Somit gilt

$$I_{h,1}^{\text{PD}} \approx \overline{Q_2} + \lambda_2 \cdot \frac{T_{\text{Test},1}}{2} = \lambda_2 \cdot \frac{T_{\text{Test},2}}{2} + \lambda_2 \cdot \frac{T_{\text{Test},1}}{2} = \lambda_2 \frac{T_{\text{Test},1} + T_{\text{Test},2}}{2}$$

und

$$I_{h,2}^{\text{PD}} \approx \lambda_1 \cdot \frac{T_{\text{Test},2}}{2} + \overline{Q_1} = \lambda_1 \cdot \frac{T_{\text{Test},2}}{2} + \lambda_1 \cdot \frac{T_{\text{Test},1}}{2} = \lambda_1 \frac{T_{\text{Test},1} + T_{\text{Test},2}}{2}$$

D.2.4 Berechnung für Markov-Modelle

Aufgrund der oben erwähnten Eigenschaft, dass die partiellen Ableitungen der Nichtverfügbarkeit bzw. Unzuverlässigkeit gleich der Wahrscheinlichkeit sind, dass sich das System in einem Zustand befindet, von dem aus es bei Eintritt von Ereignis x in einen Ausfallzustand gelangt, ist die partiellen Ableitung nach Q_x bzw. F_x gleich der Summe der (mittleren) Aufenthaltswahrscheinlichkeiten aller m_x Zustände, von denen aus eine Kante des Basis-Ereignisses x zu einem Ausfallzustand führt:

$$I_{Q,x}^{\text{B}} = \sum_{j=1}^{m_x} \overline{p_j} \quad (101)$$

D.3 Risk-Reduction (RR)

Das Risiko-Reduzierung-Potenzial (engl. Risk Reduction, RR) gibt an, wie sehr $\overline{Q}$, $F(T)$ oder $\overline{h}$ reduziert würden, wenn Basis-Ereignis BE_x nie eintreten würde, also Komponente x nicht ausfallen könnte.

$$I_{Q,x}^{\text{RR}} = Q_{\text{Sys}}(\mathbf{Q}) - Q_{\text{Sys}}(Q_x := 0) \quad (102)$$

$$I_{F,x}^{\text{RR}} = F_{\text{Sys}}(\mathbf{F}) - F_{\text{Sys}}(F_x := 0) \quad (103)$$

Das Verbesserungs-Potential kann man unmittelbar auch auf die System-Ausfallrate anwenden, da aufgrund der Definition unerheblich ist, durch welche Größe die Qualität eines Basis-Ereignisses definiert ist – oder durch welche Kombination von Größen. Allerdings muss man dann sinnvollerweise gleichzeitig $h_x = 0$ und $Q_x = 0$ setzen:

$$I_{h,x}^{\text{RR}} = h_{\text{Sys}}(\mathbf{h}, \mathbf{Q}) - h_{\text{Sys}}(h_x := 0, Q_x := 0) \quad (104)$$

D.4 Risk-Reduction-Worth (RRW)

Der Risiko-Reduzierungs-Wert (engl. Risk-Reduction-Worth, RRW) gibt an, wie sehr $\overline{Q}$, $F(T)$ oder $\overline{h}$ relativ reduziert würden, wenn Komponente x nicht ausfallen würde:

$$I_{Q,x}^{\text{RRW}} = \frac{Q_{\text{Sys}}(\mathbf{Q}) - Q_{\text{Sys}}(Q_x := 0)}{Q_{\text{Sys}}(Q_x := 0)} = \frac{Q_{\text{Sys}}(\mathbf{Q})}{Q_{\text{Sys}}(Q_x := 0)} - 1 \quad (105)$$

$$I_{F,x}^{\text{RRW}} = \frac{F_{\text{Sys}}(\mathbf{F}) - F_{\text{Sys}}(F_x := 0)}{F_{\text{Sys}}(F_x := 0)} = \frac{F_{\text{Sys}}(\mathbf{F})}{F_{\text{Sys}}(F_x := 0)} - 1 \quad (106)$$

$$I_{h,x}^{\text{RRW}} = \frac{h_{\text{Sys}}(\mathbf{h}, \mathbf{Q}) - h_{\text{Sys}}(h_x := 0, Q_x := 0)}{h_{\text{Sys}}(h_x := 0, Q_x := 0)} = \frac{h_{\text{Sys}}(\mathbf{h}, \mathbf{Q})}{h_{\text{Sys}}(h_x := 0, Q_x := 0)} - 1 \quad (107)$$

Die Risk-Reduction-Worth kann offensichtlich beliebig große Werte annehmen. Je größer, desto wirksamer ist die Verbesserung von Komponente x . Ein Wert von ≈ 0 hingegen bedeutet, dass die Komponente x praktisch keinen Einfluss hat. Achtung: Der Summand -1 wird häufig weggelassen.

D.5 Fussell-Vesely-Importanz (FV)

Dividiert man das Risiko-Reduzierungs-Potenzial durch die ursprüngliche System-Größe, so ergibt sich die Fussell-Vesely-Importanz:

$$I_{Q,x}^{\text{FV}} = \frac{I_{Q,x}^{\text{RR}}}{Q_{\text{Sys}}(\mathbf{Q})} = \frac{Q_{\text{Sys}}(\mathbf{Q}) - Q_{\text{Sys}}(Q_x := 0)}{Q_{\text{Sys}}(\mathbf{Q})} \quad (108)$$

$$I_{F,x}^{\text{FV}} = \frac{I_{F,x}^{\text{RR}}}{F_{\text{Sys}}(\mathbf{F})} = \frac{F_{\text{Sys}}(\mathbf{F}) - F_{\text{Sys}}(F_x := 0)}{F_{\text{Sys}}(\mathbf{F})} \quad (109)$$

$$I_{h,x}^{\text{FV}} = \frac{I_{h,x}^{\text{RR}}}{h_{\text{Sys}}(\mathbf{h}, \mathbf{Q})} = \frac{h_{\text{Sys}}(\mathbf{h}, \mathbf{Q}) - h_{\text{Sys}}(h_x := 0, Q_x := 0)}{h_{\text{Sys}}(\mathbf{h}, \mathbf{Q})} \quad (110)$$

Die Fussell-Vesely-Importanz lässt sich sehr leicht auf Basis von Minimalschnitten berechnen: $Q_{\text{Sys}}(Q_x := 0)$ ist der Anteil der System-Nichtverfügbarkeit, der von den Minimalschnitten geliefert wird, die Basis-Ereignis x *nicht* enthalten. $Q_{\text{Sys}}(\mathbf{Q}) - Q_{\text{Sys}}(Q_x := 0)$ ist folglich der Anteil der System-Nichtverfügbarkeit, der von den Minimalschnitten geliefert wird, die Basis-Ereignis x enthalten. Damit gilt näherungsweise (für kleine Q_{MCS}):

$$I_{Q,x}^{\text{FV}} \approx \frac{\sum_{i=1}^{n_{\text{MCS}}} \begin{cases} 0 & \text{wenn } \text{BE}_x \notin \text{MCS}_i \\ Q_{\text{MCS},i} & \text{wenn } \text{BE}_x \in \text{MCS}_i \end{cases}}{Q_{\text{Sys}}(\mathbf{Q})} \quad (111)$$

Entsprechendes gilt für $I_{F,x}^{\text{FV}}$ und $I_{h,x}^{\text{FV}}$. Die Fussell-Vesely-Importanz ist somit die Wahrscheinlichkeit, dass mindestens ein Minimalschnitt, der Komponente x enthält, zum Systemausfall geführt hat, wenn das System ausgefallen ist.

Alternativ kann man auch die Formel von Esary-Proschan (54)

$$Q_{\text{sys}}(t) \lesssim 1 - \prod_{i=1}^{n_{\text{MCS}}} (1 - Q_{\text{MCS},i}(t)) \quad (112)$$

anwenden, dann erhält man

$$I_{Q,x}^{\text{FV}} \approx \frac{1 - \prod_{i=1}^{n_{\text{MCS}}} \begin{cases} 1 & \text{wenn } \text{BE}_x \notin \text{MCS}_i \\ 1 - Q_{\text{MCS},i} & \text{wenn } \text{BE}_x \in \text{MCS}_i \end{cases}}{Q_{\text{Sys}}(\mathbf{Q})} \quad (113)$$

D.6 Risk-Achievement (RA)

Für die Nichtverfügbarkeit und die Unzuverlässigkeit ist die Risiko-Erreichung (engl. Risk-Achievement, RA) wie folgt definiert:

$$I_{Q,x}^{\text{RA}} = Q_{\text{Sys}}(Q_x := 1) - Q_{\text{Sys}}(\mathbf{Q}) \quad (114)$$

$$I_{F,x}^{\text{RA}} = F_{\text{Sys}}(F_x := 1) - F_{\text{Sys}}(\mathbf{F}) \quad (115)$$

Mit der zuvor eingeführten Definition der partiellen Ableitung (Birnbaum-Importanz) und dem Risiko-Reduzierungs-Potenzial gilt unmittelbar:

$$\begin{aligned} I_{Q,x}^{\text{RA}} + I_{Q,x}^{\text{RR}} &= (Q_{\text{Sys}}(Q_x := 1) - Q_{\text{Sys}}(\mathbf{Q})) + (Q_{\text{Sys}}(\mathbf{Q}) - Q_{\text{Sys}}(Q_x := 0)) \\ &= Q_{\text{Sys}}(Q_x := 1) - Q_{\text{Sys}}(Q_x := 0) \\ &= I_{Q,x}^{\text{PD}} \end{aligned} \quad (116)$$

Für die System-Ausfallrate h lässt sich keine RA angeben, da die Ausfallrate einer Komponente (bzw. allgemein: Die Eintrittsrates eines Ereignisses) nicht dimensionslos ist und daher auch keinen oberen Grenzwert h_{max} kennt, und es somit auch keinen oberen Grenzwert $h_{\text{Sys}}(h_{\text{max},x})$ gibt.

D.7 Risk-Achievement-Worth (RAW)

Setzt man die RA ins Verhältnis zur ursprünglichen Systemgröße, so erhält man den Faktor, um den sich das Risiko vergrößern würde, wenn die Komponente x immer ausgefallen wäre (engl. Risk-Achievement-Worth, RAW):

$$I_{Q,x}^{\text{RAW}} = \frac{Q_{\text{Sys}}(Q_x := 1) - Q_{\text{Sys}}(\mathbf{Q})}{Q_{\text{Sys}}(\mathbf{Q})} = \frac{Q_{\text{Sys}}(Q_x := 1)}{Q_{\text{Sys}}(\mathbf{Q})} - 1 \quad (117)$$

$$I_{F,x}^{\text{RAW}} = \frac{F_{\text{Sys}}(F_x := 1) - F_{\text{Sys}}(\mathbf{F})}{F_{\text{Sys}}(\mathbf{F})} = \frac{F_{\text{Sys}}(F_x := 1)}{F_{\text{Sys}}(\mathbf{F})} - 1 \quad (118)$$

Achtung: Der Summand -1 wird häufig weggelassen.

Wie für die RA ist auch die RAW für Ausfallraten nicht anwendbar, da im Allgemeinen kein Grenzwert existiert.

D.8 Kritikalitäts-Importanz (CRI)

Die Kritikalitäts-Importanz (engl. Criticality Importance, CRI) ist definiert als das Verhältnis der relativen Änderung der Systemgröße zur relativen Änderung der Komponentengröße:

$$I_{Q,x}^{\text{CRI}} = \frac{\frac{\partial Q_{\text{Sys}}}{\partial Q_x}}{\frac{Q_{\text{Sys}}}{Q_x}} = \frac{Q_{\text{Sys}}(\mathbf{Q} + \partial Q_x) - Q_{\text{Sys}}(\mathbf{Q})}{Q_{\text{Sys}}(\mathbf{Q})} \cdot \frac{Q_x}{\partial Q_x} = I_{Q,x}^{\text{PD}} \cdot \frac{Q_x}{Q_{\text{Sys}}(\mathbf{Q})} \quad (119)$$

$$I_{F,x}^{\text{CRI}} = \frac{\frac{\partial F_{\text{Sys}}}{\partial F_x}}{\frac{F_{\text{Sys}}}{F_x}} = \frac{F_{\text{Sys}}(\mathbf{F} + \partial F_x) - F_{\text{Sys}}(\mathbf{F})}{F_{\text{Sys}}(\mathbf{F})} \cdot \frac{F_x}{\partial F_x} = I_{F,x}^{\text{PD}} \cdot \frac{F_x}{F_{\text{Sys}}(\mathbf{F})} \quad (120)$$

Sie kann auf die Ausfallrate erweitert werden, indem man wie bei der partiellen Ableitung die Komponentengrößen h_x und Q_x als Funktion der Ausfallrate der Komponente beschreibt:

$$I_{h,x}^{\text{CRI}} = \frac{\frac{\partial h_{\text{Sys}}}{\partial \lambda_x}}{\frac{h_{\text{Sys}}}{\lambda_x}} = \frac{h_{\text{Sys}}(\boldsymbol{\lambda} + \partial \lambda_x) - h_{\text{Sys}}(\boldsymbol{\lambda})}{h_{\text{Sys}}(\boldsymbol{\lambda})} \cdot \frac{\lambda_x}{\partial \lambda_x} = I_{h,x}^{\text{PD}} \cdot \frac{\lambda_x}{h_{\text{Sys}}(\boldsymbol{\lambda})} \quad (121)$$

Sie ist die Wahrscheinlichkeit, dass Komponente x zum Ausfall geführt hat, wenn das System ausgefallen ist. Sie gibt damit einen Hinweis, wo man zuerst nach dem Fehler suchen sollte, wenn das System ausgefallen ist. Oder anders gesagt: Je größer die Kritikalitäts-importanz, umso stärkere Auswirkung hat eine relative Verbesserung der Komponente. Sie wird daher manchmal auch Upgrading Importance genannt.

D.9 Importanzen für generische Basis-Ereignisse

Interessant ist auch die Frage, wie sehr sich die System-Eigenschaft Q_{Sys} , F_{Sys} bzw. h_{Sys} ändert, wenn man eine Komponente verändert, die mehrfach verwendet wird. Es wird also nicht die Importanz eines einzelnen Ereignisses betrachtet, sondern die Importanz aller Ereignisse, die sich auf dasselbe generische Basis-Ereignis (GBE) beziehen, einschließlich möglicherweise vorhandener Common-Cause-Faktoren β . Dies ist im folgenden Abschnitt für Beispiel 3 enthalten.

Insbesondere die Importanzen I^{PD} und I^{CRI} sind bezüglich der generischen Basis-Ereignisse wichtig, denn sie geben an, wie sehr sich die Systemgröße absolut bzw. relativ ändert, wenn sich die Basisgröße ändert – etwa weil sie nicht genau bekannt ist.

Für Fehlerbäume berechnet sich die partielle Ableitung nach dem generischen Basisereignis x_{gen} für die System-Nichtverfügbarkeit mit der Näherungsformel (53) zu

$$I_{Q,x_{\text{gen}}}^{\text{PD}} \approx \frac{\partial \sum_{i=1}^{n_{\text{MCS}}} \left(\prod_{j=1}^{m_{\text{Lit},i}} Q_j(t) \right)}{\partial Q_{x_{\text{gen}}}} = \sum_{i=1}^{n_{\text{MCS}}} \begin{cases} 0 & \text{wenn GBE } x \notin \text{MCS}_i \\ a Q_{x_{\text{gen}}}^{a-1} \prod_{j=1, j \neq x_{\text{gen}}}^{m_{\text{Lit},i}} Q_j & \text{wenn GBE } x \in \text{MCS}_i \end{cases} \quad (122)$$

Dabei meint a die Anzahl der Basis-Ereignisse im Minimalschnitt i , die sich auf dasselbe generische Basis-Ereignis x_{gen} beziehen. Der Ausdruck $j \neq x_{\text{gen}}$ meint, dass alle Basis-Ereignisse, die auf das generische Basis-Ereignis x_{gen} verweisen, ignoriert werden sollen, unabhängig von ihrem Index im Minimalschnitt.

Die partielle Ableitung für die System-Ausfallrate berechnet sich basierend auf Minimal-schnitten zu

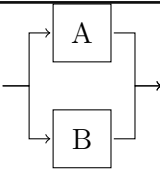
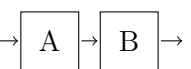
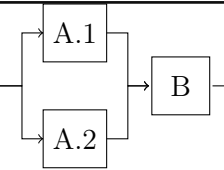
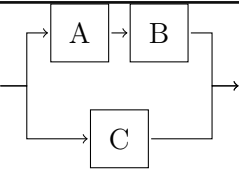
$$I_{h, \text{gen}x}^{\text{PD}} \approx \sum_{i=1}^{n_{\text{MCS}}} \begin{cases} 0 & \text{wenn GBE } x \notin \text{MCS}_i \\ a^2 \lambda_{x_{\text{gen}}}^{a-1} \left(\frac{\partial Q_{x_{\text{gen}}}}{\partial \lambda_{x_{\text{gen}}}} \right)^{a-1} \cdot \prod_{j=1, j \neq x_{\text{gen}}}^{m_{\text{Lit},i}} Q_j & \text{wenn GBE } x \in \text{MCS}_i \\ + a \lambda_{x_{\text{gen}}}^{a-1} \left(\frac{\partial Q_{x_{\text{gen}}}}{\partial \lambda_{x_{\text{gen}}}} \right)^a \cdot \sum_{j=1, j \neq x_{\text{gen}}}^{m_{\text{Lit},i}} \left(h_j \cdot \prod_{k=1, k \neq j, k \neq x}^{m_{\text{Lit},i}} Q_k \right) & \end{cases} \quad (123)$$

D.10 Beispielhafte Importanzen für die System-Nichtverfügbarkeit

Für einige einfache Architekturen sind die Importanzen bezüglich Q_{sys} in der folgenden Tabelle erwähnt. In Beispiel 3 werden zwei gleichartige Ereignisse A.1 und A.2 UND-verknüpft. Somit sind hier auch die in Abschnitt D.9 eingeführten Importanzen bezüglich des zugrundeliegenden generischen Basis-Ereignisses (A) interessant. Diese werden hier mit $I_{Q, \text{gen}A}$ bezeichnet, wohingegen $I_{Q,A}$ jeweils die Importanz des Einzelereignisses A.1 oder A.2 bezeichnet. Ein Common-Cause-Faktor zwischen A.1 und A.2 wurde nicht angenommen ($\beta_A = 0$).

Hinweis: Bei den Berechnungen wurden immer die Mittelwerte $\overline{Q_x}$ verwendet, also etwa $\overline{Q_{A.1}} \cdot \overline{Q_{A.2}}$ anstatt $1/T \cdot \int_0^T Q_{A.1}(t) \cdot Q_{A.2}(t) dt$. Außerdem wurde die Näherungsformel (41) für die Nichtverfügbarkeiten der Einzelereignisse verwendet.

Tabelle 6: Importanzen für Q_{sys} für einfache Architekturen

Wert	Beispiel 1	Beispiel 2	Beispiel 3	Beispiel 4
Block-diagramm				
Minimal-schnitte	{A & B}	{A}, {B}	{A.1 & A.2}, {B}	{A & C}, {B & C}
λ_A	1E-4/h	1E-4/h	1E-4/h	1E-4/h
$T_{\text{Test},A}$	1000 h	1000 h	1000 h	1000 h
λ_B	1E-3/h	1E-3/h	1E-6/h	1E-5/h
$T_{\text{Test},B}$	10 h	10 h	10 h	10 h
λ_C				1E-3/h

Fortsetzung auf nächster Seite

Wert	Beispiel 1	Beispiel 2	Beispiel 3	Beispiel 4
$T_{\text{Test},C}$				50 h
$\overline{Q_A}$	0,050 000	0,050 000	0,050 000	0,050 000
$\overline{Q_B}$	0,005 000	0,005 000	0,000 005	0,000 050
$\overline{Q_C}$				0,025 000
Q_{sys}	$Q_A \cdot Q_B$	$Q_A + (1 - Q_A) \cdot Q_B$	$Q_B + (1 - Q_B) \cdot Q_{A.1} \cdot Q_{A.2}$	$Q_C \cdot (Q_A + (1 - Q_A) \cdot Q_B)$
$\overline{Q_{\text{sys}}}$	0,000 250 00	0,054 750 00	0,002 504 99	0,001 251 19
$Q_{\text{sys}}(Q_A := 0)$	0,000 000 00	0,005 000 00	0,000 005 00	0,000 001 25
$Q_{\text{sys}}(Q_A := 1)$	0,005 000 00	1,000 000 00	0,050 004 75	0,025 000 00
$Q_{\text{sys}}(Q_B := 0)$	0,000 000 00	0,050 000 00	0,002 500 00	0,001 250 00
$Q_{\text{sys}}(Q_B := 1)$	0,050 000 00	1,000 000 00	1,000 000 00	0,025 000 00
$Q_{\text{sys}}(Q_C := 0)$				0,000 000 00
$Q_{\text{sys}}(Q_C := 1)$				0,050 047 50
$Q_{\text{sys}}(Q_{\text{genA}} := 0)$	0,000 000 00	0,005 000 00	0,000 005 00	0,000 001 25
$Q_{\text{sys}}(Q_{\text{genA}} := 1)$	0,005 000 00	1,000 000 00	1,000 000 00	0,025 000 00
I^{PD} über Ableitung:				
$I_{Q,A}^{\text{PD}}$	0,005 000 00	0,995 000 00	0,050 000 00	0,024 998 75
$I_{Q,B}^{\text{PD}}$	0,050 000 00	0,950 000 00	0,997 500 00	0,023 750 00
$I_{Q,C}^{\text{PD}}$				0,050 047 50
$I_{Q,\text{genA}}^{\text{PD}}$	0,005 000 00	0,995 000 00	0,099 999 50	0,024 998 75
I^{PD} über $Q_{\text{sys}}(Q_x := 1) - Q_{\text{sys}}(Q_x := 0)$				
$I_{Q,A}^{\text{PD}}$	0,005 000 00	0,995 000 00	0,049 999 75	0,024 998 75
$I_{Q,B}^{\text{PD}}$	0,050 000 00	0,950 000 00	0,997 500 00	0,023 750 00
$I_{Q,C}^{\text{PD}}$				0,050 047 50
$I_{Q,\text{genA}}^{\text{PD}}$	0,005 000 00	0,995 000 00	0,999 995 00 (f)	0,024 998 75
$I^{\text{RR}} = Q_{\text{sys}} - Q_{\text{sys}}(Q_x := 0)$				
$I_{Q,A}^{\text{RR}}$	0,000 250 00	0,049 750 00	0,002 499 99	0,001 249 94
$I_{Q,B}^{\text{RR}}$	0,000 250 00	0,004 750 00	0,000 004 99	0,000 001 19
$I_{Q,C}^{\text{RR}}$				0,001 251 19
$I_{Q,\text{genA}}^{\text{RR}}$	0,000 250 00	0,049 750 00	0,002 499 99	0,001 249 94
$I^{\text{RRW}} = Q_{\text{sys}}/Q_{\text{sys}}(Q_x := 0) - 1$				
$I_{Q,A}^{\text{RRW}}$	∞	9,950 000 00	499,997 500 00	999,950 000 00
$I_{Q,B}^{\text{RRW}}$	∞	0,095 000 00	0,001 995 00	0,000 950 00
$I_{Q,C}^{\text{RRW}}$				∞
$I_{Q,\text{genA}}^{\text{RRW}}$	∞	9,950 000 00	499,997 500 00	999,950 000 00
I^{FV} über $I^{\text{RR}}/Q_{\text{sys}}$:				
$I_{Q,A}^{\text{FV}}$	1,000 000 00	0,908 675 80	0,998 003 98	0,999 000 95
$I_{Q,B}^{\text{FV}}$	1,000 000 00	0,086 757 99	0,001 991 03	0,000 949 10
$I_{Q,C}^{\text{FV}}$				1,000 000 00

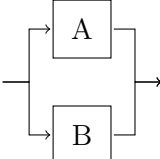
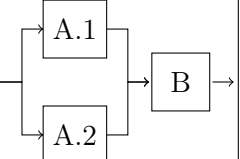
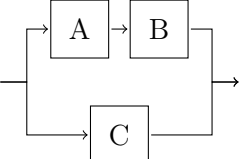
Fortsetzung auf nächster Seite

Wert	Beispiel 1	Beispiel 2	Beispiel 3	Beispiel 4
$I_{Q, \text{genA}}^{\text{FV}}$	1,000 000 00	0,908 675 80	0,998 003 98	0,999 000 95
I^{FV} über $1 - Q_{\text{sys}}(x := 0)/Q_{\text{sys}}$:				
$I_{Q, A}^{\text{FV}}$	1,000 000 00	0,908 675 80	0,998 003 98	0,999 000 95
$I_{Q, B}^{\text{FV}}$	1,000 000 00	0,086 757 99	0,001 991 03	0,000 949 10
$I_{Q, C}^{\text{FV}}$				1,000 000 00
$I_{Q, \text{genA}}^{\text{FV}}$	1,000 000 00	0,908 675 80	0,998 003 98	0,999 000 95
I^{FV} über Minimalschnitte:				
$I_{Q, A}^{\text{FV}}$	1,000 000 00	0,913 242 01	0,998 008 97	0,999 050 90
$I_{Q, B}^{\text{FV}}$	1,000 000 00	0,091 324 20	0,001 996 02	0,000 999 05
$I_{Q, C}^{\text{FV}}$				1,000 000 00
$I_{Q, \text{genA}}^{\text{FV}}$	1,000 000 00	0,913 242 01	0,998 008 97	0,999 050 90
$I^{\text{RA}} = Q_{\text{sys}}(x := 1) - Q_{\text{sys}}$:				
$I_{Q, A}^{\text{RA}}$	0,004 750 00	0,945 250 00	0,047 499 76	0,023 748 81
$I_{Q, B}^{\text{RA}}$	0,049 750 00	0,945 250 00	0,997 495 01	0,023 748 81
$I_{Q, C}^{\text{RA}}$				0,048 796 31
$I_{Q, \text{genA}}^{\text{RA}}$	0,004 750 00	0,945 250 00	0,997 495 01 (f)	0,023 748 81
$I^{\text{RA}} = I^{\text{PD}} - I^{\text{RR}}$:				
$I_{Q, A}^{\text{RA}}$	0,004 750 00	0,945 250 00	0,047 500 01	0,023 748 81
$I_{Q, B}^{\text{RA}}$	0,049 750 00	0,945 250 00	0,997 495 01	0,023 748 81
$I_{Q, C}^{\text{RA}}$				0,048 796 31
$I_{Q, \text{genA}}^{\text{RA}}$	0,004 750 00	0,945 250 00	0,097 499 51	0,023 748 81
$I^{\text{RAW}} = Q_{\text{sys}}(x := 1)/Q_{\text{sys}} - 1$:				
$I_{Q, A}^{\text{RAW}}$	19,000 000 00	17,264 840 18	18,962 075 66	18,981 018 03
$I_{Q, B}^{\text{RAW}}$	199,000 000 00	17,264 840 18	398,203 588 84	18,981 018 03
$I_{Q, C}^{\text{RAW}}$				39,000 000 00
$I_{Q, \text{genA}}^{\text{RAW}}$	19,000 000 00	17,264 840 18	398,203 588 84	18,981 018 03
$I^{\text{CRI}} = I^{\text{PD}} \cdot Q_x/Q_{\text{sys}}$:				
$I_{Q, A}^{\text{CRI}}$	1,000 000 00	0,908 675 80	0,998 008 97	0,999 000 95
$I_{Q, B}^{\text{CRI}}$	1,000 000 00	0,086 757 99	0,001 991 03	0,000 949 10
$I_{Q, C}^{\text{CRI}}$				1,000 000 00
$I_{Q, \text{genA}}^{\text{CRI}}$	1,000 000 00	0,908 675 80	1,996 007 96	0,999 000 95

D.11 Beispielhafte Importanzen für die System-Ausfallrate

Für einige einfache Architekturen sind die Importanzen bezüglich h_{sys} in der folgenden Tabelle erwähnt. In Beispiel 3 werden zwei gleichartige Ereignisse A.1 und A.2 UND-verknüpft. Somit sind hier auch die in Abschnitt D.9 eingeführte Importanzen bezüglich des zugrundeliegenden generischen Basis-Ereignisses interessant. Diese werden hier mit $I_{h, \text{genA}}$ bezeichnet, wohingegen $I_{h, A}$ jeweils die Importanz des Einzelereignisses A.1 oder A.2 bezeichnet.

Tabelle 7: Importanzen für h_{sys} für einfache Architekturen

Wert	Beispiel 1	Beispiel 2	Beispiel 3	Beispiel 4
Block-diagramm				
Minimal-schnitte	{A & B}	{A}, {B}	{A.1 & A.2}, {B}	{A & C}, {B & C}
λ_A $T_{\text{Test},A}$ λ_B $T_{\text{Test},B}$ λ_C $T_{\text{Test},C}$	1E-4/h 1000 h 1E-3/h 10 h 1E-3/h 50 h	1E-4/h 1000 h 1E-3/h 10 h 	1E-4/h 1000 h 1E-6/h 10 h 	1E-4/h 1000 h 1E-5/h 10 h 1E-3/h 50 h
$\overline{Q_A}$ $\overline{Q_B}$ $\overline{Q_C}$	0,050 000 0,005 000 	0,050 000 0,005 000 	0,050 000 0,000 005 	0,050 000 0,000 050 0,025 000
h_{sys}	$h_A \cdot Q_B + h_B \cdot Q_A$	$h_A + h_B$	$h_{A.1} \cdot Q_{A.2} + h_{A.2} \cdot Q_{A.1} + h_B$	$h_A \cdot Q_C + h_C \cdot Q_A + h_B \cdot Q_C + h_C \cdot Q_B$
h_{sys}	$h_A \cdot h_B \cdot (T_A + T_B)/2$	$h_A + h_B$	$h_A \cdot h_A \cdot T_A + h_B$	$h_A \cdot h_C \cdot (T_A + T_C)/2 + h_B \cdot h_C \cdot (T_B + T_C)/2$
h_{sys}	5,0500E-5/h	1,1000E-3/h	1,1000E-5/h	5,2800E-5/h
$h_{\text{sys}}(A:=0)$ $h_{\text{sys}}(B:=0)$ $h_{\text{sys}}(C:=0)$ $h_{\text{sys}}(\text{genA}:=0)$	0,000 000 00/h 0,000 000 00/h 0,000 000 00/h	0,001 000 00/h 0,000 100 00/h 0,001 000 00/h	0,000 001 00/h 0,000 010 00/h 0,000 001 00/h	0,000 000 30/h 0,000 052 50/h 0,000 000 00/h 0,000 000 30/h
IPD über Ableitung:				
$\partial h_{\text{sys}}/\partial \lambda_A$	$h_B \cdot (T_A + T_B)/2$	1	$h_A \cdot T_A$	$h_C \cdot (T_A + T_C)/2$
$\partial h_{\text{sys}}/\partial \lambda_B$	$h_A \cdot (T_A + T_B)/2$	1	1	$h_C \cdot (T_B + T_C)/2$
$\partial h_{\text{sys}}/\partial \lambda_C$				$h_A \cdot (T_A + T_C)/2 + h_B \cdot (T_B + T_C)/2$
$\partial h_{\text{sys}}/\partial \lambda_{\text{genA}}$	$h_B \cdot (T_A + T_B)/2$	1	$2h_A \cdot T_A$	$h_C \cdot (T_A + T_C)/2$
$I_{h,A}^{\text{PD}}$ $I_{h,B}^{\text{PD}}$ $I_{h,C}^{\text{PD}}$ $I_{h,\text{genA}}^{\text{PD}}$	0,505 000 00 0,050 500 00 0,505 000 00	1,000 000 00 1,000 000 00 1,000 000 00	0,100 000 00 1,000 000 00 0,200 000 00	0,525 000 00 0,030 000 00 0,052 800 00 0,525 000 00

Fortsetzung auf nächster Seite

Wert	Beispiel 1	Beispiel 2	Beispiel 3	Beispiel 4
$I^{RR} = h_{sys} - h_{sys}(x := 0)$:				
$I_{h,A}^{RR}$	0,000 050 50/h	0,000 100 00/h	0,000 010 00/h	0,000 052 50/h
$I_{h,B}^{RR}$	0,000 050 50/h	0,001 000 00/h	0,000 001 00/h	0,000 000 30/h
$I_{h,C}^{RR}$				0,000 052 80/h
$I_{h,genA}^{RR}$	0,000 050 50/h	0,000 100 00/h	0,000 010 00/h	0,000 052 50/h
$I^{RRW} = h_{sys}/h_{sys}(x := 0) - 1$:				
$I_{h,A}^{RRW}$	∞	0,100 000 00	10,000 000 00	175,000 000 00
$I_{h,B}^{RRW}$	∞	10,000 000 00	0,100 000 00	0,005 714 29
$I_{h,C}^{RRW}$				∞
$I_{h,genA}^{RRW}$	∞	0,100 000 00	10,000 000 00	175,000 000 00
I^{FV} über I^{RR}/h_{sys} :				
$I_{h,A}^{FV}$	1,000 000 00	0,090 909 09	0,909 090 91	0,994 318 18
$I_{h,B}^{FV}$	1,000 000 00	0,909 090 91	0,090 909 09	0,005 681 82
$I_{h,C}^{FV}$				1,000 000 00
$I_{h,genA}^{FV}$	1,000 000 00	0,090 909 09	0,909 090 91	0,994 318 18
I^{FV} über $1 - h_{sys}(x := 0)/h_{sys}$:				
$I_{h,A}^{FV}$	1,000 000 00	0,090 909 09	0,909 090 91	0,994 318 18
$I_{h,B}^{FV}$	1,000 000 00	0,909 090 91	0,090 909 09	0,005 681 82
$I_{h,C}^{FV}$				1,000 000 00
$I_{h,genA}^{FV}$	1,000 000 00	0,090 909 09	0,909 090 91	0,994 318 18
I^{FV} über Minimalschnitte:				
$I_{h,A}^{FV}$	1,000 000 00	0,090 909 09	0,909 090 91	0,994 318 18
$I_{h,B}^{FV}$	1,000 000 00	0,909 090 91	0,090 909 09	0,005 681 82
$I_{h,C}^{FV}$				1,000 000 00
$I_{h,genA}^{FV}$	1,000 000 00	0,090 909 09	0,909 090 91	0,994 318 18
$I^{CRI} = I^{PD} \cdot h_x/h_{sys}$:				
$I_{h,A}^{CRI}$	1,000 000 00	0,090 909 09	0,909 090 91	0,994 318 18
$I_{h,B}^{CRI}$	1,000 000 00	0,909 090 91	0,090 909 09	0,005 681 82
$I_{h,C}^{CRI}$				1,000 000 00
$I_{h,genA}^{CRI}$	1,000 000 00	0,090 909 09	1,818 181 82	0,994 318 18

E Glossar und Abkürzungsverzeichnis

Tabelle 8: *Begriffe und Abkürzungen*

Begriff	Bedeutung
$\approx$	Etwa gleich, aber sicher kleiner. Meint hier immer: Die Formel ist geringfügig konservativ, aber für korrekt ausgelegte Systeme (Nichtverfügbarkeiten nur etwas größer als null) praktisch gut verwendbar.
Ausfalldichte	$f(t)$, Ableitung der $\rightarrow$ Unzuverlässigkeit.
Ausfallrate	$\rightarrow$ Eintrittsrates
Basis-Ereignis	Ein Ereignis eines $\rightarrow$ Elements.
Bedingung	Bedingungs-Ereignis. Ein Basis-Ereignis, das nur durch eine Wahrscheinlichkeit (typ. Nichtverfügbarkeit) beschrieben wird, nicht durch eine Eintrittsrates.
β	Der Common-Cause-Faktor bezüglich der Eintrittsrates oder der Nichtverfügbarkeit mehrerer nicht völlig unabhängiger Ereignisse.
Eintrittsrates	Die bezüglich der Verfügbarkeit bzw. Zuverlässigkeit bedingte Rate des Eintritts eines Ereignisses. Wenn sie sich auf ein Ereignis bezieht, das einen Ausfall beschreibt, auch Ausfallrate oder Fehlerrates genannt.
Element	Eine beliebige $\rightarrow$ Komponente, menschliches Verhalten oder eine Umgebungsbedingung, die das Verhalten des Systems beeinflusst.
Ereignis	Eine Situation oder ein Zustand, welche(n) ein Element oder ein System annehmen kann.
EUC	Equipment under Control, Begriff aus [IEC 61508], hier immer mit „Prozess“ bezeichnet.
$f(t)$	$\rightarrow$ Ausfalldichte.
$F(t)$	$\rightarrow$ Unzuverlässigkeit.
Fehlerrates	$\rightarrow$ Eintrittsrates
FT	Fault Tree, Fehlerbaum
FTA	Fault tree analysis, Fehlerbaumanalyse
h	$\rightarrow$ Eintrittsrates.
HR	Hazard Rate, tatsächliche oder berechnete (also geschätzte) Eintrittsrates einer Gefährdung
Kante	Die Repräsentation eines Basis-Ereignisses in einem Markov-Modell.
Komponente	Eine technische Einheit, die meist mit unterschiedlichen Fehlermodi ausfallen kann.

Fortsetzung auf nächster Seite

Tabelle 8: *Begriffe und Abkürzungen*

Begriff	Bedeutung
MRT	Mean repair time, mittlere Zeit zwischen Detektion eines Ausfalls und der Reparatur, falls der Prozess (in [IEC 61508]: die „EUC“) im Fall eines erkannten Ausfalls noch weiterbetrieben wird.
MTTD	Mean time to detect, mittlere Zeit bis zum Entdecken des Ausfalls.
MTTF	Mean time to failure, mittlere Zeit bis zum Ausfall und auch mittlere Betriebszeit zwischen zwei Ausfällen.
MTTR	Mean time to restoration. Wiederherstellungszeit. Enthält die $\rightarrow$ MTTD und die $\rightarrow$ MRT.
Nichtverfügbarkeit	Wahrscheinlichkeit $Q(t)$, dass ein $\rightarrow$ Element nicht funktionieren wird, wenn es zur Zeit t angefordert wird.
PFD	Probability of Failure on Demand, Bezeichnung der mittleren $\rightarrow$ Nichtverfügbarkeit $\bar{Q}$ in [IEC 61508].
PFH	Probability of Failure per Hour, Bezeichnung der mittleren bedingten Ausfallhäufigkeit (Ausfallrate) $\rightarrow \bar{h}$ eines Systems in [IEC 61508].
PFTZ	Prozess-Fehlertoleranzzeit, auch Prozess-Sicherheitszeit, Zeit die ein Prozess mit falschen Stellgrößen betrieben werden kann, ohne in einen unsicheren Zustand zu geraten.
PI	Prime Implicant, Prim-Implikant. Äquivalent eines Minimalchnitts im Fall von inkohärenten Fehlerbäumen. Bei kohärenten Fehlerbäumen sind die Prim-Implikanten identisch zu Minimalchnitten.
Q	$\rightarrow$ Nichtverfügbarkeit
Sub-tree	Ein Teil-Fehlerbaum, der mittels Transfer-Gatter in einem höheren Fehlerbaum als $\rightarrow$ Zweig referenziert wird.
System-Lebenszeit	Die geplante Einsatzzeit des betrachteten Systems. Wird benötigt, wenn die interessierende Größe nicht konstant und nicht periodisch ist, insbesondere also zur Bestimmung der $\rightarrow$ Unzuverlässigkeit, oder wenn Basis-Ereignisse des Modells „Nicht-Wiederherstellbar“ im Fehlerbaum oder Markov-Modell enthalten sind.
THR	Tolerable Hazard Rate, akzeptierbare Gefährdungsrate, Zielgröße für Sicherheitsfunktionen mit kontinuierlicher oder häufiger Anforderung.
TPFD	Tolerable Probability of Failure on Demand, akzeptierbare Nichtverfügbarkeit, Zielgröße für Sicherheitsfunktionen mit seltener Anforderung.

Fortsetzung auf nächster Seite

Tabelle 8: *Begriffe und Abkürzungen*

Begriff	Bedeutung
TPFH	Tolerable Probability of Failure per Hour, akzeptierbare Ausfallrate, Zielgröße für Sicherheitsfunktionen mit kontinuierlicher oder häufiger Anforderung. Mathematisch und auf oberster Ebene auch logisch identisch zur $\rightarrow$ THR.
Unzuverlässigkeit	Wahrscheinlichkeit $F(t_1, t_2)$ dass ein $\rightarrow$ Element während der Zeitspanne $t_1 \dots t_2$ ausfällt.
$w(t)$	Eintritts-Häufigkeit im Fall von test- und/oder reparierbaren Ereignissen. Im Gegensatz zu $h(t)$ ist $w(t)$ nur bezüglich des letzten Tests oder Tauschs bedingt, nicht bezüglich der Verfügbarkeit des Systems. $w(t)$ wird daher oft auch als <u>unbedingte Eintrittsfrequenz</u> bezeichnet. Im Gegensatz zu $f(t)$ ist $w(t)$ bezüglich des letzten Tests oder Tauschs bedingt, sein Integral über die Zeit kann also größer als 1 werden.
Zweig	Der Teil eines Fehlerbaums, der sich unterhalb eines bestimmten Gatters befindet einschließlich des Gatters selbst. Spezialfall: Auch ein Basis-Ereignis ist ein Zweig.

F Literaturverzeichnis

Literatur

- [EN 50126] *Railway Applications – The Specification and Demonstration of Reliability, Availability, Maintainability and Safety (RAMS)*, CENELEC (2017)
- [EN 50129] *Railway Applications – Communication, signalling and processing systems – Safety related electronic systems for signalling*, CENELEC (2017)
- [EN 61025] *Fault tree analysis (FTA)*, IEC (2006)
- [EN 61165] *Application of Markov techniques*, IEC (2006)
- [IEC 61508] *Functional safety of electrical/electronic/programmable electronic safety-related systems*, IEC (2010)
- [IEC 61508-6] *Functional safety of electrical/electronic/programmable electronic safety-related systems, Part 6: Guidelines on the application of IEC 61508-2 and IEC 61508-3*, IEC (2010)
- [ISO 13849] *Safety of machinery – Safety-related parts of control systems*, ISO (2015)
- [ISO 26262] *Road vehicles – Functional safety*, ISO (2011)
- [NUREG] *The Fault Tree Handbook*, NUREG-0492, US Nuclear Regulatory Commission (1981)
- [U. Weber] *Theory and Execution of Fault Tree Analyses for Railway Applications*, TÜV SÜD Rail GmbH (2010)
- [ASTRA TM] *ASTRA 3.x: Theoretical Manual*, Sergio Contini and Vaidas Matuzas, EUR 25052 EN 2011
- [PSA IMP] *An overview of PSA importance measures*, M. van der Borst, H. Schoonakker, Reliability Engineering and System Safety 72 (2001) 241-245

G Änderungsverzeichnis

Tabelle 9: Änderungsverzeichnis

Datum	Änderung
06.03.2019	Formel (49) korrigiert In Beispiel 6.2 Referenz auf Formel (49) geändert Abbildung 28 geändert
26.03.2019	Einleitung zu Abschnitt 3 geändert, Überschrift 3.1 ergänzt, Beispiel 3.1 ergänzt. Einleitung zu Abschnitt 6 geändert. Abschnitt 4.8 entfernt, da redundant zu 6. Abschnitt 4.9 nach 6.3 verschoben. Anhang A geändert.
17.06.2019	Anhang D ergänzt. Abschnitte 5, 6, 7: Untere Indexgrenze in einigen Summen und Produkten von 0 auf 1 korrigiert.
25.02.2020	Formeln (44), (45), (48) und (49) sowie entsprechende Verweise in Anhang A korrigiert
29.11.2025	Formel (62) durch Formel (28) ersetzt