

Adressen:

Herausgeber: Bundesamt für Konjunkturfragen (BfK)
Belpstrasse 53
3003 Bern
Tel.: 031/61 21 39
Fax: 031/61 20 57

Geschäftsstelle: RAVEL
c/o Amstein+Walthert AG
Leutschenbachstrasse 45
8050 Zürich
Tel.: 01/30591 11
Fax: 01/305 92 14

Ressortleiter: Jean Marc Chuard
Enerconom AG
Hochfeldstrasse 34
3012 Bern
Tel.: 031/23 97 23
Fax: 031/24 63 53

Autor: Prof. Dr. Alessandro Birolini
Professur für Zuverlässigkeitstechnik
ETH Zentrum
8092 Zürich
Tel.: 01/256 51 48

Diese Studie gehört zu einer Reihe von Untersuchungen, welche zu Handen des Impulsprogrammes RAVEL von Dritten erarbeitet wurde. Das Bundesamt für Konjunkturfragen und die von ihm eingesetzte Programmleitung geben die vorliegende Studie zur Veröffentlichung frei. Die inhaltliche Verantwortung liegt bei den Autoren und der zuständigen Ressortleitung.

Copyright Bundesamt für Konjunkturfragen
3003 Bern, Mai 1992

Auszugweiser Nachdruck unter Quellenangabe erlaubt. Zu beziehen bei der Eidg. Drucksachen- und Materialzentrale, Bern (Best. Nr. 724.397.13.56 d)

Form. **724.397.13.56 d 5.92 1000**

RAVEL - Materialien zu RAVEL

Zuverlässigkeit und Energieverbrauch von elektronischen Geräten und Systemen



Alessandro Birolini

Impulsprogramm RAVEL
RAVEL - Materialien zu RAVEL

Bundesamt für Konjunkturfragen

Vorwort

Zuverlässigkeit ist eine Eigenschaft von Geräten und Systemen, deren Bedeutung immer mehr zunimmt. Ihre Sicherstellung erfolgt primär in der Entwicklungsphase mit der Durchführung bestimmter Engineerings- und Managementaktivitäten, zu welchen die Festlegung der Ziele, die Wahl und Qualifikation von Bauteilen und Stoffen, die Durchführung von Analysen und Prüfungen, das Konfigurationsmanagement sowie die Hebung der Zuverlässigkeit in der Fertigungs- und Nutzungsphase gehören. Zusammen mit der Instandhaltbarkeit bestimmt die Zuverlässigkeit auch die Verfügbarkeit eines Geräts oder Systems. Hohe Zuverlässigkeits- oder Verfügbarkeitsforderungen lassen sich oft nur mit Hilfe von Redundanz erreichen. Diese kann prinzipiell als heisse (parallel), warme (leicht belastete) oder kalte (Standby) Redundanz realisiert werden, was eine direkte Auswirkung auf den Energieverbrauch hat. Auch stellt sich die Frage des Einflusses auf die Zuverlässigkeit vom Ein- und Ausschalten eines Geräts oder Systems. Zu hohe oder zu pauschal formulierte Zuverlässigkeitsforderungen können zur Verdoppelung oder Verdreifachung grosser Teile einer Anlage (z.B. EDV inkl. Klima) und damit zu einem Energieverbrauch führen, der nicht unbedingt vertretbar bleibt. Redundanz kann oft gezielt auf Modul- oder Geräteebene eingesetzt werden; in manchen Fällen kann man durch geeignete Ein- und Abschaltstrategien sogar auf sie verzichten.

Diese Monographie zeigt den Zusammenhang zwischen Redundanz, Zuverlässigkeit, Verfügbarkeit und Energieverbrauch bei elektronischen Geräten und Systemen auf, allgemein in den Kapiteln 1, 3, 5 und an konkreten Beispielen im Kapitel 4. Sie ist speziell für Projektleiter und Mitarbeiter des mittleren und oberen Kaders als Grundlage bei der Festlegung von Zuverlässigkeits- und Verfügbarkeitsforderungen in Zusammenhang mit Überlegungen zum Energieverbrauch gedacht. Viele Abschnitte enthalten auch Detailinformationen für Entwicklungsingenieure oder Mitarbeiter der Qualitäts- und Zuverlässigkeitssicherung. Das Rohmaterial für die Beispiele der Abschnitte 4.2 und 4.3 wurde von den Herren G.A. Zardini und L. Strozzi geliefert. Herr Dr. F. Bonzanigo hat für die graphische Darstellung der wichtigsten Resultate gesorgt. Die Monographie wurde mit den Herren Dr. B. Aebischer, J.-M. Chuard und Dr. R. Walthert bereinigt. Allen diesen Herren gilt den aufrichtigen Dank für die gute Zusammenarbeit.

Zürich, März 1992

A. Birolini

Inhaltsverzeichnis

1.	Einleitung, Grundbegriffe	
1. 1	Einleitung	1.1
1.2	Grundbegriffe	1.1-1.8
2.	Zuverlässigkeit von elektronischen Geräten und Systemen	
2.1	Vorausgesagte Zuverlässigkeit	2.1
2.2	Analyse der Art und Auswirkung von Ausfällen	2.19-2.20
3.	Zuverlässigkeit und Verfügbarkeit reparierbarer Geräte und Systeme	
3.1	Einzelelement	3.2
3.2	Geräte und Systeme ohne Redundanz	3.5
3.3	Redundanz 1 aus 2	3.7
3.4	Redundanz k aus n	3.11
3.5	Einfache Serie-Parallelstrukturen	3.14
3.6	Näherungsformeln für grosse reparierbare Serien-/Parallelstrukturen	3.16-3.21
4.	Einfluss von Redundanz und Ein- Ausschalten auf Zuverlässigkeit und Energieverbrauch von Geräten und Systemen	
4.1	Allgemeine Betrachtungen	4.1
4.2	Multicomputer-System	4.2
4.3	Unterbrechungsfreie Stromversorgung	4.4-4.6
5.	Festlegung und Durchsetzung von Zuverlässigkeitsforderungen	
5.1	Kundenforderungen	5.1
5.2	Wirtschaftlichkeitsbetrachtungen	5.3
5.3	Energieverbrauch	5.5
5.4	Festlegung von Zuverlässigkeitsforderungen	5.5
5.5	Durchsetzung von Zuverlässigkeitsforderungen	5.7-5.23
Anhang A I: Definitionen und Begriffserklärungen		A1.1-A1.12
Anhang A2: Mathematische Grundlagen		A2. 1 -A2.20
Literaturverzeichnis		L 1

1 Einleitung, Grundbegriffe

1.1 Einleitung

Von einem modernen, leistungsfähigen Gerät oder System erwartet man heutzutage eine hohe Zuverlässigkeit und Verfügbarkeit. Diese Eigenschaften sind in der Entwicklungsphase in ein Gerät hineinzuentwickeln. Hohe Zuverlässigkeits- oder Verfügbarkeitsforderungen können aber oft nur mit Hilfe von Redundanz erreicht werden. Diese kann als heisse (parallel), warme (leicht belastete) oder kalte (standby) Redundanz ausgelegt werden und damit einen grossen Einfluss auf den Energieverbrauch eines Geräts oder Systems haben. Um Redundanz einzuführen, bleibt für den Anwender praktisch nur die Verdoppelung (oder Verdreifachung) grosser Teile seiner Anlage (z.B. EDV inkl. Klima) offen. Für den Hersteller ist hingegen oft möglich, Redundanz auf Modulebene einzusetzen und in manchen Fällen durch geeignete Ein- bzw. Abschaltstrategien und entsprechende Schaltungsauslegung auf Redundanz (mit kleinen Einbussen) zu verzichten. Nach Einführung der Grundbegriffe werden in Kapitel 2 die wichtigsten Einflussfaktoren und Berechnungsmethoden der Zuverlässigkeit elektronischer Geräte und Systeme (ohne Berücksichtigung der Reparatur) kurz dargelegt, insbesondere wird auf die Berechnung der Ausfallrate (vorausgesagte Zuverlässigkeit) und auf die Auswirkung von Ausfällen (Ausfallanalyse) eingegangen. Der Einfluss der Instandsetzbarkeit auf die Zuverlässigkeit und Verfügbarkeit von Geräten und Systemen mit Redundanz wird in Kapitel 3 gezeigt. Analysiert werden die wichtigsten redundanten Strukturen, so weit wie möglich mit graphischer Darlegung der Resultate. In Kapitel 4 werden zwei konkrete Beispiele untersucht, ein Multi-computer-System und eine unterbrechungsfreie Stromversorgung. Die Wechselwirkung zwischen Zuverlässigkeit, Verfügbarkeit, Redundanz und Energieverbrauch wird gezeigt. Auf die Problematik der Festlegung und Durchsetzung von Zuverlässigkeitsforderungen wird in Kapitel 5 eingegangen. Zur laufenden Beurteilung und Überprüfung der Forderungen eignen sich insbesondere die Entwurfsüberprüfungen (Design Reviews); umfassende Fragenkataloge zur Erstellung von Checklisten für Entwurfsüberprüfungen werden angegeben. Die Anhänge A I und A2 schliessen diese Monographie mit einer umfassenden Liste von Definitionen und einem Auszug aus der Wahrscheinlichkeitsrechnung.

1.2 Grundbegriffe

In diesem Kapitel werden die wichtigsten Begriffe zur Qualitäts- und Zuverlässigkeitssicherung von Geräten eingeführt (vgl. Anhang A 1 für eine ausführliche Darlegung).

1.1

1.2.1 Zuverlässigkeit

Die Zuverlässigkeit ist ein Mass für die Fähigkeit einer Betrachtungseinheit, funktionstüchtig zu bleiben. Sie wird mit R bezeichnet und durch die Wahrscheinlichkeit ausgedrückt, dass die geforderte Funktion unter vorgegebenen Arbeitsbedingungen während einer festgelegten Zeitdauer T ausfallfrei ausgeführt wird. Wie aus der Definition hervorgeht, gibt also die Zuverlässigkeit die Wahrscheinlichkeit an, dass in der Zeitspanne T kein Ausfall auftreten wird, der die Erfüllung der geforderten Funktion (auf Ebene Betrachtungseinheit) beeinträchtigt. Dies bedeutet aber nicht, dass redundante Teile nicht ausfallen dürfen. Solche Teile können ausfallen und (ohne Betriebsunterbrechung auf Ebene Betrachtungseinheit) instandgesetzt werden. Man stellt auch fest, dass mit einer numerischen Angabe der Zuverlässigkeit (z. B. $R = 0.9$) stets auch die geforderte Funktion, die Arbeitsbedingungen und die Missionsdauer T definiert werden müssen. Ebenso muss festgelegt werden, ob zu Beginn der Mission die Betrachtungseinheit als neuwertig angesehen werden kann.

Unter einer Betrachtungseinheit versteht man eine Anordnung beliebiger Komplexität (Stoff, Bauteil, Unterbaugruppe, Baugruppe, Gerät, Anlage, System), welche für Untersuchungen oder Analysen als Einheit interpretiert wird. Dabei kann es sich um eine Funktions- oder Konstruktionseinheit handeln. Zur Vereinfachung werden im folgenden komplexe Betrachtungseinheiten als System bezeichnet (Indizes S in den Formeln).

Die geforderte Funktion spezifiziert die Aufgabe der Betrachtungseinheit. Für gegebene Eingänge dürfen beispielsweise die Ausgänge vorgeschriebene Toleranzbänder nicht verlassen. Die Festlegung der geforderten Funktion ist der Ausgangspunkt jeder Zuverlässigkeitsanalyse, weil damit auch der Ausfall definiert wird. Dies kann allerdings für komplexe Betrachtungseinheiten kompliziert oder zumindest aufwendig werden.

Die Arbeitsbedingungen haben einen direkten Einfluss auf die Zuverlässigkeit und müssen genau spezifiziert werden. Die Erfahrung zeigt z. B., dass sich die Ausfallrate elektronischer Bauteile verdoppelt, wenn die Umgebungstemperatur um 10 bis 20°C erhöht wird. Geforderte Funktion und Arbeitsbedingungen können auch zeitabhängig sein. In solchen Fällen ist ein Anforderungsprofil zu definieren, auf welches alle Zuverlässigkeitsangaben bezogen werden. Ein repräsentatives Anforderungsprofil und die entsprechenden Zuverlässigkeitsziele sind im Pflichtenheft festzulegen. Oft interessiert in den praktischen Anwendungen der zeitliche Verlauf der Zuverlässigkeit R , wenn die Missionsdauer variiert wird, d.h. die Zuverlässigkeitsfunktion $R(t)$.

1.2.2 Ausfall

Ein Ausfall tritt auf, wenn eine Betrachtungseinheit aufhört, ihre geforderte Funktion auszuführen. Die Betriebszeit kann dabei sehr kurz gewesen sein, denn Ausfälle können auch durch transiente Vorgänge beim Einschalten verursacht werden. Bei der Beurteilung eines Ausfalls wird davon ausgegangen, dass zum Beanspruchungsbeginn die Betrachtungseinheit fehlerfrei war. Die Bewertung erfolgt dann unter folgenden Gesichtspunkten:

1. Art: Es wird unterschieden zwischen Sprungausfall (Kurzschluss, Unterbrechung, Funktionsfehler für elektronische Bauteile bzw. Fließen, Sprödbbruch, Fressen usw. für mechanische Bauteile), Driftausfall und intermittierendem Ausfall.

2. Ursache: Eine übliche Einteilung unterscheidet zwischen Anwendungsfehlerausfällen (Fehler in der Entwicklung, Fertigung oder Bedienung), inhärenten Ausfällen, Verschleissausfällen, Primärausfällen und Folgeausfällen.
3. Auswirkung: Abhängig davon, ob man sich auf die direkt betroffene oder auf eine übergeordnete Betrachtungseinheit bezieht, wird unterschieden zwischen keiner Auswirkung, Teilausfall, Vollaussfall und überkritischem Ausfall. Bei überkritischen Ausfällen ist die Sicherheit nicht mehr gewährleistet, der Einfluss auf die geforderte Funktion kann dabei verschieden gross sein.

1.2.3 Ausfallrate

Die Ausfallrate spielt in Zuverlässigkeitsanalysen eine wichtige Rolle. Sie wird im Anhang A2 (Gl. (A2.20) definiert und in diesem Abschnitt auf eine mehr heuristische Weise eingeführt. Zur Zeit $t = 0$ es seien N statistisch identische (unabhängige) Betrachtungseinheiten unter den gleichen Bedingungen in Betrieb gesetzt worden. Zur Zeit t seien $n(t)$ Betrachtungseinheiten noch nicht ausgefallen, vgl. Bild 1.1.

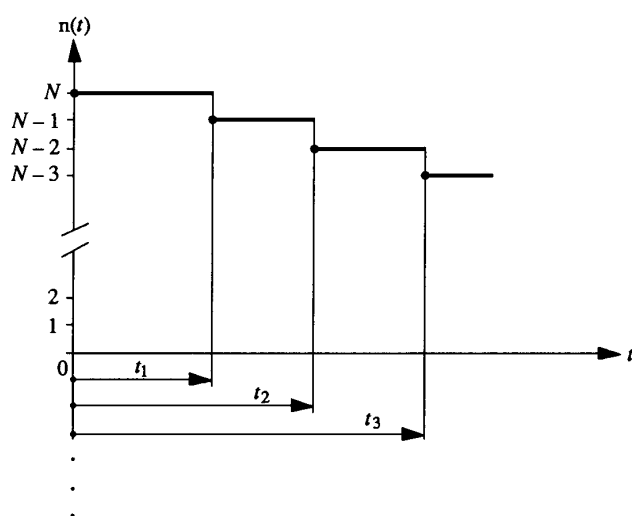


Bild 1.1 Anzahl $n(t)$ der Betrachtungseinheiten, welche zur Zeit t noch nicht ausgefallen sind

$t_1, \dots, t_N$ sind die beobachteten ausfallfreien Arbeitszeiten der N Betrachtungseinheiten. Gemäss obiger Voraussetzung sind sie unabhängige Realisierungen der als Zufallsgrösse τ betrachteten ausfallfreien Arbeitszeit der Betrachtungseinheit. Die Funktion

$$\hat{R}(t) = \frac{n(t)}{N} \quad (1.1)$$

ist die empirische Zuverlässigkeitsfunktion. Sie konvergiert für $N \rightarrow \infty$ gegen die (wahre) Zuverlässigkeitsfunktion $R(t)$. Als empirische Ausfallrate wird die Grösse

$$\hat{\lambda}(t) = \frac{n(t) - n(t + \delta t)}{n(t) \delta t} \quad (1.2)$$

definiert. $\hat{\lambda}(t) \delta t$ ist gleich dem Verhältnis der Anzahl Ausfälle im Intervall $(t, t+\delta t]$ zur Anzahl Betrachtungseinheiten, die zur Zeit t noch nicht ausgefallen sind. Mit Hilfe der Gl. (1.1) folgt aus Gl. (1.2)

$$\hat{\lambda}(t) = \frac{\hat{R}(t) - \hat{R}(t + \delta t)}{\delta t \hat{R}(t)}. \quad (1.3)$$

Für $N \rightarrow \infty$ und $\delta t \rightarrow 0$ konvergiert $\hat{\lambda}(t)$ gegen die *Ausfallrate*

$$\lambda(t) = -\frac{1}{R(t)} \cdot \frac{dR(t)}{dt}. \quad (1.4)$$

Gleichung (1.4) stellt einen wichtigen Zusammenhang dar. Sie zeigt, dass die Ausfallrate $\lambda(t)$ die *Zuverlässigkeitsfunktion* $R(t)$ vollständig bestimmt. Mit $R(0) = 1$ folgt aus Gl. (1.4)

$$R(t) = e^{-\int_0^t \lambda(x) dx}. \quad (1.5)$$

In vielen praktischen Anwendungen trifft der Fall zu, in welchem die Ausfallrate als näherungsweise konstant angenommen werden kann. Hier wird

$$\lambda(t) = \lambda \quad (1.6)$$

gesetzt. Aus Gl. (1.5) folgt dann

$$R(t) = e^{-\lambda t}. \quad (1.7)$$

Der Mittelwert der ausfallfreien Arbeitszeit τ wird allgemein mit *MTTF* (Mean Time To Failure) bezeichnet und gemäss $MTTF = E[\tau] = \int_0^\infty R(t) dt$ bestimmt. Für den Fall einer konstanten Ausfallrate λ gilt speziell $E[\tau] = \int_0^\infty e^{-\lambda t} dt = 1/\lambda$. Es ist üblich

$$MTBF = \frac{1}{\lambda} \quad (1.8)$$

zu setzen, wobei *MTBF* für Mean Time Between Failures steht (vgl. Anhang A1 für eine Gegenüberstellung der Begriffe *MTBF* und *MTTF*).

Zahlreiche Versuche zeigen, dass im allgemeinen Fall die Ausfallrate einer (unendlich grossen) Grundgesamtheit statistisch identischer Betrachtungseinheiten typisch den Verlauf gemäss Bild 1.2 aufweist. Dieser setzt sich zusammen aus drei charakteristischen Phasen:

1. *Phase der Frühausfälle* (Frühausfallphase): $\lambda(t)$ nimmt (in der Regel) rasch ab; Ausfälle in dieser Phase lassen sich auf Materialschwächen, Qualitätsschwankungen in der Fertigung oder *Anwendungsfehler* (Dimensionierung, Prüfung, Bedienung) zurückführen.
2. *Phase der Ausfälle mit konstanter Ausfallrate*: $\lambda(t)$ ist näherungsweise konstant und gleich λ ; in dieser Phase treten die Ausfälle meistens plötzlich und rein zufällig auf.
3. *Phase der Verschleissausfälle* (Spätausfallphase): $\lambda(t)$ steigt mit zunehmender Betriebszeit immer schneller an; Ausfälle in dieser Phase sind auf *Alterung*, *Abnützung*, *Ermüdung* usw. zurückzuführen.

Die Dauer der einzelnen Phasen kann in der Praxis stark variieren. Bei elektronischen Röhren und elektromechanischen Bauteilen ist beispielsweise eine ausgeprägte Phase der Verschleiss-

ausfälle feststellbar, während eine solche bei den meisten Halbleiterbauteilen in weiten Grenzen nicht auftritt. Die Phase der Frühausfälle kann je nach der Komplexität der Betrachtungseinheiten und der Reife des Herstellungsprozesses praktisch nicht vorhanden sein oder bis zu einigen wenigen tausend Betriebsstunden dauern. Ganz allgemein hängt die Ausfallrate stark von den Arbeitsbedingungen ab (vgl. Bild 2.2 und Bild 2.3)

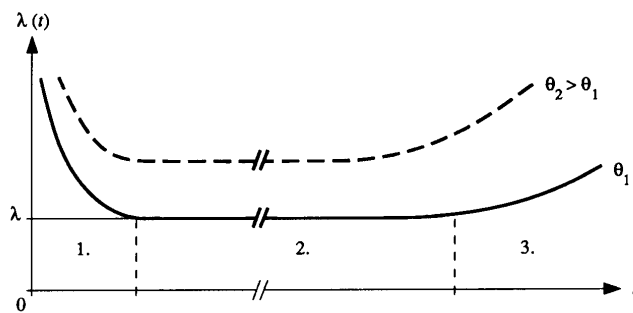


Bild 1.2 Typischer Verlauf der Ausfallrate einer (unendlich grossen) Grundgesamtheit statistisch identischer Betrachtungseinheiten (gestrichelt angegeben ist die prinzipielle Verschiebung der Kurve, wenn die Belastung, z. B. Umgebungstemperatur θ_A für elektronische Bauteile, erhöht wird)

1.2.4 Instandhaltbarkeit

Für viele Geräte und Systeme ist eine Instandhaltung möglich. Unter Instandhaltung versteht man alle Aktivitäten zur Erhaltung oder Wiederherstellung des Sollzustands einer Betrachtungseinheit. Man unterscheidet zwischen Wartung, d. h. periodischen Arbeiten zur Kontrolle des Funktionszustands und zur Entdeckung von verborgenen Ausfällen sowie zur Vermeidung von Drift- bzw. Verschleissausfällen, und Instandsetzung, d. h. Aktivitäten zur Lokalisierung und Behebung eines Ausfalls. Es ist üblich, Instandsetzung auch als Reparatur zu bezeichnen (zwischen Instandsetzung und Reparatur wird im folgenden nicht mehr unterschieden).

Die Instandhaltbarkeit ist hingegen ein Mass für die Fähigkeit einer Betrachtungseinheit, funktionstüchtig gehalten werden zu können. Sie wird ausgedrückt durch die Wahrscheinlichkeit, dass der Zeitaufwand für eine Reparatur bzw. für eine Wartung kleiner als eine gegebene Zeitspanne ist, wenn die Instandhaltung unter definierten materiellen und personellen Bedingungen erfolgt. Der Mittelwert der Reparaturzeit wird mit M7TR (Mean Time To Repair) und jener der Zeit für eine Wartung mit MITPM (Mean Time To Preventive Maintenance) bezeichnet.

Die Instandhaltbarkeit muss in ein Gerät oder System durch ein Instandhaltungskonzept während der Entwicklungsphase hineinentwickelt werden. Wegen des direkten Einflusses auf die Verfügbarkeit, der starken Zunahme der Instandhaltungskosten und der Schwierigkeiten, qualifiziertes Personal für die Durchführung von Instandhaltungsarbeiten zu finden, wird der Instandhaltbarkeit eine immer grössere Bedeutung zugemessen. Wie die Erfahrung zeigt, hängt aber die im Betrieb erreichte Instandhaltbarkeit auch von der Installation der Betrachtungseinheit sowie von der Organisation, Ausrüstung und Ausbildung des Instandhaltungspersonals, d. h. ganz allgemein von der logistischen Unterstützung ab.

1.2.5 Logistische Unterstützung

Mit logistischer Unterstützung (Logistik) werden alle Massnahmen und Aktivitäten bezeichnet, die mit dem Ziel ausgeführt werden, eine wirksame und wirtschaftliche Verwendung einer Betrachtungseinheit während der Nutzungsphase zu ermöglichen. Die logistische Unterstützung beginnt in der Entwicklungsphase (Teil des Instandhaltungskonzepts) und schliesst den Kundendienst ein.

1.2.6 Verfügbarkeit

Die Verfügbarkeit (Punkt-Verfügbarkeit) ist ein Mass für die Fähigkeit einer Betrachtungseinheit, zu einem gegebenen Zeitpunkt funktionstüchtig zu sein. Sie wird als Wahrscheinlichkeit ausgedrückt, dass die Betrachtungseinheit zu einem bestimmten Zeitpunkt die geforderte Funktion unter vorgegebenen Arbeitsbedingungen ausführt, unabhängig davon, ob sie bis zu diesem Zeitpunkt ausgefallen oder noch nicht ausgefallen ist (dies im Gegensatz zur Zuverlässigkeit, für welche der Zeitpunkt des ersten Ausfalls auf Ebene Betrachtungseinheit massgebend ist). Ihre Berechnung ist in der Regel kompliziert, weil neben der Zuverlässigkeit und Instandhaltbarkeit auch die logistische Unterstützung und die menschlichen Faktoren zu berücksichtigen sind. Oft geht man von der Annahme idealer logistischer Unterstützung und menschlicher Faktoren aus, so dass die Verfügbarkeit nur noch eine Funktion der Zuverlässigkeit und Instandhaltbarkeit wird. Im Falle eines Dauerbetriebs (die Betrachtungseinheit pendelt zwischen Arbeits- und Reparaturzustand) konvergiert die Punkt-Verfügbarkeit schnell gegen den Ausdruck $MTBF / (MTBF + MTTR)$. Dieser Wert ist auch gleich dem asymptotischen und stationären Wert der durchschnittlichen Verfügbarkeit (Tab. 3. 1). Je nach Anwendung können andere Verfügbarkeitsarten definiert werden.

1.2.7 Sicherheit

Die Sicherheit ist ein Mass für die Fähigkeit einer Betrachtungseinheit, weder Menschen, Sachen noch Umwelt zu gefährden. Ihre Untersuchung muss unter zwei Gesichtspunkten erfolgen: Sicherheit, wenn die Betrachtungseinheit korrekt funktioniert und betrieben wird, und Sicherheit, wenn die Betrachtungseinheit oder ein Teil davon ausgefallen ist. Der erste Aspekt wird durch die Unfallverhütung abgedeckt, die vielfach durch gesetzliche Vorschriften geregelt ist. Der zweite Aspekt ist Gegenstand der technischen Sicherheit und wird mit den Methoden der Zuverlässigkeitstheorie untersucht. Dabei müssen auch die Einwirkungen äusserer Einflüsse (Katastrophe, Sabotage usw.) berücksichtigt werden. Prinzipiell soll jedoch zwischen technischer Sicherheit und Zuverlässigkeit unterschieden werden. Während die Sicherheitstheorie Massnahmen untersucht, die es gestatten, bei einem Ausfall die Betrachtungseinheit in einen sicheren Zustand zu bringen (fail safe), untersucht die Zuverlässigkeitstheorie Massnahmen, um ganz allgemein die Anzahl Ausfälle zu vermindern. Die Sicherheit einer Betrachtungseinheit bestimmt weitgehend das Auftreten von Produkthaftpflichtfällen.

1.2.8 Kosten- bzw. Systemwirksamkeit

Alle bisher vorgestellten Begriffe sind miteinander verkoppelt. Ihr Zusammenhang lässt sich am besten anhand des Begriffs Kostenwirksamkeit zeigen und ist im Bild 1.3 dargelegt. Unter Kostenwirksamkeit versteht man ein Mass für die Fähigkeit einer Betrachtungseinheit, die geforderte Funktion mit dem bestmöglichen Verhältnis von Nutzen zu Lebenslaufkosten zu erfüllen. Es ist üblich, Kostenwirksamkeit auch als Systemwirksamkeit zu bezeichnen. Als Lebenslaufkosten wird die Summe der Anschaffungs-, Betriebs-, Instandhaltungs- und Ausscheidungskosten definiert. Diese Kosten hängen stark von der Zuverlässigkeit und Instandhaltbarkeit ab. Sie werden für grosse technische Systeme mit hohen Zuverlässigkeitsanforderungen oft bis zu 90% in der Definitions- und Entwicklungsphase bestimmt. Für solche Systeme liegt die Anwendungsdauer in der Regel über 10 Jahren und die Anschaffungskosten sollten als Richtwert etwa 50% der Lebenslaufkosten ausmachen.

Aus Bild 1.3 ist die zentrale Funktion der Qualitätssicherung ersichtlich. Sie fasst alle Sicherungsaktivitäten zusammen und stützt sich im wesentlichen auf folgende Aspekte:

1. Konfigurationsmanagement: Verfahren zur Festlegung, Beschreibung, Prüfung und Genehmigung der Konfiguration einer Betrachtungseinheit sowie zu ihrer Steuerung und Überwachung bei Änderungen und Modifikationen (es ist üblich, das Konfigurationsmanagement in Beschreibung, Überprüfung, Steuerung und Überwachung der Konfiguration zu unterteilen, vgl. Abschnitt 5.5.4).
2. Qualitätsprüfung: Planung, Durchführung und Auswertung aller notwendigen Prüfungen, um sicherzustellen, dass die Betrachtungseinheit den gestellten Anforderungen genügt (zu diesen Prüfungen gehören Eingangsprüfungen, Qualifikationsprüfungen, Zwischenprüfungen und Endprüfungen samt Zuverlässigkeits-, Instandhaltbarkeits- und Sicherheitsprüfungen).
3. Qualitätssteuerung in der Fertigung: Steuerung der Fertigungsprozesse und -abläufe mit dem Ziel, das geforderte Qualitätsniveau einer Betrachtungseinheit sicherzustellen.
4. Qualitätsdatensystem: Ein oft computerunterstütztes System zur raschen und wirksamen Erfassung, Analyse und Korrektur aller Defekte und Ausfälle, die während der Herstellung, Prüfung und Nutzung einer Betrachtungseinheit auftreten sowie zur Verdichtung, Speicherung, Auswertung und Rückkopplung der entsprechenden Qualitäts- und Zuverlässigkeitsdaten zu allen Linienstellen.

Das System (und damit auch der Aufwand) für die Sicherstellung der Qualität und Zuverlässigkeit muss der Bedeutung des hergestellten Geräts oder Systems angepasst und in einem entsprechenden Qualitätssicherungs-Handbuch dokumentiert werden. Ein solches Handbuch beschreibt das vorhandene System und gibt an, wie in der betreffenden Firma die Aufgaben standardmässig gelöst werden (Verfahren). Es soll in Zusammenarbeit mit den Linienstellen hergestellt werden, Kompetenzen bzw. Verantwortungen regeln und dynamisch bleiben. Es ist von Vorteil das Qualitäts- und Zuverlässigkeits-Handbuch in drei Teile aufzugliedern: Allgemeine Richtlinien und Verfahren (wird den Kunden abgegeben), Detailrichtlinien und -verfahren (kann von den Kunden eingesehen werden) und Motivation und Schulung.

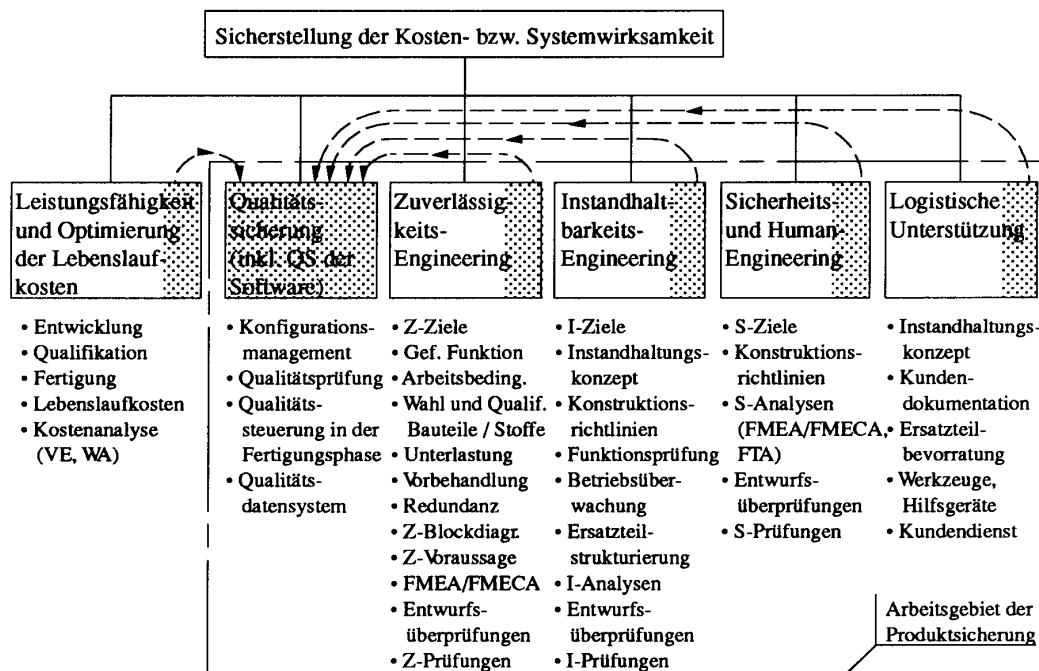
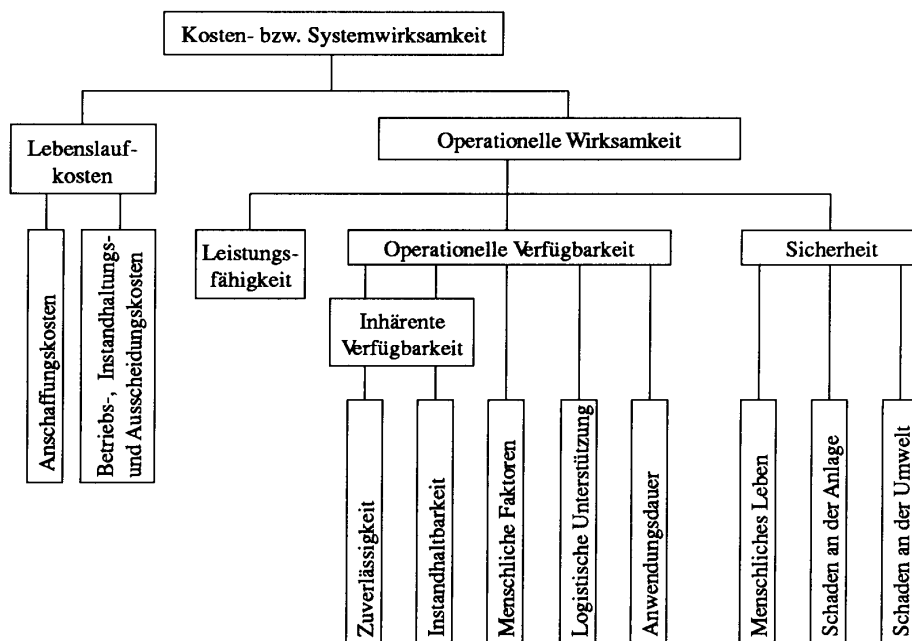


Bild 1.3 Kosten- bzw. Systemwirksamkeit bei komplexen Geräten und Systemen (Z = Zuverlässigkeit, I = Instandhaltbarkeit, S = Sicherheit), vgl. Anhang AI für Begriffe

2. Zuverlässigkeit von elektronischen Geräten und Systemen

Die Sicherstellung der Zuverlässigkeit (und für reparierbare Geräte und Systeme auch der Instandhaltbarkeit und Verfügbarkeit) muss hauptsächlich in der Entwicklungsphase erfolgen. Die entsprechenden Aktivitäten fallen in den Kompetenzbereich des Entwicklungsingenieurs und beinhalten Analysen (zur rechtzeitigen Erkennung und Beseitigung von Schwachstellen und zur Durchführung von Vergleichsstudien), die Untersuchung der Art und Auswirkung von Defekten und Ausfällen und die Überprüfung der Erfüllung von Entwicklungs- und Konstruktionsrichtlinien. In diesem Abschnitt werden die Analysemethoden der Zuverlässigkeitstechnik kurz dargelegt.

2.1 Vorausgesagte Zuverlässigkeit

2.1.1 Einleitung

Die vorausgesagte Zuverlässigkeit, ist jene Zuverlässigkeit, die anhand der Struktur der Betrachtungseinheit und der Zuverlässigkeit ihrer Elemente rechnerisch bestimmt wird. Die Voraussage ist wichtig, um Schwachstellen frühzeitig zu erkennen, Alternativlösungen zu untersuchen, Zusammenhänge zwischen Zuverlässigkeit, Instandhaltbarkeit, logistischer Unterstützung und Verfügbarkeit quantitativ zu erfassen sowie Zuverlässigkeitsforderungen an Unterlieferanten stellen zu können. Infolge der getroffenen Vereinfachungen bei der Festlegung der Modelle, der ungenügenden Berücksichtigung der Auswirkung innerer und äußerer Störeinflüsse (Schaltvorgänge, EMV usw.), der Vernachlässigung von Entwicklungs-, Konstruktions- und Fertigungsfehlern sowie der Unsicherheit der verwendeten Daten kann die vorausgesagte Zuverlässigkeit lediglich eine Schätzung der wahren Zuverlässigkeit darstellen. Letztere kann nur mit Hilfe von Zuverlässigkeitsprüfungen ermittelt werden. Bezogen auf die Ausfallrate eines Geräts oder Systems kann mit der nötigen Erfahrung, und falls in der Fertigungsphase geeignete Massnahmen getroffen werden, die Abweichung zwischen der vorausgesagten und der wahren Zuverlässigkeit klein gehalten werden (etwa Faktor 2, wobei in der Regel die Voraussage pessimistisch ist). Im Rahmen von Vergleichsstudien spielt zudem die absolute Genauigkeit keine wesentliche Rolle, weil man an relativen Werten verschiedener Konzepte interessiert ist.

Zu den theoretisch orientierten Überlegungen, die zur vorausgesagten Zuverlässigkeit führen, braucht man für die Analysen eingehende Kenntnisse der Betrachtungseinheit und der konkreten Möglichkeiten zur prinzipiellen Verbesserung der Zuverlässigkeit. Von der Betrachtungseinheit sind vor allem die Wirkungsweise und die Arbeitsbedingungen jedes Elements sowie die gegenseitigen Beziehungen und Beeinflussungen der verschiedenen Elemente wichtig (Eingang/Ausgang, Aufteilung der Last, Auswirkung der Ausfälle, Transiente usw.). Zu den konkreten Möglichkeiten zur Zuverlässigkeitsverbesserung gehören die Reduktion der thermischen, elektrischen und mechanischen Belastungen, die Verwendung qualitativ besserer bzw.

2.1

geeigneter Bauteile und Materialien, die Vereinfachung von Entwurf und Konstruktion, die Vorbehandlung der kritischen Bauteile und Baugruppen, das Hinzufügen von Redundanz sowie umfassende Qualifikationsprüfungen mit systematischer Analyse und Behebung der festgestellten Zuverlässigkeitsschwachstellen zu diesem Zweck sind spezielle Entwicklungs- und Konstruktionsrichtlinien bekannt [2, 1991].

Die Berechnung der vorausgesagten Zuverlässigkeit elektronischer und elektromechanischer Betrachtungseinheiten erfolgt gemäß folgenden Schritten:

1. Definition der geforderten Funktion und des Anforderungsprofils.
2. Aufstellung des Zuverlässigkeitsblockdiagramms (ZBD).
3. Bestimmung der Arbeitsbedingungen für jedes Element im ZBD.
4. Bestimmung der Ausfallrate für jedes Element im ZBD.
5. Berechnung der vorausgesagten Zuverlässigkeit für jedes Element im ZBD.
6. Berechnung der vorausgesagten Zuverlässigkeit der Betrachtungseinheit.
7. Behebung der Schwachstellen und Wiederholung der Schritte 2 bis 6, falls die Zuverlässigkeitsziele nicht erreicht worden sind.

2.1.2 Geforderte Funktion, Anforderungsprofil

Die geforderte Funktion spezifiziert die Aufgabe der Betrachtungseinheit. Ihre Festlegung bildet den Ausgangspunkt jeder Analyse, weil damit auch der Ausfall definiert wird (dabei ist es für die praktische Realisierbarkeit notwendig, für alle Parameter Toleranzbereiche statt fester Werte vorzuschreiben). Zusätzlich zur geforderten Funktion müssen auch die Umweltbedingungen festgelegt werden. Zu diesen gehören Umgebungstemperatur, Lagertemperatur, Feuchtigkeit, Staub, korrosive Atmosphäre, Vibrationen, Schocks, Lärm, Netzschwankungen usw. Aus den Umweltbedingungen und den internen Belastungen können die Arbeitsbedingungen der einzelnen Elemente der Betrachtungseinheit und damit ihre Ausfallraten bestimmt werden. Geforderte Funktion und Umweltbedingungen sind oft zeitabhängig. Daraus ergibt sich ein Anforderungsprofil. Es ist üblich, ein repräsentatives Anforderungsprofil ins Pflichtenheft aufzunehmen, damit die Zuverlässigkeitsziele und -berechnungen darauf bezogen werden können. Für komplexe Betrachtungseinheiten wird im Pflichtenheft oft mit einer groben Beschreibung der geforderten Funktion begonnen, welche im Laufe der Entwicklungsphase verfeinert wird.

2.1.3 Zuverlässigkeitsblockdiagramm

Das Zuverlässigkeitsblockdiagramm ist ein Ereignisdiagramm. Es gibt Antwort auf die Frage: Welche Elemente müssen zur Erfüllung der geforderten Funktion funktionieren, und welche dürfen ausfallen (Redundanz)? Die Aufstellung erfolgt, indem man die Betrachtungseinheit in Elemente zerlegt, die eine klar umrissene Aufgabe erfüllen. Diese Elemente werden dann zu einem Blockdiagramm derart zusammengefügt, dass die für die Funktionserfüllung notwendigen Elemente in Serie- und redundante Elemente in Parallelschaltung erscheinen (vgl. Abschnitt 2.1.7.2). Mit dem Zuverlässigkeitsblockdiagramm wird stets beim höchsten Integrationsniveau begonnen. Für jede tiefere Stufe wird dann die eigene geforderte Funktion formuliert und das eigene Zuverlässigkeitsblockdiagramm aufgestellt; dies hinunter bis zum Niveau, auf welchem die Zuverlässigkeitsangaben jedes Elements bekannt sind oder berechnet werden können (vgl. Bild

2. 1). Da es sich beim Zuverlässigkeitsblockdiagramm um ein Ereignisdiagramm handelt, dürfen

für jedes Element nur zwei Zustände (gut/ausgefallen) und nur eine Ausfallart (die dominierende, z. B. Kurzschluss oder Unterbrechung) angenommen werden.

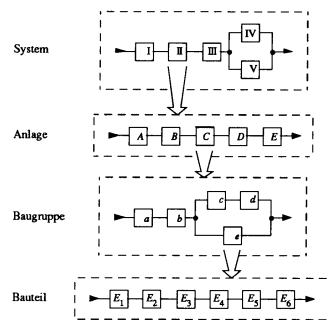
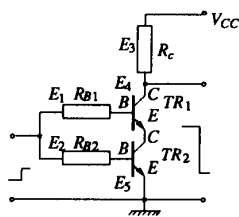


Bild 2.1 Top-down Aufstellung des Zuverlässigkeitsblockdiagramms

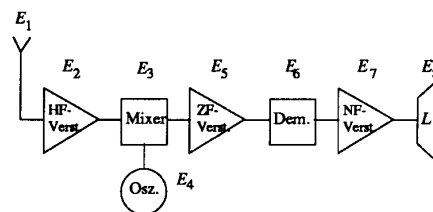
Beispiel 2.1 zeigt, dass das Zuverlässigkeitsblockdiagramm prinzipiell verschieden vom Funktionsdiagramm ist. Es wird auch deutlich, dass die Reihenfolge der in Serie geschalteten Elemente keine Rolle spielt.

Beispiel 2.1

Man stelle die Zuverlässigkeitsblockdiagramme folgender Schaltungen auf



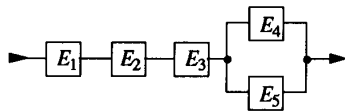
a) Elektronischer Schalter mit Redundanz



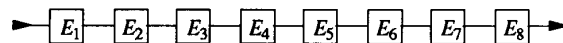
b) Einfacher Rundfunkempfänger

Lösung

Im Fall b) ist keine Redundanz vorhanden; für die Erfüllung der geforderten Funktion müssen alle Elemente funktionieren. Im Fall a) sind die Transistoren E_4 und E_5 in Redundanz, falls als Ausfallart ein Kurzschluss zwischen Emittor und Kollektor angenommen wird (die Ausfallart der Widerstände ist in der Regel eine Unterbrechung). Aus diesen Überlegungen folgt für die Zuverlässigkeitsblockdiagramme:



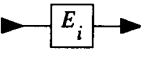
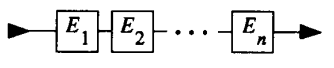
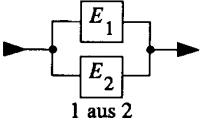
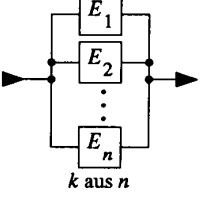
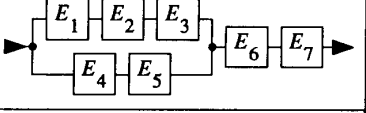
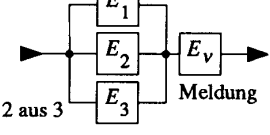
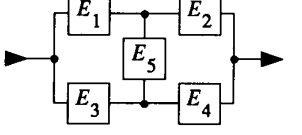
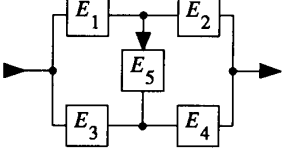
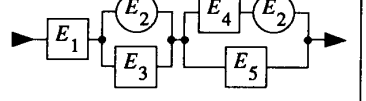
a) Elektronischer Schalter mit Redundanz



b) Einfacher Rundfunkempfänger

Tabelle 2.1 fasst die typischen Strukturen von Zuverlässigkeitsblockdiagrammen zusammen. Sie gibt auch die Formeln zur Berechnung der entsprechenden Zuverlässigkeit für den Fall einer (bis zum Systemausfall) nichtreparierbaren Betrachtungseinheit mit heisser Redundanz und unabhängigen Elementen an (vgl. Abschnitt 2.1.7).

Tabelle 2.1 Typische Strukturen von Zuverlässigkeitsblockdiagrammen und entsprechende Zuverlässigkeitsfunktionen (nichtreparierbar bis zum Systemausfall, heisse Redundanz, unabhängige Elemente)

Zuverlässigkeitsblockdiagramm	Zuverlässigkeitsfunktion ($R_S = R_S(t)$, $R_i = R_i(t)$)	Bemerkungen
	$R_S = R_i$	Einzelelement, für $\lambda(t) = \lambda \rightarrow R_i(t) = e^{-\lambda_i t}$
	$R_S = \prod_{i=1}^n R_i$	Seriemodell $\lambda_S(t) = \lambda_1(t) + \dots + \lambda_n(t)$
	$R_S = R_1 + R_2 - R_1 R_2$	Redundanz 1 aus 2, für $R_1(t) = R_2(t) = e^{-\lambda t}$ gilt $R_S(t) = 2 e^{-\lambda t} - e^{-2 \lambda t}$
	$R_1 = \dots = R_n = R$ $R_S = \sum_{i=k}^n \binom{n}{i} R^i (1-R)^{n-i}$	Redundanz k aus n, für $k = 1$ gilt $R_S = 1 - (1-R)^n$
	$R_S = (R_1 R_2 R_3 + R_4 R_5 - R_1 R_2 R_3 R_4 R_5) R_6 R_7$	Serie-/Parallelstruktur
	$R_1 = R_2 = R_3 = R$ $R_S = (3 R^2 - 2 R^3) R_v$	Majoritäts-Redundanz (allgemeiner Fall $n + 1$ aus $2n + 1$)
	$R_S = R_5 (R_1 + R_2 - R_1 R_2) (R_3 + R_4 - R_3 R_4) + (1 - R_5) (R_1 R_3 + R_2 R_4 - R_1 R_2 R_3 R_4)$	Brückenschaltung mit Zweiwegverbindung
	$R_S = R_4 (R_2 + R_1 (R_3 + R_5 - R_3 R_5) - (R_1 R_2) (R_3 + R_5 - R_3 R_5)) + (1 - R_4) R_1 R_3$	Brückenschaltung mit gerichteter Verbindung
	$R_S = R_2 R_1 (R_4 + R_5 - R_4 R_5) + (1 - R_2) R_1 R_3 R_5$	das Element E_2 erscheint zweimal im Zuverlässigkeitsblockdiagramm

2.1.4 Arbeitsbedingungen

Die Arbeitsbedingungen jedes Elements im Zuverlässigkeitsblockdiagramm haben einen direkten Einfluss auf die Zuverlässigkeit der Betrachtungseinheit und müssen deshalb sorgfältig ermittelt werden. Sie werden durch die *Umweltbedingungen* und die *internen Belastungen* der Elemente bestimmt. Grundsätzlich wird angenommen, dass die Bauteile keinesfalls überlastet werden. Es ist dabei zu berücksichtigen, dass für viele Parameter die *Belastbarkeit* mit steigender Umgebungstemperatur abnimmt. Dies gilt insbesondere für die Leistung, aber auch für die Spannung im Falle von Kondensatoren und für den Strom im Falle von Dioden. Es ist üblich, den *Belastungsfaktor* S folgendermassen zu definieren

$$S = \frac{\text{tatsächliche Belastung}}{\text{maximale Belastbarkeit bei } 25^\circ\text{C}}. \quad (2.1)$$

Eine absichtliche Nichtausnützung der maximalen Belastbarkeit bei gegebener Umgebungstemperatur θ_A nennt man *Unterlastung* (Derating). Der Zusammenhang zwischen dem Belastungsfaktor S und der Ausfallrate ist für einige wichtige elektronische Bauteile aus Bild 2.2 ersichtlich.

2.1.5 Ausfallrate

Für eine gegebene Betrachtungseinheit ist die Ausfallrate gleich der Wahrscheinlichkeit, bezogen auf δt , im Intervall $(t, t+\delta t]$ auszufallen, unter der Bedingung, dass sie bis zur Zeit t nicht ausgefallen ist. Es ist üblich, die Ausfallrate mit $\lambda(t)$ zu bezeichnen. Ihr typischer Verlauf für eine (unendlich grosse) Grundgesamtheit statistisch identischer Betrachtungseinheiten setzt sich aus den Perioden der Frühausfälle, der Ausfälle mit konstanter Ausfallrate und der Verschleissausfälle zusammen (Bild 1.2). Durch eine gezielte Vorbehandlung können die Frühausfälle provoziert werden, so dass zu Beginn der Nutzungsphase eine konstante (oder näherungsweise konstante) Ausfallrate angenommen werden kann. In diesem Fall wird

$$\lambda(t) = \lambda \quad (2.2)$$

gesetzt. Die Ausfallrate eines neuen Bauteils kann nur experimentell ermittelt werden. Für etablierte Bauteile liegen entsprechende Werte in speziellen Ausfallratenkatalogen vor. Wichtige Kataloge für elektronische und elektromechanische Bauteile sind das MIL-HDBK-217, der CNET-Ausfallratenkatalog und die DIN 40039. Für die Berechnung der vorausgesagten Zuverlässigkeit werden die Ausfallraten anhand solcher Ausfallratenkataloge bestimmt. Dabei geht man von folgenden Grundmodellen aus

$$\lambda = \lambda_b \pi_E \pi_Q \pi_A \quad (2.3)$$

für diskrete Bauteile, und

$$\lambda = \pi_Q \pi_L (C_1 \pi_T \pi_V + C_2 \pi_E) \quad (2.4)$$

für ICs. λ_b berücksichtigt die Umgebungstemperatur θ_A und die elektrische Belastung S , π_E die Umweltbedingungen, π_Q die Fertigungsqualität und die Vorbehandlung, π_A die Bauteileigenschaften und die Anwendung (es können mehrere Faktoren sein), C_1 die Komplexität, C_2 die

Anzahl Anschlüsse (Pins) und den Gehäusotyp, π_T die Chiptemperatur θ_J und die Technologie, π_V die Spannungsbelastung (CMOS) und π_L die Reife des Herstellungsprozesses. Der Wert von λ liegt zwischen etwa 10^{-9} h^{-1} für einfache passive Bauteile und 10^{-7} h^{-1} für VLSI-ICs (Bilder 2.2 und 2.3). Es ist üblich, die Einheit 10^{-9} h^{-1} mit *FIT* (Failures In Time) zu bezeichnen, wobei 1000 FIT eine Ausfallquote m von etwa 1% pro Jahr entspricht vgl. Gl (5.2).

λ_b steigt in der Regel exponentiell mit der *Umgebungstemperatur* θ_A an. Als Beispiel zeigt Bild 2.3 den Verlauf von λ_b für einige wichtige diskrete Bauteile gemäss MIL-HDBK-217E. Der Einfluss des Belastungsfaktors S ist in Bild 2.2 ebenfalls ersichtlich.

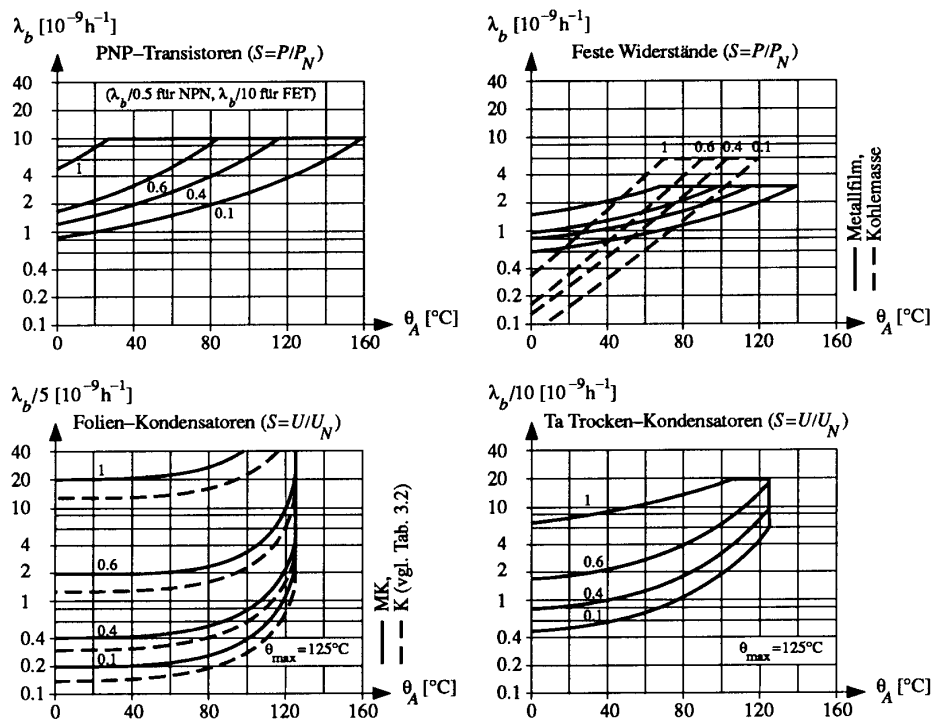


Bild 2.2 λ_b in 10^{-9} h^{-1} (FIT) als Funktion der Umgebungstemperatur θ_A , mit $S = 0.1, 0.4, 0.6$ und 1 als Parameter (MIL-HDBK-217E)

Für ICs zeigt Bild 2.3 den Verlauf von π_T als Funktion der *Chiptemperatur* θ_J , ebenfalls gemäss MIL-HDBK-217E. Die Darstellung rechts im Bild 2.4, mit dem linearen Massstab, bringt die grosse Abhängigkeit des Faktors π_T von der Chiptemperatur deutlicher zum Ausdruck (aus dieser Darstellung lässt sich die Entwicklungsrichtlinie $\theta_J \leq 85^{\circ}\text{C}$ begründen). Der Unterschied zwischen Kunststoff- und Keramik- (bzw. Cerdip-) Gehäusen muss nicht so gross sein, wie Bild 2.3 zeigt. In den meisten Anwendungen kann der Wert von π_T für Keramik-Gehäuse auch für Kunststoff-Gehäuse verwendet werden.

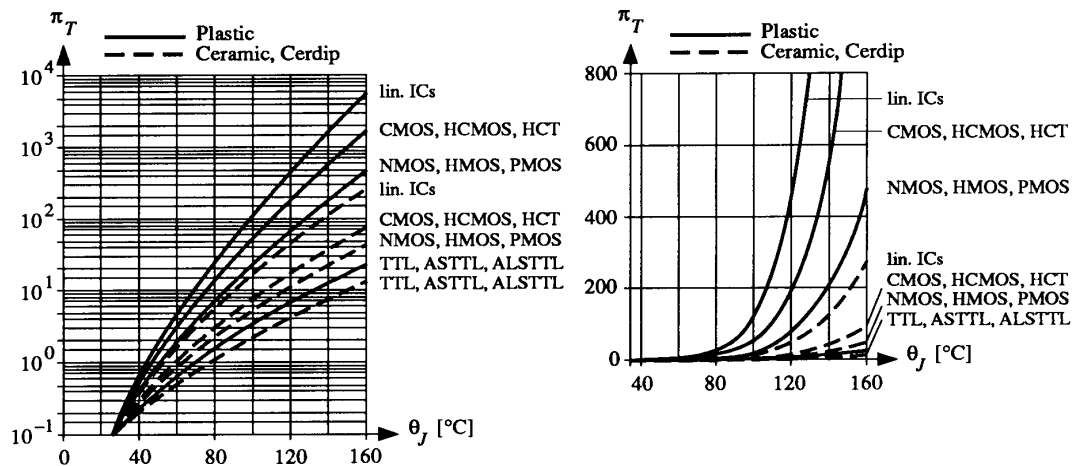


Bild 2.3 Zusammenhang zwischen dem Faktor π_T und der Chiptemperatur θ_J für ICs (MIL-HDBK-217E)

Der Einfluss der *Umweltbedingungen* (klimatisch und mechanisch) wird mit dem Faktor π_E berücksichtigt. Tabelle 2.3 fasst einige wichtige Anwendungsklassen zusammen und gibt die entsprechenden π_E -Faktoren an. G_B darf auch für übliche Industrieanwendungen verwendet werden (G_M entspricht z. B. einem fahrenden Jeep auf einer Landstrasse).

Tabelle 2.2 Wichtige Umweltbedingungen nach CNET-Ausfallratenkatalog und entsprechende π_E -Faktoren (MIL-HDBK-217E)

Bezeichnung	Belastung					π_E - Faktor				
	Vibra- tionen	Lärm dB	Staub	RH %	mech. Schocks	ICs		d. H.	R	C
G_B (Ground benign)	0	40 - 70	s	20 - 70	0	0.38	0.2	1	1	1
G_F (Ground fixed)	2 – 60 Hz 0.5 g	40 - 70	m	20 - 90	10 – 15 g / 11 ms	2.5	0.78	2 - 5.5	1.5 - 2.9	1.4 - 2.4
G_M (Ground mobile)	2 – 300 Hz 0 – 5 g	40 - 70	m	20 - 100	500 g / 0.5 ms 15 – 50 g / 11 ms	4.2	2.2	4.9 - 17	7.8 - 11	7.8 - 12

RH = relative Feuchtigkeit; d. H. = diskrete Halbleiter; R = Widerstände; C = Kondensatoren; s = schwach; m = mässig; $g = 9.81 \text{ m/s}^2$

Für den π_Q -Faktor werden verschiedene *Qualitätsklassen* unterschieden. Tabelle 2.3 gibt einige Beispiele davon. Der π_Q -Faktor berücksichtigt den Einfluss der Fertigungsqualität und der *Vorbehandlung*. Die Erfahrung zeigt, dass marktübliche Bauteile guter Qualität bei den Klassen B-1, JAN und M in Tab. 2.3 liegen, unabhängig vom Gehäusotyp (Kunststoff). Eine *Vorbehandlung* drängt sich nur für kritische Bauteile auf (neuere VLSI-ICs, Kundensaltungen sowie einige lineare ICs und Leistungsbauteile).

Tabelle 2.3 π_Q -Faktoren (MIL-HDBK-217E)

Monolithisch integrierte Schaltungen	S 0.25	$S - 1$ 0.75	B 1.0	$B - 1$ 2.0	$B - 2$ 5.0	D 10	$D - 1$ 20
Hybride integrierte Schaltungen	S 0.25	$S - 1$ 0.4	B 0.5	$B - 1$ 1.0		D 20	
Diskrete Halbleiterbauteile		JANTXV 0.5 - 0.7	JANTX 1.0	JAN 1.8 - 5.0		tief 2.5 - 25	Kunststoff 8 - 50
Widerstände (ohne Potentiometer)	S 0.03	R 0.1	P 0.3	M 1	NE 5	tief 15	
Kondensatoren	S 0.03	R 0.1	P 0.3	M 1	L 1.5	NE 3	tief 10

2.8

2.1.6 Vorausgesagte Zuverlässigkeit des Einzelements

Für das (nichtreparierbare) Einzelement wird oft eine konstante Ausfallrate $\lambda(t) = \lambda$ angenommen. In diesem Fall gilt für die Zuverlässigkeitsfunktion

$$R(t) = e^{-\lambda t}. \quad (2.5)$$

Der Mittelwert der ausfallfreien Arbeitszeit ist dann gleich $1 / \lambda$. Es ist üblich

$$\frac{1}{\lambda} = MTBF \quad (2.6)$$

zu setzen, wobei MTBF für *Mean Time Between Failures* steht. Für nicht konstante Ausfallraten wird mit der Verteilungsfunktion der ausfallfreien Arbeitszeit τ

$$F(t) = \Pr\{\tau \leq t\}$$

operiert (vgl. Tab. A2.1). In diesem Fall gilt

$$R(t) = \Pr\{\tau > t\} = 1 - F(t) \quad (2.7)$$

und für den Mittelwert der ausfallfreien Arbeitszeit MTTF

$$MTTF = E[\tau] = \int_0^{\infty} R(t) dt, \quad (2.8)$$

wobei MTTF für *Mean Time To Failure* steht. Die Gl. (2.8) stellt eine fundamentale Beziehung dar. Sie gilt nicht nur für ein Einzelement, sondern bleibt gleich, auch wenn $R(t)$ die Zuverlässigkeitsfunktion einer beliebigen Betrachtungseinheit (eines beliebigen Systems) ist. Um dies zu betonen, wird im folgenden mit $R_S(t)$ und $MTTF_S$ operiert

$$MTTF_S = \int_0^{\infty} R_S(t) dt. \quad (2.9)$$

Darüber hinaus kann Gl. (2.8) bzw. (2.9) auch für reparierbare Betrachtungseinheiten (Systeme) verwendet werden. Nimmt man z. B. an, dass nach einem Ausfall das System als neuwertig betrachtet werden kann, so läuft nach der Erneuerung wieder eine ausfallfreie Arbeitszeit τ mit der gleichen Verteilungsfunktion und damit mit dem gleichen Erwartungswert an. Im Falle wo die Betrachtungseinheit eine auf T_{UL} beschränkte *Anwendungsdauer* aufweist, springt $R(t)$ bei $t = T_{UL}$ auf 0, so dass die obere Grenze des Integrals in Gl. (2.8) bzw. (2.9) gleich T_{UL} gesetzt werden kann; im folgenden wird stets $T_{UL} = \infty$ angenommen.

2.1.7 Berechnung der vorausgesagten Zuverlässigkeit einer Betrachtungseinheit

Die Berechnung der vorausgesagten Zuverlässigkeit einer Betrachtungseinheit (eines Systems) erfolgt für viele praktischen Anwendungen mit Hilfe der in Tab. 2.1 angegebenen Grundmodelle.

2.1.7.1 Betrachtungseinheiten ohne Redundanz (Seriemodell)

Eine Betrachtungseinheit (ein System) ist im Sinne der Zuverlässigkeit ohne Redundanz, falls für die Erfüllung der geforderten Funktion alle Elemente $(E_1, \dots, E_n)$ funktionieren müssen. In diesem Fall besteht das Zuverlässigkeitsblockdiagramm aus einer Serieschaltung (Tab. 2.1). Für die Untersuchungen geht man oft von der Annahme aus, dass die ausfallfreien Arbeitszeiten $\tau_1, \dots, \tau_n$ die Elemente $E_1, \dots, E_n$ unabhängige Zufallsgrößen sind (unabhängige Arbeitsweise aller Elemente). Die Berechnungen sind dann besonders einfach. Mit e_i sei das Ereignis

das Element E_i arbeitet ausfallfrei im Intervall $(0, t]$

bezeichnet. Die Wahrscheinlichkeit für das Eintreten dieses Ereignisses ist gleich der Zuverlässigkeitsfunktion $R_i(t)$ des Elements E_i . Es gilt also

$$\Pr\{e_i\} = \Pr\{\tau_i > t\} = R_i(t). \quad (2.10)$$

Die Betrachtungseinheit arbeitet nun im Intervall $(0, t]$ ausfallfrei, wenn gleichzeitig alle Elemente $E_1, \dots, E_n$ im Intervall $(0, t]$ ausfallfrei arbeiten. Folglich gilt für die Zuverlässigkeitsfunktion der Betrachtungseinheit (des Systems)

$$R_S(t) = \Pr\{e_1 \cap \dots \cap e_n\}. \quad (2.11)$$

Der Index S steht hier und im folgenden für System und bezieht sich auf die Betrachtungseinheit. Wegen der vorausgesetzten Unabhängigkeit der Arbeitsweise aller Elemente folgt

$$R_S(t) = \prod_{i=1}^n R_i(t). \quad (2.12)$$

Für die Ausfallrate des Systems gilt dann

$$\lambda_S(t) = \sum_{i=1}^n \lambda_i(t). \quad (2.13)$$

wobei $\lambda_i(t)$ die Ausfallrate des Elements E_i ist. Die Gl. (2.13) erlaubt folgende wichtige Schlussfolgerung:

Die Ausfallrate eines Systems ohne Redundanz, bestehend aus Elementen, deren ausfallfreie Arbeitszeiten unabhängig sind, ist gleich der Summe der Ausfallraten seiner Elemente.

Der Erwartungswert der ausfallfreien Arbeitszeit des Systems folgt aus Gl. (2.9). Der Spezialfall, in welchem alle Elemente eine konstante Ausfallrate aufweisen ($\lambda_i(t) = \lambda_i$), führt zu

$$R_S(t) = e^{-(\lambda_1 + \dots + \lambda_n)t}, \quad \lambda_S(t) = \lambda_S = \sum_{i=1}^n \lambda_i, \quad MTBF_S = \frac{1}{\lambda_S}. \quad (2.14)$$

2.1.7.2 Der Begriff der Redundanz

Die hohen Forderungen bezüglich Zuverlässigkeit, Verfügbarkeit und/oder Sicherheit, die in vielen Anwendungen gestellt werden, lassen sich oft nur mit Hilfe von Redundanz erfüllen. Redundanz entsteht in ihrer einfachsten Form durch das Einführen von zusätzlichen Elementen (Reserve-Elemente), welche dieselbe Funktion wie andere Elemente ausführen und dadurch die Zuverlässigkeit, die Verfügbarkeit und/oder die Sicherheit der Betrachtungseinheit (des Systems) erhöhen. Im Zuverlässigkeitsblockdiagramm wird eine Redundanz als Parallelschaltung erscheinen. Es werden prinzipiell drei Redundanzarten unterschieden:

1. Heisse Redundanz (aktive oder parallele Redundanz): Das Redundanzelement ist von Anfang an der gleichen Belastung wie das arbeitende Element ausgesetzt; die Ausfallrate im Reservezustand ist gleich der Ausfallrate im Arbeitszustand.
2. Warme Redundanz (leicht belastete Redundanz): Das Redundanzelement ist bis zum Ausfall des arbeitenden Elements oder bis zu seinem eigenen Ausfall einer kleineren Belastung ausgesetzt; die Ausfallrate im Reservezustand ist kleiner als die Ausfallrate im Arbeitszustand.
3. Kalte Redundanz (Standby- oder unbelastete Redundanz): Das Redundanzelement ist bis zum Ausfall des arbeitenden Elements keiner Belastung ausgesetzt; die Ausfallrate im Reservezustand wird gleich Null gesetzt.

Aus obigen Darlegungen könnte man schliessen, dass eine Redundanz lediglich aus der Verdopplung der schwachen Teile besteht. Das ist nur teilweise richtig, denn oft werden raffiniertere Entwicklungs- und/oder Konstruktionsmassnahmen (wie beispielsweise fehlerkorrigierende Codierung oder Parallelausführung der gleichen Funktion aber mit anderen Elementen) verwendet. Es kann auch vorkommen, dass die Redundanz nicht genau die gleiche Funktion ausführt wie das Arbeitselement (Pseudo-Redundanz). Ausserdem bedingt die Redundanz oft eine Auf

teilung der Last und/oder das Einführen von Zusatzeinrichtungen zur Überwachung bzw. Umschaltung, die in den Zuverlässigkeitsberechnungen berücksichtigt werden müssen (im Falle einer Aufteilung der Last, kann die Untersuchung oft mit Hilfe von stochastischen Prozessen erfolgen, vgl. Tab. 3.3 bis 3.5, mit $\#u = 0$ für den nichtreparierbaren Fall). Einige wichtige (bis zum Systemausfall nichtreparierbaren) Redundanzstrukturen mit unabhängigen Elementen in heisser Redundanz werden in den Abschnitten 2.1.7.3 bis 2.1.7.6 untersucht.

Ein Parallelmodell besteht aus n Elementen, die im Sinne der Zuverlässigkeit parallel geschaltet sind. Für die Erfüllung der geforderten Funktion sind k Elemente notwendig. Die übrigen $n - k$ Elemente bilden die Reserve. Eine solche Struktur wird als *Redundanz k aus n* bezeichnet.

Es sei zuerst der Fall einer *heissen Redundanz 1 aus 2* untersucht (Tab. 2.1). Die geforderte Funktion ist erfüllt, wenn im Intervall $(0, t]$ mindestens eines der Elemente E_1 oder E_2 ausfallfrei arbeitet. Mit der gleichen Bezeichnung wie im Abschnitt 2.1.7.1 folgt

$$R_S(t) = \Pr\{e_1 \cup e_2\} = \Pr\{e_1\} + \Pr\{e_2\} - \Pr\{e_1 \cap e_2\}. \quad (2.15)$$

Sind die ausfallfreien Arbeitszeiten der Elemente E_1 und E_2 unabhängig (unabhängige Arbeitsweise der Elemente E_1 und E_2), dann gilt für die Zuverlässigkeitsfunktion des Systems

$$R_S(t) = R_1(t) + R_2(t) - R_1(t)R_2(t). \quad (2.16)$$

Der Mittelwert der ausfallfreien Arbeitszeit ($MTTF_S$) lässt sich aus Gl. (2.9) berechnen. Der Spezialfall gleicher Elemente mit konstanter Ausfallrate ($R_1(t) = R_2(t) = e^{-\lambda t}$) führt zu

$$R_S(t) = 2e^{-\lambda t} - e^{-2\lambda t} \quad (2.17)$$

und

$$MTTF_S = \frac{2}{\lambda} - \frac{1}{2\lambda} = \frac{3}{2\lambda}. \quad (2.18)$$

Die Verallgemeinerung zur *heissen Redundanz k aus n* (Tab. 2.1) ist leicht möglich. Für den Fall gleicher Elemente gilt

$$R_S(t) = \sum_{i=k}^n \binom{n}{i} R^i(t) (1 - R(t))^{n-i}. \quad (2.19)$$

Der verhältnismässig kleine Gewinn für die $MTTF_S$ gemäss Gl. (2.18) wird viel grösser, wenn eine Reparatur ohne Betriebsunterbrechungen (auf Niveau Betrachtungseinheit) zugelassen wird (Tab. 3.3). Für kurze Missionen ($t \ll 1/\lambda$) ist aber auch im nichtreparierbaren Fall der Gewinn durch die Redundanz gross, wie Bild 2.4 zeigt.

2.1.7.3 Parallelmodelle

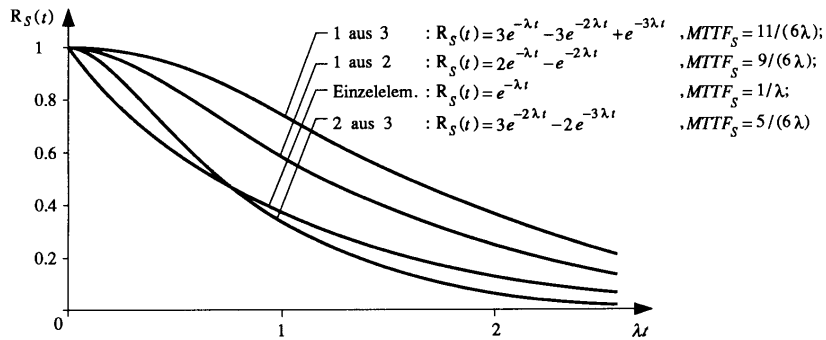
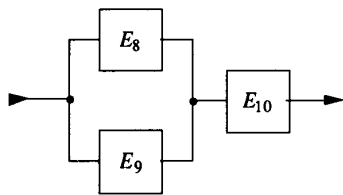


Bild 2.4 Zuverlässigkeitsfunktion für das Einzelelement und für eine heiße Redundanz 1 aus 2, 1 aus 3 und 2 aus 3 (nichtreparierbar bis zum Systemausfall, gleiche, unabhängige Elemente, konstante Ausfallrate λ)

2.1.7.4 Serie-/Parallelstrukturen

Serie-/Parallelstrukturen lassen sich durch sukzessive Anwendung der Formeln für das Serie- und das Parallelmodell untersuchen. Das gilt insbesondere im Fall eines nicht-reparierbaren Systems mit heißer Redundanz und unabhängigen Elementen. Zur Illustration des Vorgehens wird das 5. Modell in Tab. 2.1 behandelt.

1. Schritt: Die Serieschaltungen von E_1, E_2 und E_3 bzw. von E_4 und E_5 bzw. von E_6 und E_7 werden durch E_8 bzw. E_9 bzw. E_{10} ersetzt. Man erhält



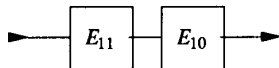
mit

$$R_8(t) = R_1(t)R_2(t)R_3(t)$$

$$R_9(t) = R_4(t)R_5(t)$$

$$R_{10}(t) = R_6(t)R_7(t)$$

2. Schritt: Die Parallelschaltung von E_8 und E_9 wird durch E_{11} ersetzt. Man erhält



mit

$$R_{11}(t) = R_8(t) + R_9(t) - R_8(t)R_9(t)$$

3. Schritt: Für die Zuverlässigkeitsfunktion des Systems folgt aus Schritt 1 und Schritt 2

$$R_S = R_{11} R_{10} = (R_1 R_2 R_3 + R_4 R_5 - R_1 R_2 R_3 R_4 R_5) R_6 R_7, \quad (2.20)$$

mit $R_S = R_S(t)$ und $R_i = R_i(t)$, $i = 1, \dots, 7$.

Für den Mittelwert der ausfallfreien Arbeitszeit gilt stets Gl. (2.9) mit $R_S(t)$ aus Gl. (2.20). Weisen alle Elemente eine konstante Ausfallrate (λ_1 bis λ_7) auf, so folgt

$$R_S(t) = e^{-(\lambda_1 + \lambda_2 + \lambda_3 + \lambda_6 + \lambda_7)t} + e^{-(\lambda_4 + \lambda_5 + \lambda_6 + \lambda_7)t} - e^{-(\lambda_1 + \lambda_2 + \lambda_3 + \lambda_4 + \lambda_5 + \lambda_6 + \lambda_7)t}$$

und damit

$$MTTF_S = \frac{1}{\lambda_1 + \lambda_2 + \lambda_3 + \lambda_6 + \lambda_7} + \frac{1}{\lambda_4 + \lambda_5 + \lambda_6 + \lambda_7} - \frac{1}{\lambda_1 + \lambda_2 + \lambda_3 + \lambda_4 + \lambda_5 + \lambda_6 + \lambda_7}. \quad (2.21)$$

Wichtig für die Überlegungen im Zusammenhang mit der Wechselwirkung zwischen Zuverlässigkeit und Energieverbrauch ist der Vergleich zwischen verschiedenen Grundformen für Serie-/Parallenstrukturen. Die Bilder 2.5 und 2.6 zeigen einen solchen Vergleich. Verglichen werden in Bild 2.5 ein Einzelement E_1 (Ausfallrate λ_1) mit der heissen Redundanz 1 aus 2 inkl. Serie-Element E_2 zur Überwachung und Umschaltung (Ausfallrate λ_2) und in Bild 2.6 die heisse Redundanz 1 aus 2 von Bild 2.5 mit der Struktur, die man erhält, wenn auch für das Element E_2 eine heisse Redundanz gefordert wäre (offenbar muss $\lambda_1 < \lambda_2 < \lambda_3$ gelten, die Grenzfälle $\lambda_1 = \lambda_2$ bzw. $\lambda_1 = \lambda_2 = \lambda_3$ sind nur zur Information angeführt). Der obere Teil der Bilder 2.5 und 2.6 zeigt den Verlauf der Zuverlässigkeitsfunktion; man erkennt, dass der Gewinn vor allem für sehr kleine Werte von $\lambda_1 t$ gross ist und dass ohne zu grosse Einbusse auf die Zuverlässigkeit $\lambda_2 \approx 0.05\lambda_1$ (bzw. $\lambda_3 \approx 0.1\lambda_2$ für Bild 2.6) noch akzeptiert werden kann. Dieses Resultat kommt auch im unteren Teil der Bilder 2.5 und 2.6 zum Ausdruck (Verhältnis der Mittelwerte der ausfallfreien Arbeitszeit). Der Energieverbrauch ist praktisch proportional zur Anzahl der jeweiligen Elemente; der Gewinn an Zuverlässigkeit ist stark vom Verhältnis der Ausfallraten abhängig. Dieser Vergleich ist noch wichtiger im Falle reparierbarer Betrachtungseinheiten, vgl. dazu Abschnitt 3.4.

2.1.7.5 Majoritätsredundanz

Die *Majoritätsredundanz* ist eine spezielle Ausführung der Redundanz k aus n , welche vor allem bei *digitalen Schaltungen* zur Anwendung gelangt. Für die Erfüllung der geforderten Funktion werden $2n + 1$ gleiche Elemente in heisser Redundanz verwendet. Die $2n + 1$ Ausgänge werden durch ein *Vergleichselement* miteinander verglichen, wobei am Ausgang dieses Elements das Signal erscheint, welches gleich wie die Mehrzahl der $2n + 1$ Eingangssignale ist. Die Untersuchung erfolgt mit Hilfe der Prozedur für die Serie-/Parallelstrukturen. Für $n = 1$ ergibt sich eine heisse Redundanz 2 aus 3 in Serie mit dem Vergleichselement E_V (Tab. 2.1). Die Majoritätsredundanz realisiert damit in einfacher Weise eine *störungstolerante Struktur* mit automatischer Weiterführung der geforderten Funktion (ohne Umschalteneinrichtung) bei einem Ausfall ($n = 1$). Das *Vergleichselement* für eine Majoritätsredundanz 2 aus 3 besteht aus drei NAND-Gattern mit je 2 Eingängen und einem NAND-Gatter mit 3 Eingängen pro Bit. Ebenso einfach ist auch der *Alarmgeber* (drei EX-NOR-Gatter mit 2 Eingängen und ein NAND-Gatter mit 3 Eingängen).

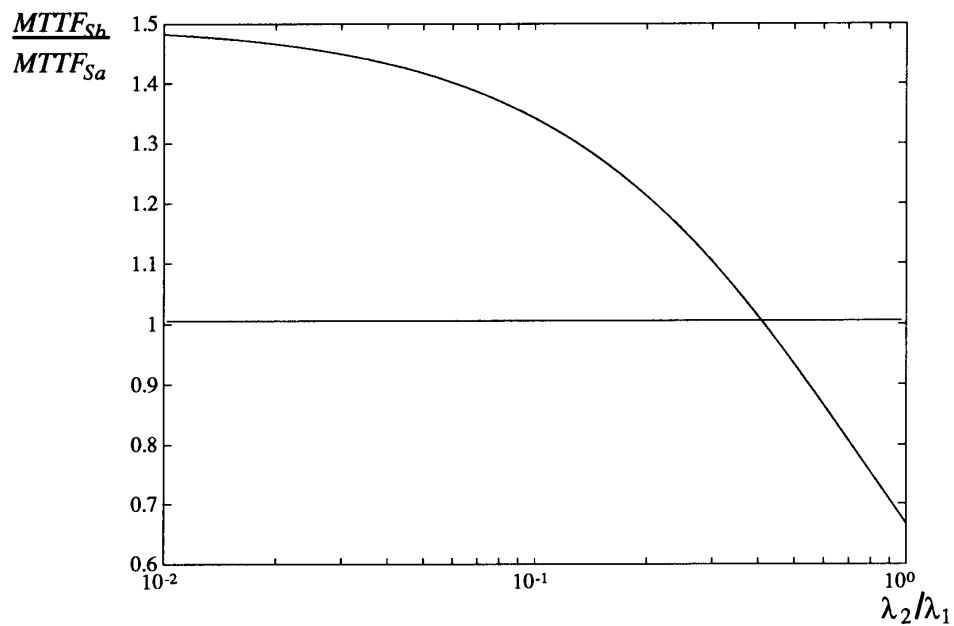
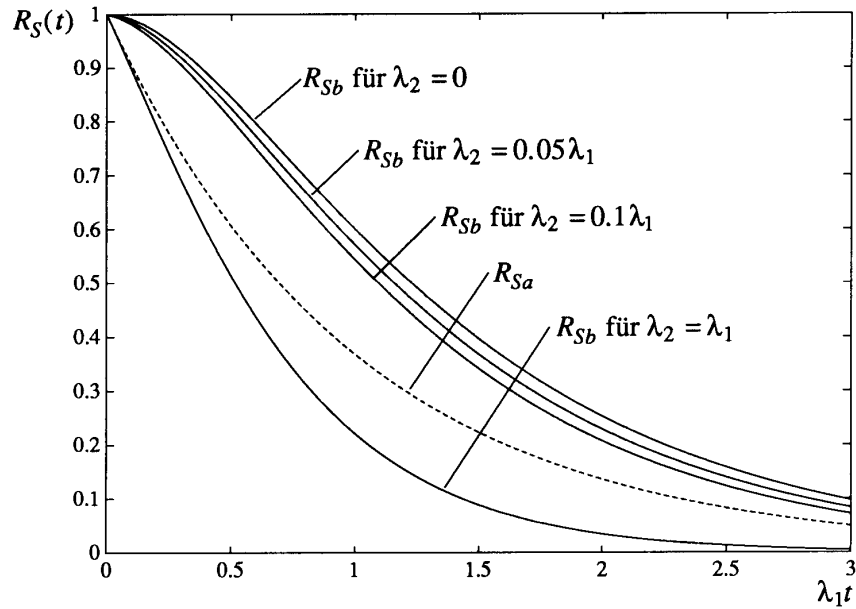
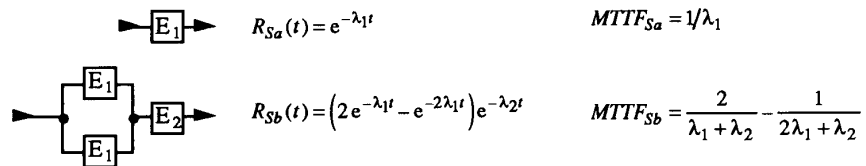


Bild 2.5 Vergleich zwischen Grundstrukturen (nichtreparierbar bis Systemausfall, heisse Redundanz, unabhängige Elemente, konstanten Ausfallarten)

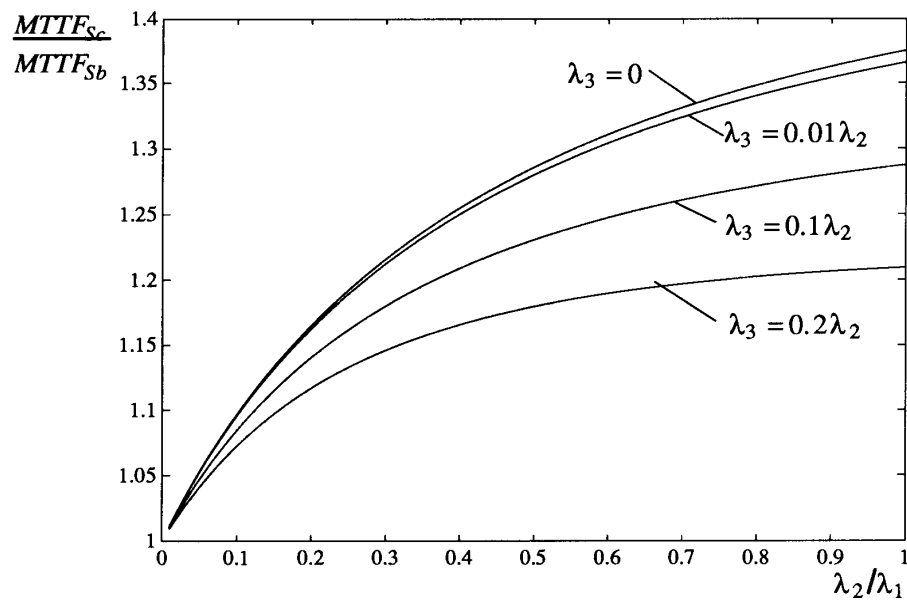
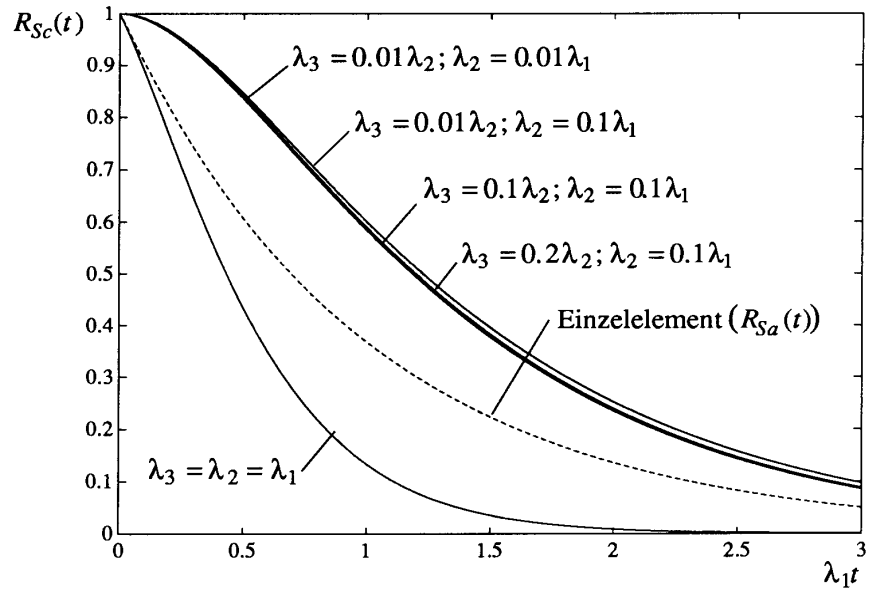
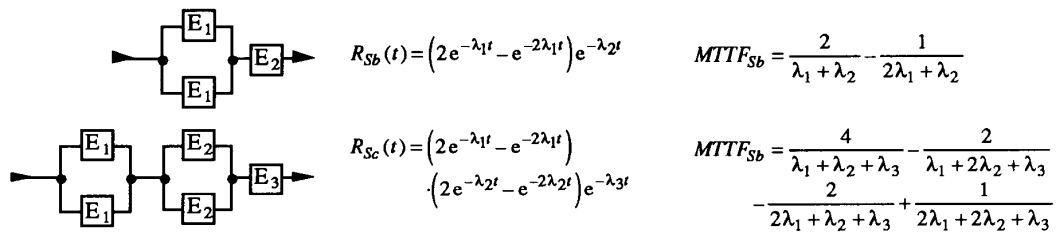


Bild 2.6 Vergleich zwischen Grundstrukturen (nichtreparierbar bis Systemausfall, heisse Redundanz, unabhängige Elemente, konstanten Ausfallarten)

2.1.7.6 Systeme mit komplexerer Struktur

Unter den Annahmen heisser Redundanz, unabhängiger Elemente und nicht reparierbar bis Systemausfall können Systeme beliebiger Struktur untersucht werden. Für manuelle Auswertungen kennt man die Methode des Schlüsselements (Satz der totalen Wahrscheinlichkeit, vgl. die letzten 3 Beispiele der Tab. 2.1), der erfolgreichen Pfade, des Zustandsraums oder allg. der Boole'schen Funktionen. Computerunterstützte Methoden sind auch bekannt, vgl. Abschnitt 2.1.9.

2.1.7.7 Elemente mit mehr als einer Ausfallart

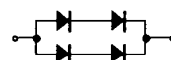
In den vorangehenden Abschnitten wurde angenommen, jedes Element weise nur eine Ausfallart auf, z. B. Kurzschluss oder Unterbrechung. In Wirklichkeit ist es aber so, dass viele Bauteile verschiedene Ausfallarten aufweisen. Um sich im Falle elektronischer Bauteile gleichzeitig gegen eine Unterbrechung oder einen Kurzschluss zu schützen (störungstolerant für mindestens einen Ausfall), wird die sogenannte Quadredundanz verwendet, mit oder ohne Querverbindung, je nachdem, ob die Unterbrechungen oder die Kurzschlüsse überwiegen. Für die Quadredundanz gilt

$$\bar{R}_S = 2\bar{R}_U^2 - \bar{R}_U^4 + (2\bar{R}_K - \bar{R}_K^2)^2 \quad \text{bzw.} \quad (2.22)$$



bzw.

$$\bar{R}_S = 2\bar{R}_K^2 - \bar{R}_K^4 + (2\bar{R}_U - \bar{R}_U^2)^2 \quad (2.23)$$



Dabei sind $\bar{R}_U$ und $\bar{R}_K$ die Wahrscheinlichkeiten für einen Ausfall als Unterbrechung bzw. als Kurzschluss ($\bar{R} = 1 - R = \bar{R}_U + \bar{R}_K$)

2.1.8 Störungstolerante Betrachtungseinheiten

Für spezielle Anwendungen (hohe Zuverlässigkeits-, Verfügbarkeits- und/oder Sicherheitsforderungen) muss die Betrachtungseinheit (das System) störungstolerant konzipiert werden. Die Grundidee dabei ist, das Auftreten einer Störung (Ausfall, Defekt) sofort zu erkennen und die Betrachtungseinheit derart zu rekonfigurieren, dass sie sicher bleibt und ihre geforderte Funktion mit minimaler Leistungseinbusse weiterführen kann. Methoden zur Untersuchung ausfalltoleranter Systeme sind in den Abschnitten 2.1.7.2 bis 2.1.7.7 angegeben worden; vgl. insbesondere die Majoritätsredundanz und die Quadredundanz. Letztere ist eine der wenigen Strukturen, welche mindestens einen Ausfall irgendwelchen Art tolerieren kann. Der Preis dafür ist aber eine Vervielfachung der Anzahl Elemente. Im Zusammenhang mit ICs wäre es jedoch nicht sinnvoll, eine Quadredundanz mit vier Elementen auf dem gleichen Chip zu realisieren, weil das Risiko einer Wechselwirkung mit darauffolgendem Totalausfall (Single-point failure) zu gross ist.

Die Berücksichtigung der möglichen Ausfallarten ist im Rahmen der Untersuchung von ausfalltoleranten Systemen ausschlaggebend. Dies erfolgt mit den Methoden zur Ausfallartenanalyse (FMEA/FMECA, FTA usw.), die in Abschnitt 2.2 vorgestellt werden. Die Ausfallartenanalyse ermöglicht u. a. die Erfassung von Abhängigkeiten und Wechselwirkungen zwischen den

Elementen der Betrachtungseinheit (des Systems) und damit das Planen geeigneter Massnahmen zur Verhinderung von Folgeausfällen. Folgeausfälle können durch Schutz der Ein- und/oder Ausgänge der kritischen Teile oft vermieden werden. Weitere mögliche Massnahmen sind die Realisierung redundanter Elemente mit verschiedenen Bauteilen und das Einfügen von StandbyElementen zu einer 2-aus-3-Majoritätsredundanz, welche erst beim Ausfall zweier aktiver Elemente in Betrieb genommen werden. Diese und ähnliche Methoden werden erfolgreich in Raumfahrtprojekten und allgemein in Geräten und Systemen angewandt, an die hohe Zuverlässigkeits- und/oder Sicherheitsanforderungen gestellt werden. Alle Teile, welche an mehreren Funktionen beteiligt sind, müssen besonders sorgfältig ausgelegt bzw. dimensioniert werden. Zu diesen gehören beispielsweise Clock-Schaltungen, BusNahtstellen sowie eingebaute Prüfungen und Überwachungen (Paritätsprüfung, Watchdog, Zustandsprüfung, Betriebsüberwachung usw.). Der Erkennung und Lokalisierung verborgener Ausfälle und der Vermeidung falscher Meldungen ist grosse Bedeutung beizumessen. Falsche Meldungen können z. B. bereits durch Synchronisationsprobleme in der Überwachung verursacht werden. Obige Darlegungen gelten primär für Ausfälle, lassen sich aber auf Defekte der Hardware oder der Software übertragen. Die Untersuchungen bei der Software sind allerdings schwieriger, weil entsprechende Modelle noch nicht ausgereift sind (Wichtige Entwicklungs- und Konstruktionsrichtlinien sind in [2, 1991] gegeben).

2.1.9 Computerunterstützte Zuverlässigkeits- und Verfügbarkeitsanalyse komplexer Systeme

Die Berechnung der Zuverlässigkeit komplexer Geräte und Systeme mit vielen Bauteilen oder mit vernetzter Zuverlässigkeitsstruktur kann zeitraubend werden. Der Einsatz von Computerprogrammen für solche Anwendungen liegt auf der Hand. Von einer solchen Software erwartet man insbesondere grosse Benutzerfreundlichkeit, Robustheit und Übertragbarkeit. Spezielle Leistungsmerkmale und Eigenschaften sind:

- 1 . Das Programm soll direkt vom Entwicklungsingenieur am Arbeitsplatz verwendet werden können (PC oder Terminal), verwaltet wird es allerdings von der zentralen Stelle für die Qualitäts- und Zuverlässigkeitssicherung.
- 2 . Das Programm muss beliebig komplexe Zuverlässigkeitsblockdiagramme (bis zu 10 Niveaus mit frei wählbarer Bezeichnung) verarbeiten können; gleiche Elemente dürfen mehrmals auftreten; die Struktur darf vernetzt sein; Serie- und Parallelmodelle sollen automatisch editiert werden.
- 3 . Bauteile sollen nach der kommerziellen oder einer firmeninternen Bezeichnung gespeichert werden (bis 100'000); nicht anwendungsspezifische Daten sollen fest gespeichert werden können (Datenbank); anwendungsspezifische Daten werden abgefragt, hier sollen folgende Erleichterungen vorhanden sein:

- die Umgebungstemperatur auf Bauteilebene soll automatisch aus der Umgebungstemperatur auf Ebene Betrachtungseinheit (System) und den frei wählbaren Temperaturdifferenzen zwischen den Baugruppen berechnet werden
- der Einsatzfaktor (Duty - cycle) soll von Baugruppenebene an aufwärts frei gewählt werden können und automatisch berücksichtigt werden ($\lambda = d \lambda_1 + (1-d) \lambda_0$, λ_0 bzw. λ_1 ist die Ausfallrate im- bzw. im eingeschalteten Zustand)

- der Umweltfaktor (π_E), der Belastungsfaktor (S) der Qualitätsfaktor (π_Q) und die Felddaten sollen für ganze Baugruppen oder auch nur für bestimmte Bauteilfamilien pauschal geändert werden können (SET-Anweisungen).
4. Die Berechnung der Ausfallraten von Bauteilen soll nach verschiedenen Ausfallratenkatalogen erfolgen (Vergleichsstudien).
 5. Das Programm muss die bestehenden CA-Werkzeuge des Entwicklers ergänzen und somit über geeignete Interfaces mit diesen kommunizieren können (Übernahme von Bauteildaten bzw. Einlesen von Stücklisten).
 6. Durch eine umfassende, flexible und intelligente Datenausgabe (verschiedene Ausgabemedien, Sortiermöglichkeiten und Darstellungsformen) soll die Lokalisierung von globalen und lokalen Schwachstellen (im Sinne der Zuverlässigkeit) gefördert werden.
 7. Das Programm muss für Erweiterungen offen und weitgehend unabhängig von proprietären Hilfsmitteln sein (dies impliziert z.B. die Errichtung eines eigenen Datenbanksystems); die Portabilität muss gewährleistet sein, sowohl bezüglich des Rechners und Betriebssystems als auch der Programmiersprache.

Das Software-Paket CARP (Computer-Aided Reliability Prediction) der Professur für Zuverlässigkeitstechnik der ETHZ erfüllt alle obigen Voraussetzungen und ist in weit über 100 Firmen eingeführt worden (das Programm ist frei zugänglich). Die Berechnung der vorausgesagten Zuverlässigkeit für nicht reparierbare Systeme basiert auf dem MIL-HDBK-217E. Die Erweiterung durch andere Modelle ist vorgesehen und die Vergleichsberechnung anhand von Felddaten ist implementiert. Darüber hinaus ist in CARP eine auf Erfahrungen im Umgang mit dem MIL-HDBK-217E basierende Berechnungsmethode integriert (MIL korrigiert), in welcher die Faktoren π_E und π_Q fest vorgegeben sind und für π_T und C_1 spezielle Regeln verwendet werden (vgl. Abschnitt 2.1.5). CARP ist sowohl für MS/DOS-PC wie auch für VAX (unter VMS) mit deutscher und englischer Benutzeroberfläche erhältlich. Die Erweiterung zur Untersuchung der Zuverlässigkeit und der Verfügbarkeit reparierbarer Systeme ist in CARAP realisiert. CARAP (Computer-Aided Reliability and Availability Prediction) setzt konstante (zeitunabhängige) Ausfallwerte für jedes Element im System voraus; falls auch die Reparaturrate konstant ist, darf das System über beliebig viele Reparaturmannschaften verfügen (Markoff-Modelle). Im Falle beliebiger Reparaturrate, wird nur eine Reparaturmannschaft zugelassen (semi-regenerative Prozesse). Die Reparaturstrategie (Prioritätsliste) ist wählbar. Ausgehend vom Zuverlässigkeitsblockdiagramm (mit Ausfallrate und Reparaturrate für jedes Element), erzeugt das Programm algorithmisch das Diagramm der Zustandsübergänge und generiert automatisch die Gleichungen für den Mittelwert der ausfallfreien Arbeitszeit $MTTF_S$ und den stationären Wert der Punktverfügbarkeit PA_S . Die Zuverlässigkeitsfunktion $R_S(t)$ kann für einfache Strukturen exakt berechnet werden. Für grosse Serie-/Parallelstrukturen und für komplexe Strukturen verlangt das Programm die zugelassene Abweichung von den genauen Resultaten und entfernt dann automatisch vom Diagramm der Zustandsübergänge die Zustände mit mehr als k Ausfällen ($k = 1, 2, \dots$ je nach der geforderten Genauigkeit).

2.2 Analyse der Art und Auswirkung von Ausfällen

Zur Ausfallratenanalyse (vorausgesagte Zuverlässigkeit) kommt für die kritischen Elemente einer Betrachtungseinheit (eines Systems) stets auch eine Ausfallartenanalyse hinzu. Diese wird mit FMEA (Failure Modes and Effects Analysis) bzw. mit FMECA (Failure Modes, Effects, and Criticality Analysis) bezeichnet und besteht in der systematischen Untersuchung der möglichen Ausfälle bezüglich ihrer Auswirkung auf die Funktionstüchtigkeit und Sicherheit des betreffenden Elements und der von diesem beeinflussten Elemente. Die Untersuchung berücksichtigt die verschiedenen Ausfallarten und Ausfallursachen. Sie ermöglicht die Bestimmung der potentiellen Gefahren und damit die Analyse der Vorkehrungen zur Beseitigung bzw. zur Milderung der Auswirkungen oder zur Verkleinerung der Auftretswahrscheinlichkeit der Ausfälle. Bei der FMEA/FMECA werden oft nicht nur Ausfälle, sondern auch Defekte berücksichtigt. Die Abkürzung FMEA/FMECA steht dann für Fault Modes and Effects Analysis/Fault Modes, Effects, and Criticality Analysis. Eine FMEA/FMECA wird vom Entwicklungsingenieur in Zusammenarbeit mit einem Zuverlässigkeitsingenieur bottom-up durchgeführt. In der Regel beschränkt man sich auf die Schnittstellen und auf die kritischen Elemente. Die Prozedur ist in Tab. 2.4 angegeben (für die Untersuchung mechanischer Betrachtungseinheiten stellt die FMEA/ FMECA eines der wichtigsten Analysewerkzeuge dar).

Neben der FMEA/FMECA ist auch die Fault Tree Analysis (FTA) bekannt, welche eine systematische Untersuchung der Auswirkung von Ausfällen und Fehlern erlaubt. Dabei geht man top-down vom unerwünschten Ereignis (Top event) aus und setzt es mit UND- bzw. ODER-Verknüpfungen von internen Ausfällen oder auch von externen Einflüssen zusammen. Ein Vorteil der FTA ist, dass sie auch Situationen behandeln kann, in welchen das unerwünschte Ereignis auf Ebene Betrachtungseinheit (System) durch das Zusammenwirken mehrerer Ausfälle oder Fehler zustandekommt. Sie ist aber weniger systematisch als die FMEA/FMECA und gibt weniger Gewähr, dass alle Ausfall- bzw. Fehlerarten berücksichtigt worden sind. Die Erfahrung zeigt, dass die FMEA/FMECA und die FTA sich gegenseitig ergänzen.

Die Verfahren der FMEA/FMECA und der FTA sind für die Hardware etabliert, sie lassen sich auch gemischt anwenden (Ursache/Auswirkung-Tabellen bzw. -Diagramme, wie z.B. das SHIKAWA-Diagramm) und auf die Software übertragen.

Tabelle 2.4 Prozedur für die Durchführung einer Failure Modes, Effects, and Criticality Analysis (FMECA)

1. Laufende Numerierung der Schritte.
2. Bezeichnung der betreffenden Elemente (z. B. Transistor, Speisung, Venile usw.) und Kurzbeschreibung ihrer Funktionen. Wenn möglich Referenzangabe zum Zuverlässigkeitsblockdiagramm.
3. Annahme einer möglichen Ausfallart. (Dabei ist oft zu berücksichtigen, in welcher Betriebsphase einer Mission sich die Betrachtungseinheit befindet, denn ein Ausfall oder ein Fehler in einer frühen Betriebsphase kann einen Einfluss auf die gerade untersuchte Betriebsphase haben.)
4. Kurzbeschreibung der möglichen Ursachen für die in (3) angenommene Ausfallart. Die Identifikation der Ursachen ist notwendig, um die Auftretswahrscheinlichkeit schätzen oder berechnen zu können (9) und um Verhütungs- oder Kompensationsmassnahmen zu untersuchen (7). Eine Ausfallart (Kurzschluss, Unterbrechung, Drift usw.) kann mehrere Ursachen haben; ferner kann es sich um einen Primär- oder um einen Folgeausfall handeln. Alle unabhängigen Ursachen müssen identifiziert und untersucht werden.
5. Beschreibung der Symptome, mit welchen sich die in (3) angenommene Ausfallart manifestiert, sowie der Möglichkeiten zur Lokalisierung des Ausfalls. Ferner Kurzbeschreibung der lokalen Auswirkung des Ausfalls auf das betreffende Element und auf die Elemente, die in Beziehung mit diesem stehen (beispielsweise Input/Output).
6. Kurzbeschreibung der Auswirkungen der in (3) angenommenen Ausfallart auf die ganze Betrachtungseinheit in bezug auf die Sicherheit und auf die Erfüllung der geforderten Funktion.
7. Kurzbeschreibung der Vorkehrungen, welche die Auswirkung des Ausfalls mildern, die Auftretswahrscheinlichkeit verkleinern oder die Weiterführung der Mission bzw. der geforderten Funktion erlauben.
8. Gewichtung der Auswirkung der in (3) angenommenen Ausfallart auf die Sicherheit und auf die Erfüllung der geforderten Funktion der ganzen Betrachtungseinheit. Die Bewertungsziffer wird in der Regel gemäss folgender Skala festgelegt: 1 = praktisch keine Auswirkung (sicher), 2 = Teilausfall (unkritisch), 3 = Vollausfall (kritisch), 4 = überkritischer Ausfall (katastrophal). Die Bewertung erfolgt mit Ingenieurgefühl.
9. Berechnung oder Schätzung der Auftretswahrscheinlichkeit der in (3) angenommenen Ausfallart unter Berücksichtigung der in (4) identifizierten Ausfallursachen. (Anstelle der Auftretswahrscheinlichkeit kann auch die Ausfallrate angegeben werden.) Eine für Sicherheitsanalysen übliche Abstufung der Auftretswahrscheinlichkeit ist A = häufig, B = wahrscheinlich, C = wenig wahrscheinlich, D = unwahrscheinlich, E sehr unwahrscheinlich.
10. Zusammenfassung von Bemerkungen oder Anregungen zu den Angaben der früheren Punkte, zur Einführung von Korrekturmassnahmen usw.

3 Zuverlässigkeit und Verfügbarkeit reparierbarer Geräte und Systeme

Die Untersuchung des Zeitverhaltens reparierbarer Geräte und Systeme erfolgt mit Hilfe der Theorie der stochastischen Prozesse. Dabei wird in der Regel vom Zuverlässigkeitsblockdiagramm und von den Verteilungsfunktionen der ausfallfreien Arbeitszeiten und der Instandsetzungszeiten ausgegangen. Die Zuverlässigkeit und die Verfügbarkeit werden in Funktion der Zeit berechnet. Enthält die Betrachtungseinheit keine Redundanz, bleibt die Zuverlässigkeitsfunktion dieselbe, ob man die Instandsetzbarkeit berücksichtigt oder nicht. Anders ist hingegen, wenn redundante Teile vorhanden sind. Hier kann in der Regel die geforderte Funktion auch während einer Instandsetzung weiter ausgeführt werden. Bei den Verfügbarkeitsanalysen werden im Gegensatz zu den Zuverlässigkeitsanalysen Betriebsunterbrechungen zugelassen. Je nach Art der erlaubten Unterbrechungen werden verschiedene Verfügbarkeitsarten definiert. Die wichtigsten davon werden in Abschnitt 3.1 eingeführt. Abschnitt 3.2 behandelt den Fall einer Betrachtungseinheit ohne Redundanz, die Abschnitte 3.3 und 3.4 Parallelmodelle und die Abschnitte 3.5 und 3.6 einfache Serie-Parallelstrukturen. Die wichtigsten Resultate sind in den Tabellen 3.1 bis 3.5 und Bildern 3.8 und 3.12 gegeben.

Zur Vereinheitlichung der Modelle und zur Vereinfachung der Berechnungen werden folgende Annahmen getroffen:

1. Das System wechselt ständig zwischen Arbeits- und Reparaturzustand (Dauerbetrieb); mit Ausnahme der Untersuchungen im stationären Zustand (insbesondere für die Verfügbarkeit) wird angenommen, dass zur Zeit $t = 0$ das System neu ist (Zustand Z_0); zur Verdeutlichung werden alle berechnete Zuverlässigkeitsgrößen mit den Indices S_0 versehen (S für System und 0 für Ausgangszustand Z_0 bei $t = 0$).
2. Während der Reparatur eines Ausfalls auf Systemebene können keine weiteren Ausfälle mehr auftreten.
3. Es steht nur eine Reparaturmannschaft zur Verfügung; die Reparatur wird begonnen sobald die Reparaturmannschaft frei (von einer früheren Reparatur) ist.
4. Redundante Elemente werden auf Systemebene ohne Betriebsunterbrechung repariert.
5. Nach jeder Reparatur ist das reparierte Element (im Zuverlässigkeitsblockdiagramm) neuwertig.
6. Die ausfallfreien Arbeitszeiten und Reparaturzeiten jedes Elements sind stetig, statistisch unabhängig, positiv und haben einen endlichen Mittelwert und eine endliche Varianz; der Mittelwert der ausfallfreien Arbeitszeiten wird mit $MTTF$ (Mean Time To Failure), jener der Reparaturzeiten mit $MTTR$ (Mean Time To Repair) bezeichnet.
7. Der Einfluss der Wartung und der Umschalteneinrichtungen wird vernachlässigt.

Diese Annahme treffen in der Praxis oft zu. Kritisch zu überprüfen ist Annahme 5, vor allem wenn das reparierte Element aus einer Kette von Teilen besteht, die bei einem Ausfall nur teilweise ersetzt werden (sie ist erfüllt, falls die nicht ersetzten Teile eine konstante Ausfallrate aufweisen). Die Annahme 2 vereinfacht die Berechnungen (reduziert die Anzahl Zustände im Zustandsraum); sie ist für manche rein elektronische Systeme (Computersysteme) nur beschränkt

erfüllt, liefert allerdings gute Schätzungen der realen Verhältnisse, weil die Wahrscheinlichkeit für mehrfache Ausfälle auf Systemebene in der Regel sehr klein ist.

3.1 Das Einzelement

Ein Einzelement ist eine Anordnung beliebiger Komplexität (Bauteil, Baugruppe, Gerät, System), welches für Zuverlässigkeitsuntersuchungen als eine Einheit betrachtet wird. Sein Zuverlässigkeitsblockdiagramm besteht aus einem einzi-

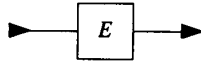


Bild 3.1 Zuverlässigkeitsblockdiagramm eines Einzelements

Das Zeitverhalten eines solchen Elements (neu zur Zeit $t=0$) ist im Bild 3.2 gegeben.

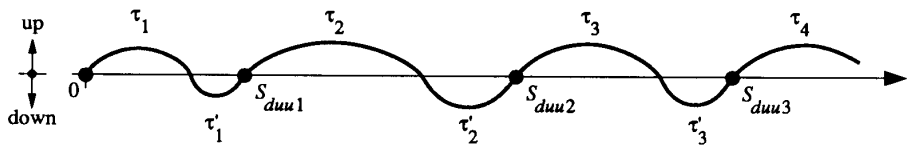


Bild 3.2 Möglicher Zeitverlauf einer reparierbaren Betrachtungseinheit bestehend aus einem einzigen Element (Instandsetzungszeiten übertrieben gross eingezeichnet).

$\tau_1, \tau_2, \dots$ sind die *ausfallfreien Arbeitszeiten*. Sie sind statistisch unabhängige Zufallsgrößen und besitzen die Verteilungsfunktion

$$F(t) = \Pr\{\tau_i \leq t\}, \quad i = 1, 2, \dots \quad (3.1)$$

$\tau'_1, \tau'_2, \dots$ sind die *Instandsetzungszeiten* (Reparaturzeiten). Sie sind ebenfalls statistisch unabhängige Zufallsgrößen; ihre Verteilungsfunktion wird mit $G(t)$ bezeichnet.

$$G(t) = \Pr\{\tau'_i \leq t\}, \quad i = 1, 2, \dots \quad (3.2)$$

Infolge der Annahme 5 ist die Betrachtungseinheit zu den Zeitpunkten $0, S_{duu1}, S_{duu2}, \dots$ neuwertig. Diese Zeitpunkte sind damit *Erneuerungspunkte*; ihr Vorhandensein erleichtert die Untersuchungen wesentlich. In der Theorie der stochastischen Prozesse nennt man den Prozess bestehend aus den Zeitpunkten $0, S_{duu1}, S_{duu2}, \dots$ einen (gewöhnlichen) Erneuerungsprozess; der Prozess bestehend aus $\tau_1, \tau'_1, \tau_2, \tau'_2, \dots$ ist ein (gewöhnlicher) alternierender Erneuerungsprozess.

1. Zuverlässigkeitsfunktion $R_{S0}(t)$: Die Zuverlässigkeitsfunktion ist gleich der Wahrscheinlichkeit, dass die Betrachtungseinheit während dem Intervall $(0, t]$ die geforderte Funktion unter den gegebenen Arbeitsbedingungen ausführt. Es ist also

$$R_{S0}(t) = \Pr\{\text{kein Ausfall in } (0, t] \mid \text{neu zur Zeit } t = 0\} \quad (3.3)$$

gen Element (Bild 3. 1).

und damit

$$R_{S0}(t) = 1 - F(t). \quad (3.4)$$

2 . Punkt-Verfügbarkeit $PA_{S0}(t)$: Die Punkt-Verfügbarkeit ist gleich der Wahrscheinlichkeit, dass die Betrachtungseinheit zum Zeitpunkt t die geforderte Funktion unter den gegebenen Arbeitsbedingungen ausführt. Es ist also

$$PA_{S0}(t) = \Pr\{\text{im Arbeitszustand zur Zeit } t \mid \text{neu zur Zeit } t = 0\}, \quad (3.5)$$

und es gilt [2, 1991]

$$PA_{S0}(t) = 1 - F(t) + \int_0^t h_{duu}(x)(1 - F(t - x)) dx. \quad (3.6)$$

$h_{duu}(t)\delta t$ gibt die Wahrscheinlichkeit an, dass S_{duu1} oder S_{duu2} oder S_{duu3} oder ... im Intervall $(t, t+\delta t]$ fällt. Man kann zeigen, dass

$$\lim_{t \rightarrow \infty} PA_{S0}(t) = PA = \frac{MTTF}{MTTF + MTTR} \quad (3.7)$$

gilt. $MTTF$ und $MTTR$ sind dabei die Erwartungswerte der ausfallfreien Arbeitszeiten $\tau_1, \tau_2, \dots$ und der Instandsetzungszeiten $\tau'_1, \tau'_2, \dots$. Gl. (3.7) gilt für jede Verteilung von τ und von τ' . Für $F(t) = 1 - e^{-\lambda t}$ und $G(t) = 1 - e^{-\mu t}$, d.h. für eine konstante Ausfallrate $\lambda = 1 / MTBF$ und Reparaturrate $\mu = 1 / MTTR$ gilt

$$PA_{\mu}(t) = \frac{\mu}{\lambda + \mu} + \frac{\lambda}{\lambda + \mu} e^{-(\lambda + \mu)t}.$$

Die Konvergenz von $PA_{S0}(t)$ gegen den Wert PA_S erfolgt mit einer Zeitkonstante $\approx MTTR$. Das Resultat der Gl. (3.7) kann mit der geometrischen Deutung der Wahrscheinlichkeit leicht interpretiert werden. PA_S ist ein asymptotischer Wert und gleichzeitig auch der Wert der Punkt-Verfügbarkeit für alle $t \geq 0$ im *stationären Zustand*. Ein stationärer Zustand liegt vor, wenn

die Betrachtungseinheit lange vor dem Zeitpunkt $t=0$ (praktisch für $t < -10 MTTR$) eingeschaltet wurde und erst ab $t=0$ beobachtet wird ($t=0$ ist ein willkürlicher Zeitpunkt).

Bei $t = 0$ ist dann die Betrachtungseinheit mit der Wahrscheinlichkeit PA_S im Arbeitszustand und mit der Wahrscheinlichkeit $1 - PA_S$ im Instandsetzungszustand.

3 . Durchschnittliche Verfügbarkeit $AA_{S0}(t)$: Die durchschnittliche Verfügbarkeit ist gleich dem erwarteten Prozentsatz der Missionsdauer t , währenddem die Betrachtungseinheit die geforderte Funktion unter den vorgegebenen Arbeitsbedingungen ausführt. Es ist also

$$AA_{S0}(t) = \frac{1}{t} E[\text{totale Betriebszeit in } (0, t] \mid \text{neu zur Zeit } t = 0], \quad (3.8)$$

und es gilt

$$AA_{S0}(t) = \frac{1}{t} \int_0^t PA_{S0}(x) dx \quad (3.9)$$

mit $PA_{S0}(t)$ aus Gl. (3.6). Für den stationären Zustand gilt

$$AA_S = PA_S = \frac{MTTF}{MTTF + MTTR} \approx 1 - \frac{MTTR}{MTTF} \quad (3.10)$$

- 4 . Gesamtverfügbarkeit OA_S :** Die Gesamtverfügbarkeit (Overall-Availability) wird als Verhältnis der totalen Betriebszeit zur Summe der totalen Betriebszeit und der totalen Instandhaltungszeit in einem Intervall $(0, T)$ für $T \rightarrow \infty$ definiert. Mit den Abkürzungen $MTTF$ = Mittelwert der ausfallfreien Arbeitszeit, MDT = Mittelwert der Ausfallzeiten (Reparatur plus Logistik), $MTTPM$ = Mittelwert der Wartungszeiten und T_{PM} = Wartungsperiode folgt für die Gesamtverfügbarkeit

$$OA_S = \frac{MTTF}{MTTF + MDT} = \frac{MTTF}{MTTF + MTTR + \frac{MTTF}{T_{PM}} MTTPM}. \quad (3.11)$$

- 5 . Intervall-Zuverlässigkeit $IR_{S0}(t, t+\theta)$:** Die Intervall-Zuverlässigkeit ist gleich der Wahrscheinlichkeit, dass die Betrachtungseinheit während dem Intervall $(t, t+\theta]$ die geforderte Funktion unter den vorgegebenen Arbeitsbedingungen ausführt. Es ist also

$$IR_{S0}(t, t+\theta) = \Pr\{\text{im Arbeitszustand in } (t, t+\theta] \mid \text{neu zur Zeit } t=0\} \quad (3.12)$$

und es gilt [2, 1991]

$$IR_{S0}(t, t+\theta) = 1 - F(t+\theta) + \int_0^t h_{duu}(x)(1 - F(t+\theta-x)) dx. \quad (3.13)$$

Für $F(t) = 1 - e^{-\lambda t}$, d.h. im Falle einer konstanten Ausfallrate λ , erhält man

$$IR_{S0}(t, t+\theta) = PA_{S0}(t) \cdot e^{-\lambda \theta}. \quad (3.14)$$

und insbesondere, im stationären Zustand

$$IR_S(\theta) = PA_S \cdot e^{-\lambda \theta} \quad (3.15)$$

mit PA_S aus Gl. (3.7). Die Produktregel, welche in den Gln. (3.14) und (3.15) zur Anwendung kommt, kann nur für den Fall der exponentiell-verteilten ausfallfreien Arbeitszeit (konstante Ausfallrate) verwendet werden.

- 6 . Missions-Verfügbarkeit $MA_{S0}(T_0, t_f)$:** Die Missions-Verfügbarkeit ist gleich der Wahrscheinlichkeit, dass eine Mission mit der totalen Betriebszeit T_0 keinen Ausfall haben wird, der nicht innerhalb der Zeitspanne t_f repariert werden kann. Man kann zeigen, [2,1991], dass für $F(t) = 1 - e^{-\lambda t}$ gilt

$$MA_{S0}(T_0, t_f) = e^{-\lambda T_0 (1-G(t_f))}. \quad (3.16)$$

Die Missions-Verfügbarkeit spielt eine Rolle bei den Fällen, wo Betriebsunterbrechungen $\leq t_f$ toleriert werden können.

Tabelle 3.1 fasst die wichtigsten Resultate für das reparierbare Einzelelement im stationären Zustand zusammen.

Tabelle 3.1 Wichtigste Resultate für das reparierbare Einzelelement im stationären Zustand

	Ausfall- und Reparaturrate		Definition Bemerkungen
	beliebig	konstant	
1. $\Pr\{\text{im Arbeitszu-}\br/> \text{stand zur Zeit}\br/> t = 0\} \quad (p)$	$\frac{MTTF}{MTTF + MTTR}$	$\frac{\mu}{\lambda + \mu}$	$MTTF = E[\tau_i], i \geq 1$ $MTTR = E[\tau_i'], i \geq 1$
2. Verteilungsfunktion der ersten ausfall- freien Arbeitszeit (ab $t = 0$)	$\frac{1}{MTTF} \int_0^\infty (1 - F(x)) dx$	$1 - e^{-\lambda t}$	Für die erste Reparaturzeit (ab $t = 0$) muss man MTTF mit MTTR und $F(x)$ mit $G(x)$ ersetzen
3. Punkt- Verfügbarkeit	$\frac{MTTF}{MTTF + MTTR}$	$\frac{\mu}{\lambda + \mu}$	$PAS = \Pr\{\text{im Arbeitszustand zur}\br/> \text{Zeit } t\}, t \geq 0$
4. Durchschnittliche Verfügbarkeit	$\frac{MTTF}{MTTF + MTTR}$	$\frac{\mu}{\lambda + \mu}$	$AAS = (1/t) E[\text{totale Betriebszeit}\br/> \text{in } (0, t)], t > 0$
5. Intervall- Zuverlässigkeit	$\frac{1}{MTTF + MTTR} \cdot \int_0^\infty (1 - F(x)) dx$	$\frac{\mu}{\lambda + \mu} e^{-\lambda \theta}$	$IRS(\theta) = \Pr\{\text{im Arbeitszustand in}\br/> (t, t+\theta]\}, t \geq 0$

λ = konstante Ausfallrate ($MTTF = 1/\lambda = MTBF$) μ = konstante Reparaturrate ($MTTR = 1/\mu$)

3.2 Geräte und Systeme ohne Redundanz

Das Zuverlässigkeitsblockdiagramm einer Betrachtungseinheit ohne Redundanz besteht in der Serieschaltung aller lebenswichtigen Elemente (Bild 3.3).



Bild 3.3 Zuverlässigkeitsblockdiagramm eines Systems ohne Redundanz (Seriemodell)

Für die Untersuchungen geht man oft von der Annahme aus, dass jedes Element eine konstante Ausfallrate λ_{0i} und eine konstante Reparaturrate μ_{i0} besitzt und unabhängig von den anderen Elementen arbeitet und ausfällt. Infolge der konstanten Ausfall- und Reparaturraten lässt sich das Zeitverhalten des Systems durch einen zeithomogenen Markoffprozess beschreiben (vgl. Anhang A2.2 für die Darlegung der Konzepte). Bild 3.4 gibt das entsprechende Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ an.

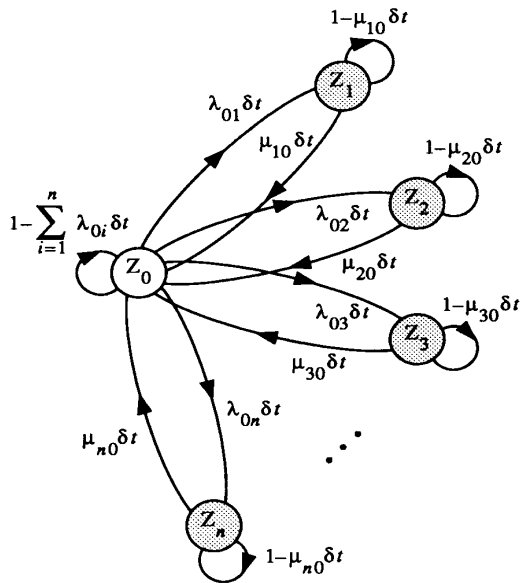


Bild 3.4 Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ eines Systems ohne Redundanz (konst. Ausfall- und Reparaturraten, kein weiterer Ausfall im Ausfallzustand auf Systemebene, t beliebig, $\delta t \rightarrow 0$)

1 . Zuverlässigkeitsfunktion $R_{S0}(t)$: Die Betrachtungseinheit wird im Intervall $(0,t)$ nur dann ausfallfrei arbeiten, wenn kein Element in diesem Zeitintervall ausfällt. Für $R_{S0}(t)$ folgt

$$R_{S0}(t) = e^{-\lambda_{S0} t} \quad \text{mit} \quad \lambda_{S0} = \sum_{i=1}^n \lambda_{0i}, \quad (3.17)$$

und für den Mittelwert der ausfallfreien Arbeitszeit $MTTF_{S0}$

$$MTTF_{S0} = \frac{1}{\lambda_{S0}}. \quad (3.18)$$

2 . Punkt-Verfügbarkeit $PA_{S0}(t)$: Ein System von Differentialgleichungen für die Berechnung von $PA_{S0}(t)$ kann mit Hilfe von Bild 3.4 und der Tab. A2.2 aufgestellt werden. Für den stationären Fall erhält man dann insbesondere

$$PA_S = AA_S = \frac{1}{1 + \sum_{i=1}^n \frac{\lambda_{0i}}{\mu_{i0}}} \approx 1 - \sum_{i=1}^n \frac{\lambda_{0i}}{\mu_{i0}}. \quad (3.19)$$

Die Näherungsformel auf der rechten Seite der Gl. (3.19) gilt für $\lambda_{0i} \ll \mu_{i0}$ (was in praktisch allen Anwendungen erfüllt ist, weil $\lambda_{0i}/\mu_{i0} = \text{MTTR}_i/\text{MTBF}_i$) und ist identisch mit jener die man bei der Berechnung der Punkt-Verfügbarkeit unter der Annahme, dass jedes Element unabhängig von den anderen arbeitet und repariert wird (eine Reparaturmannschaft pro Element), erhalten würde

$$PA_S = \prod_{i=1}^n PA_i = \prod_{i=1}^n \frac{1}{1 + \lambda_{0i} / \mu_{i0}} \approx 1 - \sum_{i=1}^n \frac{\lambda_{0i}}{\mu_{i0}}. \quad (3.20)$$

3. Intervall-Zuverlässigkeit $IR_{S0}(t, t+\theta)$: Infolge der konstanten Ausfallrate aller Elemente gilt

$$IR_{S0}(t, t+\theta) = PA_{S0}(t) e^{-\lambda_{S0} \theta}, \quad (3.21)$$

und für den asymptotischen und stationären Wert

$$IR_S(\theta) = PA_S e^{-\lambda_{S0} \theta}. \quad (3.22)$$

Tabelle 3.2 fasst die wichtigsten Resultate für ein Gerät oder System ohne Redundanz (Seriemodell) bestehend aus den Elementen E_1 bis E_n .

Tabelle 3.2 Wichtigste Resultate für reparierbare Systeme ohne Redundanz

Grösse	Ausdruck	Bemerkungen
1. Zuverlässigkeitsfunktion $R_{S0}(t)$	$\prod_{i=1}^n R_i(t)$	jedes Element arbeitet unabhängig von den anderen Elementen
2. Mittelwert der ausfallfreien Arbeitszeit $MTTF_{S0}$	$\int_0^{\infty} R_{S0}(t) dt$	$R_i(t) = e^{-\lambda_{0i} t}$ ergibt $R_{S0}(t) = e^{-\lambda_{S0} t}$, mit $\lambda_{S0} = \lambda_{01} + \dots + \lambda_{0n}$ und $MTTF_{S0} = 1 / \lambda_{S0}$
3. Ausfallrate $\lambda_{S0}(t)$	$\sum_{i=1}^n \lambda_{0i}(t)$	jedes Element arbeitet und fällt, unabhängig von den anderen Elementen aus
4. Punkt-Verfügbarkeit und durchschnittliche Verfügbarkeit im stationären Zustand (eine Reparaturmannschaft) $PA_S = AA_S$	1) $\frac{1}{1 + \sum_{i=1}^n \frac{\lambda_{0i}}{\mu_{i0}}} \approx 1 - \sum_{i=1}^n \frac{\lambda_{0i}}{\mu_{i0}}$ 2) $\frac{1}{1 + \sum_{i=1}^n \lambda_{0i} MTTR_i} \approx 1 - \sum_{i=1}^n \lambda_{0i} MTTR_i$	im Ausfallzustand kann kein weiterer Ausfall eintreten: 1) konstante Ausfallrate (λ_{0i}) und Reparaturrate (μ_{i0}) jedes Elements 2) konstante Ausfallrate (λ_{0i}) und bleibige Reparaturrate jedes Elements, $MTTR_i$ = Mittelwert der Reparaturzeit des i -ten Element
5. Intervall-Zuverlässigkeit im stationären Zustand $IR_S(\theta)$	$PA_S e^{-\lambda_{S0} \theta}$	Die Ausfallrate jedes Elements ist konstant ($\lambda_{S0} = \lambda_{01} + \dots + \lambda_{0n}$)

3.3 Redundanz 1 aus 2

3.7

Eine Redundanz 1 aus 2 stellt die einfachste Redundanzstruktur dar, die man in den Anwendungen finden kann. Sie besteht aus zwei Elementen E_1 und E_2 , eines davon im Arbeitszustand, das

andere im Reservezustand. Fällt das Arbeitselement aus, so geht dieses in den Reparaturzustand über, und das Reserveelement übernimmt die Erfüllung der geforderten Funktion. Ein Ausfall auf Systemebene tritt nur dann auf, wenn während der Reparatur eines Elements auch das andere Element noch ausfällt. Das entsprechende Zuverlässigkeitsblockdiagramm ist in Bild 3.5 gezeigt.

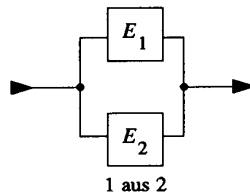


Bild 3.5 Zuverlässigkeitsblockdiagramm einer Redundanz 1 aus 2

Die Untersuchungen setzen die Annahmen 1 bis 7 voraus (Einleitung zum Kapitel 3). Insbesondere gilt, dass die Instandsetzung eines redundanten Elements *ohne Betriebsunterbrechung* auf Systemebene erfolgen kann und implizit, dass sie (theoretisch) sofort nach dem Auftreten des Ausfalls beginnt.

Infolge der konstanten Ausfall- und Reparaturrate lässt sich das Zeitverhalten des Systems gemäss Bild 3.5 mittels eines zeithomogenen Markoff-Prozesses untersuchen (vgl. Anhang A2.2 für die Darlegung der Konzepte). Die Anzahl Zustände ist fünf, wenn die Elemente E_1 und E_2 verschieden sind (Bild 3.6) und drei, wenn sie identisch sind (Bild 3.7).

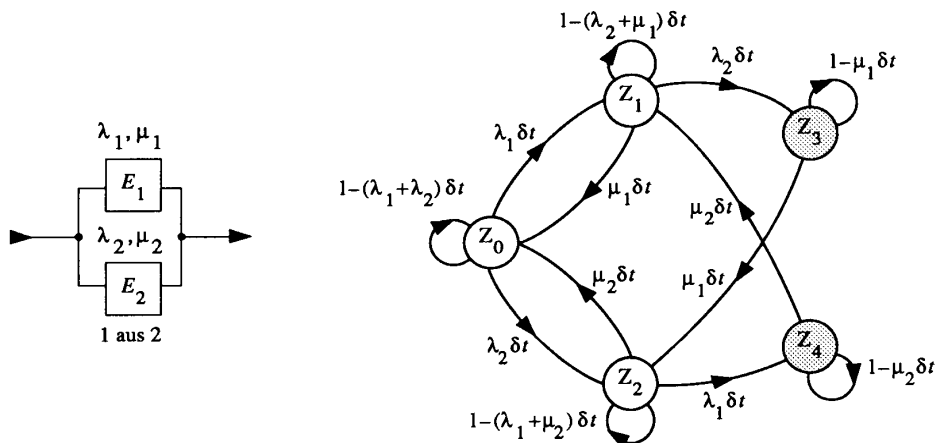


Bild 3.6 Zuverlässigkeitsblockdiagramm und Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ für eine heisse Redundanz 1 aus 2 mit 2 verschiedenen Elementen mit konstanten Ausfallraten λ_1, λ_2 und konstanten Reparaturraten μ_1, μ_2 (eine Reparaturmannschaft, t beliebig, $\delta t \rightarrow 0$)

Im Falle identischer Elemente es seien, wie im Bild 3.7b, λ die konstante Ausfallrate im Arbeitszustand, λ_r die konstante Ausfallrate im Reservezustand und μ die konstante Reparaturrate. Man beachte, dass $\lambda_r = \lambda$ der Fall der heissen (parallelen) Redundanz und $\lambda_r \equiv 0$ der Fall der kalten (standby) Redundanz darstellen. Das System ist ausgefallen (down) im Zustand Z_2 bzw. Z_2 .

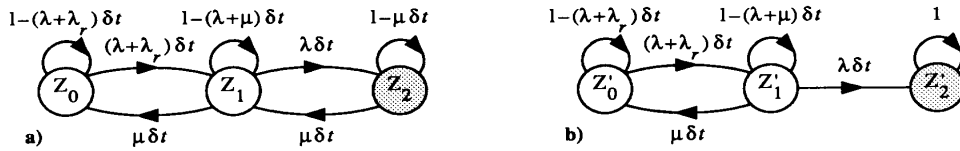


Bild 3.7 a) Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ einer warmen Redundanz 1 aus 2 (gleiche Elemente, eine Reparaturmannschaft, konstante Ausfall- und Reparaturraten (λ bzw. λ_r und μ), t beliebig, $\delta t \rightarrow 0$); b) wie Punkt a, aber zur Berechnung der Zuverlässigkeit

Aus Bild 3.7a bzw. 3.7b und mit Hilfe der Tab. A2.2 können die Differentialgleichungen für die Zustandswahrscheinlichkeiten

$$P_i(t) = \Pr\{\text{der Prozeß ist zur Zeit } t \text{ im Zustand } Z_i\}, \quad i = 0, 1, 2 \quad (3.23)$$

bzw. für $P'_i(t)$ aufgestellt werden, vgl. die Beispiele A2.8 und A2.9

1. Punkt-Verfügbarkeit $PA_{S0}(t)$: Wird in Zusammenhang mit Bild 3.7a als Anfangsbedingung $P_0(0) = 1$ und $P_1(0) = P_2(0) = 0$ gewählt, so gehen die Zustandswahrscheinlichkeiten $P_0(t)$, $P_1(t)$ und $P_2(t)$ in die bedingten Zustandswahrscheinlichkeiten $P_{00}(t) \equiv P_0(t)$, $P_{01}(t) \equiv P_1(t)$ und $P_{02}(t) \equiv P_2(t)$ über. Die Punkt-Verfügbarkeit $PA_{S0}(t)$ ist dann gegeben durch

$$PA_{S0}(t) = P_{00}(t) + P_{01}(t). \quad (3.24)$$

Für die Laplace-Transformierte von PA_{S0} folgt

$$\tilde{PA}_{S0}(s) = \frac{(s + \mu)(s + \lambda + \lambda_r + \mu) + s\lambda}{s((s + \lambda + \lambda_r)(s + \lambda + \mu) + \mu(s + \mu))}. \quad (3.25)$$

Für $t \rightarrow \infty$ gilt dann

$$\lim_{t \rightarrow \infty} PA_{S0}(t) = PA_S = \frac{\mu(\lambda + \lambda_r + \mu)}{(\lambda + \lambda_r)(\lambda + \mu) + \mu^2} \approx 1 - \frac{\lambda(\lambda + \lambda_r)}{\mu(\lambda + \lambda_r + \mu)}. \quad (3.26)$$

PA_S ist auch der Wert der Punkt-Verfügbarkeit und der durchschnittlichen Verfügbarkeit für alle $t \geq 0$ im stationären Zustand.

2. Zuverlässigkeitsfunktion $R_{S0}(t)$: Für die Berechnung der Zuverlässigkeitsfunktion muss berücksichtigt werden, dass die Redundanz 1 aus 2 im Intervall $(0, t]$ nur dann ausfallfrei arbeitet, wenn in diesem Intervall der Zustand Z_2 nicht angenommen wird. Um festzustellen, ob während des Intervalls $(0, t]$ der Zustand Z_2 angenommen wurde, wird Z_2 absorbierend gemacht. Dies führt zum Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ in Bild 3.7b. Die Berechnungen sind in derselben Weise wie für die Punkt-Verfügbarkeit durchzuführen. Falls zur Zeit $t = 0$ das System im Zustand Z_0 ist, folgt für die Zuverlässigkeitsfunktion

$$R_{S0}(t) = P'_{00}(t) + P'_{01}(t) \quad (3.27)$$

mit der Laplace-Transformierten

$$\tilde{R}_{S0}(s) = \frac{s + 2\lambda + \lambda_r + \mu}{(s + \lambda + \lambda_r)(s + \lambda) + s\mu}. \quad (3.28)$$

Der Mittelwert der ausfallfreien Arbeitszeit ist dann gegeben durch ($MTTF_{S0} = \tilde{R}_{S0}(0)$)

$$MTTF_{S0} = \frac{2\lambda + \lambda_r + \mu}{\lambda(\lambda + \lambda_r)}. \quad (3.29)$$

Für $\lambda \ll \mu$ kann man zeigen [2,1991], dass $R_{S0}(t)$ durch eine Exponentialfunktion angenähert werden kann

$$R_{S0}(t) \approx e^{-\lambda_{S0} t} \quad \text{mit} \quad \lambda_{S0} = \frac{1}{MTTF_{S0}} = \frac{\lambda(\lambda + \lambda_r)}{2\lambda + \lambda_r + \mu}. \quad (3.30)$$

Dieses Resultat ist für die Untersuchung von komplexer Strukturen (insbesondere mit Hilfe der Methode der Makrostrukturen) wichtig; es zeigt (am Beispiel der heißen Redundanz $\lambda_r = \lambda$), dass

für $\lambda \ll \mu$ eine reparierbare Redundanz 1 aus 2 mit konstanter Ausfallrate λ und Reparaturrate μ näherungsweise durch ein Einzelelement mit konstanter Ausfallrate $\approx 2\lambda^2 / (3\lambda + \mu)$ ersetzt werden kann (eine äquivalente Reparaturrate kann aus dem Vergleich der Gleichungen für den stationären Wert der Punkt-Verfügbarkeit hergeleitet werden, vgl. Tab. 3.5).

Die Gl. (3.30) enthält offenbar auch die Fälle der heißen Redundanz $\lambda_r = \lambda$ und der kalten Redundanz $\lambda_r = 0$.

3 . Intervall-Zuverlässigkeit $IR_{S0}(t, t+\theta)$: Wegen der Gedächtnislosigkeit des zeithomogenen Markoff-Prozesses lässt sich die Intervall-Zuverlässigkeit aus den Ausdrücken der bedingten Zustandswahrscheinlichkeiten und der Zuverlässigkeitsfunktionen berechnen.

$$IR_{S0}(t, t+\theta) = P_{00}(t) R_{S0}(\theta) + P_{01}(t) R_{S1}(\theta). \quad (3.31)$$

Für den asymptotischen und stationären Wert folgt dann

$$IR_S(\theta) = \frac{\mu^2 R_{S0}(\theta) + \mu(\lambda + \lambda_r) R_{S1}(\theta)}{(\lambda + \lambda_r)(\lambda + \mu) + \mu^2} \approx R_{S0}(\theta). \quad (3.32)$$

Tabelle 3.3 fasst die wichtigsten Resultate für eine reparierbare Redundanz 1 aus 2 mit identischen Elementen zusammen und Bild 3.8 zeigt den Verlauf der Unverfügbarkeit $1 - PA_S$ und des Mittelwerts der ausfallfreien Arbeitszeit $MTTF_{S0}$, normiert auf λ als Funktion von μ/λ und mit λ_r/λ als Parameter ($\lambda_r = \lambda$ stellt eine heiße und $\lambda_r \equiv 0$ eine kalte (standby) Redundanz dar).

Tabelle 3.3 Zuverlässigkeitsfunktion $R_{SO}(t)$ für $\mu \gg \lambda, \lambda_r$, Erwartungswert der ausfallfreien Arbeitszeit $MTTF_{SO}$, Verfügbarkeit $PA_S = AA_S$ und Intervall-Zuverlässigkeit $IR_S(\theta)$ einer reparierbaren Redundanz 1 aus 2 identische Elemente (konst. Ausfall- und Reparaturrate, eine Reparaturmannschaft)

	1 aus 2 heiße Redundanz ($\lambda_r = \lambda$)	1 aus 2 warme Redundanz ($\lambda_r < \lambda$)	1 aus 2 kalte Redundanz ($\lambda_r = 0$)
$R_{SO}(t)^\#$	$\approx e^{-\frac{2\lambda^2 t}{3\lambda + \mu}}$	$\approx e^{-\frac{\lambda(\lambda + \lambda_r)t}{2\lambda + \lambda_r + \mu}}$	$\approx e^{-\frac{\lambda^2 t}{2\lambda + \mu}}$
$MTTF_{SO}^\#$	$\frac{3\lambda + \mu}{2\lambda^2} \approx \frac{\mu}{2\lambda^2}$	$\frac{2\lambda + \lambda_r + \mu}{\lambda(\lambda + \lambda_r)} \approx \frac{\mu}{\lambda(\lambda + \lambda_r)}$	$\frac{2\lambda + \mu}{\lambda^2} \approx \frac{\mu}{\lambda^2}$
$PA_S = AA_S^*$	$\frac{\mu(2\lambda + \mu)}{2\lambda(\lambda + \mu) + \mu^2} \approx 1 - 2\left(\frac{\lambda}{\mu}\right)^2$	$\frac{\mu(\lambda + \lambda_r + \mu)}{(\lambda + \lambda_r)(\lambda + \mu) + \mu^2} \approx 1 - \frac{\lambda(\lambda + \lambda_r)}{\mu^2}$	$\frac{\mu(\lambda + \mu)}{\lambda(\lambda + \mu) + \mu^2} \approx 1 - \left(\frac{\lambda}{\mu}\right)^2$
$IR_S(\theta)^*$	$\approx R_{SO}(\theta)$	$\approx R_{SO}(\theta)$	$\approx R_{SO}(\theta)$

λ, λ_r = Ausfallrate (r = Reserve); μ = Reparaturrate; $^\#$ neu zur Zeit $t = 0$; * asymptotischer und stationärer Wert

Aus Tab. 3.3 ist der durch die Möglichkeit einer Reparatur ohne Betriebsunterbrechung erzielbare Gewinn für die $MTTF_{SO}$ ersichtlich. Für den Fall heißer Redundanz gilt (mit $1/\lambda = MTBF$ und $1/\mu = MTTR$)

	Einzelement	1 aus 2 nicht reparierbar	1 aus 2 reparierbar
$\frac{MTTF_S}{1/\lambda} =$	1	1.5	$\frac{\mu}{2\lambda} = \frac{MTBF}{2 MTTR}$

Obige Untersuchungen können leicht auch auf andere Zusammenstellungen der Übergangsraten bedingt durch Aufteilung der Last, Unterschiede in den Elementen, Priorität in der Reparatur usw. ausgedehnt werden.

3.4 Redundanz k aus n

Eine Redundanz k aus n besteht aus der Parallelschaltung von n Elementen, wovon k für die Erfüllung der geforderten Funktion notwendig sind und $n-k$ in Reserve stehen. Bild 3.9 zeigt das entsprechende Zuverlässigkeitsblockdiagramm, unter der Annahme einer hundertprozentig zuverlässigen Umschalteneinrichtung.

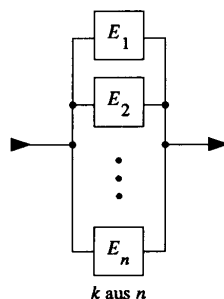


Bild 3.9 Zuverlässigkeitsblockdiagramm einer Redundanz k aus n mit idealer Umschalteneinrichtung

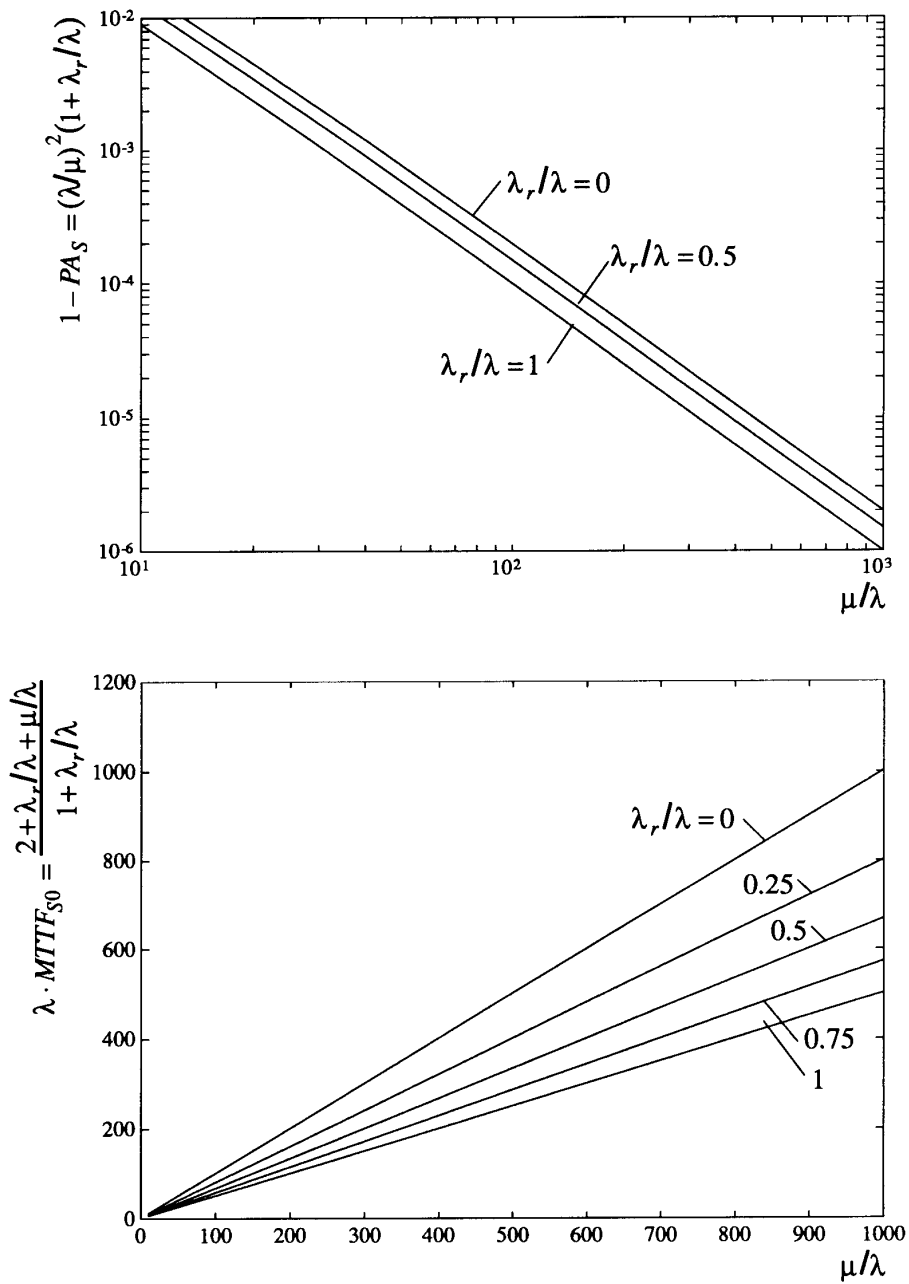


Bild 3.8 Verlauf der Unverfügbarkeit $1 - PA_S$ und des Mittelwerts der ausfallfreien Arbeitszeit $MTTF_{S0}$ normiert auf λ für eine reparierbare Redundanz 1 aus 2 (identische Elemente, konst. Ausfallrate λ und λ_r (r = Reserve), konst. Reparaturrate μ , eine Reparaturmannschaft, neu zur Zeit $t = 0$ für $MTTF_{S0}$)

Für die folgenden Untersuchungen wird von der praxisnahen Annahme ausgegangen, die n Elemente seien identisch. Untersucht wird der Fall einer warmen Redundanz. Unter der Annahme einer konstanten Ausfallrate λ , bzw. λ_r (r für Reserve) und Reparaturrate μ für jedes Element kann das Zeitverhalten der Redundanz k aus n mit Hilfe eines *Geburts- und Todesprozesses* untersucht werden. Bild 3.10 gibt das entsprechende Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ an (vgl. Anhang A2.2 für die Darlegung der Konzepte); dabei gilt für die Übergangsraten v_i

$$v_i = k\lambda + (n-k-i)\lambda_r, \quad i = 0, \dots, n-k. \quad (3.33)$$

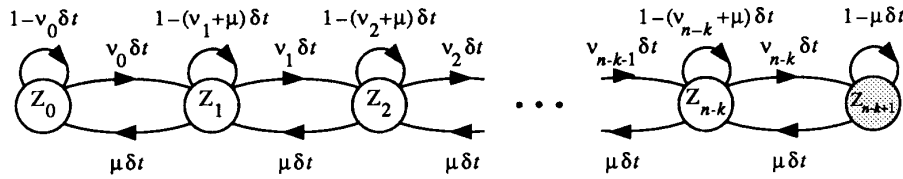


Bild 3.10 Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ einer warmen Redundanz k aus n identische Elemente (konstante Ausfall- und Reparaturraten, eine Reparaturmannschaft, kein weiterer Ausfall während einer Reparatur auf Niveau System, t beliebig, $\delta t \rightarrow 0$)

Aus Bild 3.10 und mit Hilfe der Tab. A2.2 kann man für die Zustandswahrscheinlichkeit $P_i(t)$ gemäss Gl. (3.23) ein System von Differentialgleichungen für die Berechnung der Punkt-Verfügbarkeit bzw. für die Zuverlässigkeitsfunktion aufstellen. Für den stationären Wert der Punkt-Verfügbarkeit erhält man insbesondere

$$PA_S = \sum_{j=0}^{n-k} p_j \quad (3.34)$$

mit

$$p_j = \frac{\pi_j}{\sum_{i=0}^{n-k} \pi_i}, \quad \pi_0 = 1 \quad \text{und} \quad \pi_i = \frac{v_0 \dots v_{i-1}}{\mu^i}. \quad (3.35)$$

PA_S ist auch der asymptotische und stationäre Wert der durchschnittlichen Verfügbarkeit AA_S . Tabelle 3.4 fasst die wichtigsten Resultate für eine reparierbare Redundanz k aus n mit identischen Elementen zusammen. Es zeigt sich, dass die Resultate von $n-k$ abhängen [2,1991]; Tabelle 3.4 konzentriert sich deshalb auf die Fälle $n-k=1$ und $n-k=2$, allgemein und für einige wichtige Spezialfälle.

Tabelle 3.4 Mittelwert der ausfallfreien Arbeitszeit und stationäre Werte der Verfügbarkeit und der Intervall-Zuverlässigkeit einer reparierbaren warmen Redundanz k aus n identische Elemente (konstante Ausfall- und Reparaturraten, eine Reparaturmannschaft, kein weiterer Ausfall während einer Reparatur auf Ebene System)

		Mittelwert der ausfallfreien Arbeitszeit ($MTTF_{S0}$)	Stationärer Wert der Punkt-Verfügbarkeit und der durchschnittlichen Verfügbarkeit ($PA_S = AA_S$)	Stat. Wert der Intervall-Zuverlässigkeit ($IR_S(\theta)$)
$n-k=1$	allg. Fall	$\frac{v_0 + v_1 + \mu}{v_0 v_1} \approx \frac{\mu}{v_0 v_1}$	$\frac{v_0 \mu + \mu^2}{v_0 v_1 + v_0 \mu + \mu^2} \approx 1 - \frac{v_0 v_1}{\mu^2}$	$\approx R_{S0}(\theta)$
	$n=2$ $k=1$	$\frac{2\lambda + \lambda_r + \mu}{\lambda(\lambda + \lambda_r)} \approx \frac{\mu}{\lambda(\lambda + \lambda_r)}$	$\frac{\mu(\lambda + \lambda_r + \mu)}{(\lambda + \lambda_r)(\lambda + \mu) + \mu^2} \approx 1 - \frac{\lambda(\lambda + \lambda_r)}{\mu^2}$	$\approx R_{S0}(\theta)$
	$n=3$ $k=2$	$\frac{4\lambda + \lambda_r + \mu}{2\lambda(2\lambda + \lambda_r)} \approx \frac{\mu}{2\lambda(2\lambda + \lambda_r)}$	$\frac{\mu(2\lambda + \lambda_r + \mu)}{(2\lambda + \lambda_r)(2\lambda + \mu) + \mu^2} \approx 1 - \frac{2\lambda(2\lambda + \lambda_r)}{\mu^2}$	$\approx R_{S0}(\theta)$
$n-k=2$	allg. Fall	$\frac{v_2(v_0 + v_1 + \mu)}{v_0 v_1 v_2} + \frac{\mu(v_0 + \mu) + v_0 v_1}{v_0 v_1 v_2} \approx \frac{\mu^2}{v_0 v_1 v_2}$	$\frac{v_0 v_1 \mu + v_0 \mu^2 + \mu^3}{v_0 v_1 v_2 + v_0 v_1 \mu + v_0 \mu^2 + \mu^3} \approx 1 - \frac{v_0 v_1 v_2}{\mu^3}$	$\approx R_{S0}(\theta)$
	$n=3$ $k=1$	$\frac{\lambda(2\lambda + 3\lambda_r + \mu)}{\lambda(\lambda + \lambda_r)(\lambda + 2\lambda_r)} + \frac{\mu(\lambda + 2\lambda_r + \mu)}{\lambda(\lambda + \lambda_r)(\lambda + 2\lambda_r)} + \frac{1}{\lambda} \approx \frac{\mu^2}{\lambda(\lambda + \lambda_r)(\lambda + 2\lambda_r)}$	$\frac{\mu((\lambda + 2\lambda_r)(\lambda + \lambda_r) + \mu(\lambda + 2\lambda_r) + \mu^2)}{(\lambda + 2\lambda_r)((\lambda + \lambda_r) + \mu(\lambda + \lambda_r) + \mu^2) + \mu^3} \approx 1 - \frac{\lambda(\lambda + \lambda_r)(\lambda + 2\lambda_r)}{\mu^3}$	$\approx R_{S0}(\theta)$
	$n=5$ $k=3$	$\frac{3\lambda(6\lambda + 3\lambda_r + \mu)}{(3\lambda + 2\lambda_r)(3\lambda + \lambda_r)3\lambda} + \frac{\mu(3\lambda + 2\lambda_r + \mu)}{(3\lambda + 2\lambda_r)(3\lambda + \lambda_r)3\lambda} + \frac{1}{3\lambda} \approx \frac{\mu^2}{(3\lambda + 2\lambda_r)(3\lambda + \lambda_r)3\lambda}$	$\frac{\mu(3\lambda + 2\lambda_r)(3\lambda + \lambda_r) + \mu^2}{(3\lambda + 2\lambda_r)(3\lambda + \lambda_r)(3\lambda + \mu) + \mu^2(3\lambda + 2\lambda_r) + \mu^3} \approx 1 - \frac{3\lambda(3\lambda + \lambda_r)(3\lambda + 2\lambda_r)}{\mu^3}$	$\approx R_{S0}(\theta)$

$v_i = k\lambda + (n-k-i)\lambda_r, i = 0, \dots, n-k; \lambda, \lambda_r = \text{Ausfallrate (r = Reserve)}; \mu = \text{Reparaturrate}$

3.5 Einfache Serie-/Parallelstrukturen

Eine Serie-/Parallelstruktur ist eine beliebige Kombination von Serie- und Parallelmodellen (Tab. 2.1). Die Untersuchung erfolgt fallweise und stützt sich auf die Methoden der Abschnitte 3.2 bis 3.4. Falls sich das Zeitverhalten durch zeithomogene Markoff-Prozesse beschreiben lässt, können die entsprechenden Gleichungen mit Hilfe der Tab. A2.2 aufgestellt werden.

Als erstes Beispiel sei eine heiße Redundanz 1 aus 2 mit den Elementen $E_1 = E_2 = E$ in Serie mit einem Vergleichselement E_V betrachtet. Die Ausfallrate aller Elemente sei konstant, ebenfalls die Reparaturrate. Das System verfügt über eine einzige Reparaturmannschaft und man nimmt an, dass während einer Reparatur auf Ebene System kein weiterer Ausfall auftreten kann.

Bild 3.11 gibt das Zuverlässigkeitsblockdiagramm und rechts das Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ an.

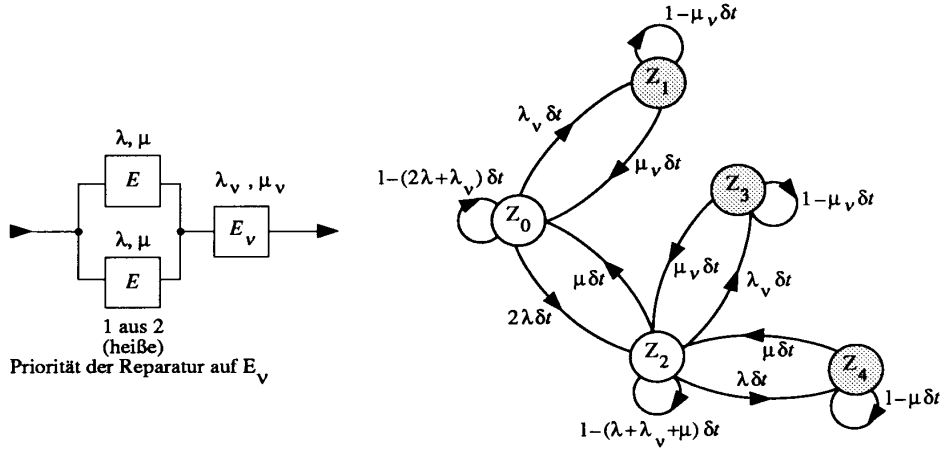


Bild 3.11 Zuverlässigkeitsblockdiagramm und Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ einer heißen Redundanz 1 aus 2 mit Vergleichselement bzw. Umschalteneinrichtung (konstanter Ausfall- und Reparaturrate, eine Reparaturmannschaft, kein weiterer Ausfall während einer Reparatur auf Systemniveau, Priorität der Reparatur auf E_v , t beliebig, $\delta t \rightarrow 0$)

Die Zuverlässigkeit des Systems gemäss Bild 3.11 kann mit Hilfe der Gleichungen aus Tab. A2.2 oder direkt gemäss folgender Überlegung gewonnen werden: Da das Serienelement E_v eine konstante Ausfallrate λ_v aufweist, gilt gemäss Gl. (3.17)

$$R_{S0}(t) = R_{1 \text{ aus } 2}(t) e^{-\lambda_v t}. \quad (3.36)$$

Die Laplace-Transformierte von $R_{S0}(t)$ folgt damit direkt aus der Laplace-Transformierten der Zuverlässigkeitsfunktion der Redundanz 1 aus 2, indem s durch $s + \lambda_v$ ersetzt wird. Für den Mittelwert der ausfallfreien Arbeitszeit $MTTF_{S0} = \tilde{R}_{S0}(0)$ folgt dann

$$MTTF_{S0} = \frac{3\lambda + \lambda_v + \mu}{(2\lambda + \lambda_v)(\lambda + \lambda_v) + \mu\lambda_v}, \quad (3.37)$$

und, insbesondere für $\lambda_v \ll \lambda \ll \mu$,

$$MTTF_{S0} \approx \frac{\mu}{2\lambda^2 + \mu\lambda_v} = \frac{1}{\lambda_v + 2\lambda(\lambda/\mu)}. \quad (3.38)$$

Unter Verwendung der Näherungsgleichung (3.30) folgt auch direkt

$$R_{S0}(t) = e^{-\lambda_{S0} t} \quad \text{mit} \quad \lambda_{S0} = \lambda_v + \frac{2\lambda^2}{3\lambda + \mu} \approx \frac{1}{MTTF_{S0}}. \quad (3.39)$$

Für den stationären Wert der Punkt-Verfügbarkeit PA_S und der durchschnittlichen Verfügbarkeit AA_S findet man [2,1991]

$$PA_S = AA_S = \frac{\mu^2 \mu_v + 2\lambda \mu \mu_v}{\mu^2 \mu_v + 2\lambda \mu \mu_v + 2\lambda(\lambda \mu_v + \lambda_v \mu) + \mu^2 \lambda_v} \approx 1 - \frac{\mu^2 \lambda_v + 2\lambda(\lambda \mu_v + \lambda_v \mu)}{\mu^2 \mu_v + 2\lambda \mu \mu_v} \quad (3.40)$$

und, insbesondere für $\lambda_v \ll \lambda \ll \mu \approx \mu_v$,

$$PA_S = AA_S \approx 1 - \frac{\lambda_v / \mu_v + 2(\lambda / \mu)^2}{1 + 2\lambda / \mu}. \quad (3.41)$$

Für den stationären Wert der Intervall-Zuverlässigkeit gilt

$$IR_S(\theta) = p_0 R_{S0}(\theta) + p_2 R_{S2}(\theta) \approx p_0 R_{S0}(\theta) \approx R_{S0}(\theta). \quad (3.42)$$

Als zweites Beispiel wird der Fall einer Majoritätsredundanz 2 aus 3 (heisse Redundanz 2 aus 3 in Serie mit einem Vergleichselement) untersucht. Die gleichen Überlegungen wie für den Fall von Bild 3.11 führen zu folgenden Resultaten für die Zuverlässigkeitsfunktion

$$R_{S0}(t) = R_{2 \text{ aus } 3}(t) e^{-\lambda_v t}, \quad (3.43)$$

für den Mittelwert der ausfallfreien Arbeitszeit

$$MTTF_{S0} = \frac{5\lambda + \lambda_v + \mu}{(3\lambda + \lambda_v)(2\lambda + \lambda_v) + \mu \lambda_v}, \quad (3.44)$$

für den stationären Wert der Punkt-Verfügbarkeit und der durchschnittlichen Verfügbarkeit

$$PA_S = AA_S = \frac{\mu(3\lambda + \lambda_v + \mu)}{(3\lambda + \lambda_v + \mu)(\lambda_v + \mu) + 3\lambda(2\lambda + \lambda_v)} \quad (3.45)$$

und für den stationären Wert der Intervall-Zuverlässigkeit

$$IR_S(\theta) = \frac{\mu(\lambda_v + \mu)R_{S0}(\theta) + 3\lambda \mu R_{S1}(\theta)}{(3\lambda + \lambda_v + \mu)(\lambda_v + \mu) + 3\lambda(2\lambda + \lambda_v)} \approx R_{S0}(\theta). \quad (3.46)$$

3.6 Näherungsformeln für grosse reparierbare Serie/Parallelstrukturen

Für grosse Serie-/Parallelstrukturen kann eine analytische Lösung aufwendig werden, auch wenn die Ausfall- und die Reparaturrate jedes Elements im Zuverlässigkeitsblockdiagramm konstant angenommen werden (hohe Anzahl von Zuständen). In solchen Fällen empfiehlt sich die Verwendung von Näherungsformeln. Prinzipiell können Näherungsformeln aus folgenden Überlegungen gewonnen werden:

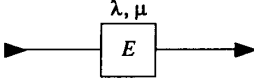
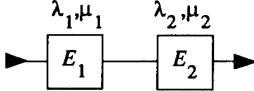
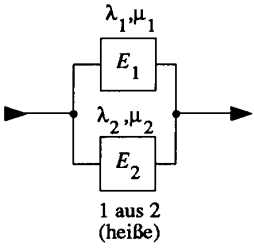
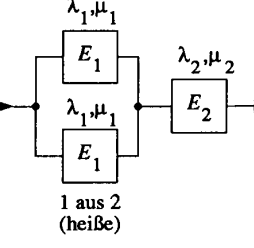
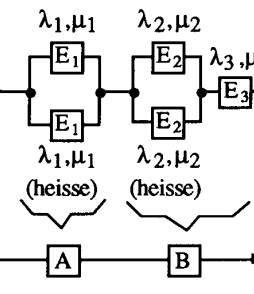
- 1. Unabhängige Arbeitsweise aller Elemente:** es wird angenommen, dass jedes Element im Zuverlässigkeitsblockdiagramm *unabhängig* von den anderen Elementen arbeitet und repariert wird (heisse Redundanz, unabhängige Elemente, eine Reparaturmannschaft pro Ele-

ment). Die Serieelemente sowie die Parallelstrukturen werden sukzessiv zusammengefasst; für jedes der so gebildeten Einzelemente wird die Ausfallrate aus dem Seriennmodell oder aus der Redundanz 1 aus 2 und die (äquivalente) Reparaturrate ($\mu = 1 / MTTR$) aus der Gleichung für die Punkt-Verfügbarkeit (PA_S) gewonnen.

2. **Bildung von Makrostrukturen:** Eine Makrostruktur ist eine Serie-, Parallel- oder einfache Serie-/Parallelstruktur, welche als eine Einheit (System) interpretiert wird und die Bedingungen der Modelle in den Abschnitten 3.1 bis 3.5 erfüllt, insbesondere eine Reparaturmannschaft und kein weiterer Ausfall nach Systemausfall. Die Prozedur ist ähnlich jenen vom Punkt 1. Tabelle 3.5 fasst die Grundmodelle zur Untersuchung von Serie-/Parallelstrukturen mit Hilfe von Makrostrukturen zusammen [2,1991].
3. **Zusammenfassung von Elementen oder Zuständen:** Elemente im Zuverlässigkeitsblockdiagramm (insbesondere Serienelemente) oder Zustände im Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ können oft zusammengefasst werden und dadurch die Anzahl Gleichungen stark reduzieren.
4. **Vernachlässigung von Zuständen:** Zustände mit mehr als k Ausfällen ($k \geq 2$) können oft auf dem Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ vernachlässigt werden, der resultierende Fehler kann unter allgemein gültigen Annahmen eingeschränkt werden [1].

Die Bilder 3.12 bis 3.14 geben den Vergleich bezüglich ausfallfreier Arbeitszeit $MTTF_{S0}$ und Unverfügbarkeit $1 - PA_S$ zwischen Grundstrukturen gemäss Tab. 3.5 ($\mu_1 = \mu_2 = \mu_3 = \mu$); man erkennt, dass (insbesondere für $\mu/\lambda_1 = MTBF_1/MTTR > 100$) das Verhältnis $\lambda_2/\lambda_1 = MTBF_1/MTBF_2$ bzw. $\lambda_3/\lambda_2 = MTBF_2/MTBF_3$ kleiner als 0.05 bleiben sollte (vgl. die Bemerkungen zu den Bildern 2.5 und 2.6 für nichtreparierbare Systeme ($\mu = 0$)).

Tabelle 3.5 Grundmodelle zur Vereinfachung der Untersuchung grosser Serie-Parallelstrukturen mit Hilfe von Makrostrukturen (heisse Redundanz, eine Reparaturmannschaft für die jeweilige Makrostruktur)

	$\lambda_S = \lambda, \quad \mu_S = \mu, \quad PA_S = \frac{1}{1 + \lambda_S / \mu_S} \approx 1 - \frac{\lambda_S}{\mu_S}$ $\Rightarrow \mu_S = \frac{\lambda_S PA_S}{1 - PA_S} \approx \frac{\lambda_S}{1 - PA_S}$
	$PA_S = \frac{1}{1 + \lambda_1 / \mu_1 + \lambda_2 / \mu_2} \approx 1 - \left(\frac{\lambda_1}{\mu_1} + \frac{\lambda_2}{\mu_2} \right)$ $\lambda_S = \lambda_1 + \lambda_2 \Rightarrow \mu_S \approx \frac{\lambda_S}{1 - PA_S} \approx \frac{\lambda_1 + \lambda_2}{\lambda_1 / \mu_1 + \lambda_2 / \mu_2}$
	$PA_S = \frac{1}{1 + \frac{\lambda_1 \lambda_2 (\mu_1^2 + \mu_2^2 + (\lambda_1 + \lambda_2)(\mu_1 + \mu_2))}{\mu_1 \mu_2 (\mu_1 \mu_2 + (\lambda_1 + \lambda_2)(\lambda_1 + \lambda_2 + \mu_1 + \mu_2))}} \approx 1 - \frac{\lambda_1 \lambda_2}{\mu_1^2 \mu_2^2} (\mu_1^2 + \mu_2^2)$ $MTTF_{S0} = \frac{(\lambda_1 + \mu_2)(\lambda_2 + \mu_1) + \lambda_1(\lambda_1 + \mu_2) + \lambda_2(\lambda_2 + \mu_1)}{\lambda_1 \lambda_2 (\lambda_1 + \lambda_2 + \mu_1 + \mu_2)} = \frac{1}{\lambda_S}$ $\rightarrow \lambda_S \approx \frac{\lambda_1 \lambda_2 (\mu_1 + \mu_2)}{\mu_1 \mu_2} \Rightarrow \mu_S \approx \frac{\lambda_S}{1 - PA_S} \approx \mu_1 \mu_2 \frac{\mu_1 + \mu_2}{\mu_1^2 + \mu_2^2}$
	$PA_S = \frac{\mu_1^2 \mu_2 + 2 \lambda_1 \mu_1 \mu_2}{\mu_1^2 \mu_2 + 2 \lambda_1 (\mu_1 \mu_2 + \lambda_1 \mu_2 + \lambda_2 \mu_1) + \mu_1^2 \lambda_2} \approx 1 - \frac{\lambda_2 + 2 \mu_2 (\lambda_1 / \mu_1)^2}{\mu_2 (1 + 2 \lambda_1 / \mu_1)}$ $MTTF_{S0} = \frac{3 \lambda_1 + \lambda_2 + \mu_1}{(2 \lambda_1 + \lambda_2)(\lambda_1 + \lambda_2) + \mu_1 \lambda_2} \approx \frac{\mu_1}{2 \lambda_1^2 + \mu_1 \lambda_2} \approx \frac{1}{\lambda_S}$ $\Rightarrow \mu_S \approx \frac{\lambda_S}{1 - PA_S} \approx \frac{\mu_2 (2 \lambda_1^2 + \mu_1 \lambda_2) (1 + 2 \lambda_1 / \mu_1)}{\mu_1 \lambda_2 + 2 \mu_2 \lambda_1^2 / \mu_1}$
	mit $\lambda_A, \lambda_B, \mu_A, \mu_B$ aus den oberen Teil dieser Tabelle folgt $MTTF_{S0} = 1/\lambda_S \quad \text{mit} \quad \lambda_S = \lambda_A + \lambda_B \approx \frac{2\lambda_1^2}{\mu_1} + \frac{2\lambda_2^2 + \mu_2 \lambda_3}{\mu_2}$ $PA_S = \frac{1}{1 + \lambda_A / \mu_A + \lambda_B / \mu_B} \approx 1 - \left(\frac{\lambda_A}{\mu_A} + \frac{\lambda_B}{\mu_B} \right)$ mit $\mu_A \approx \mu_1$ und $\mu_B \approx \frac{\mu_3 (2\lambda_2^2 + \mu_2 \lambda_3) (1 + 2\lambda_2 / \mu_2)}{\mu_2 \lambda_3 + 2\mu_3 \lambda_2^2 / \mu_2}$ folgt $PA_S \approx 1 - \left(\frac{2\lambda_1^2}{\mu_1^2} + \frac{\mu_2 \lambda_3 + 2\mu_3 \lambda_2^2 / \mu_2}{\mu_2 \mu_3 (1 + 2\lambda_2 / \mu_2)} \right)$

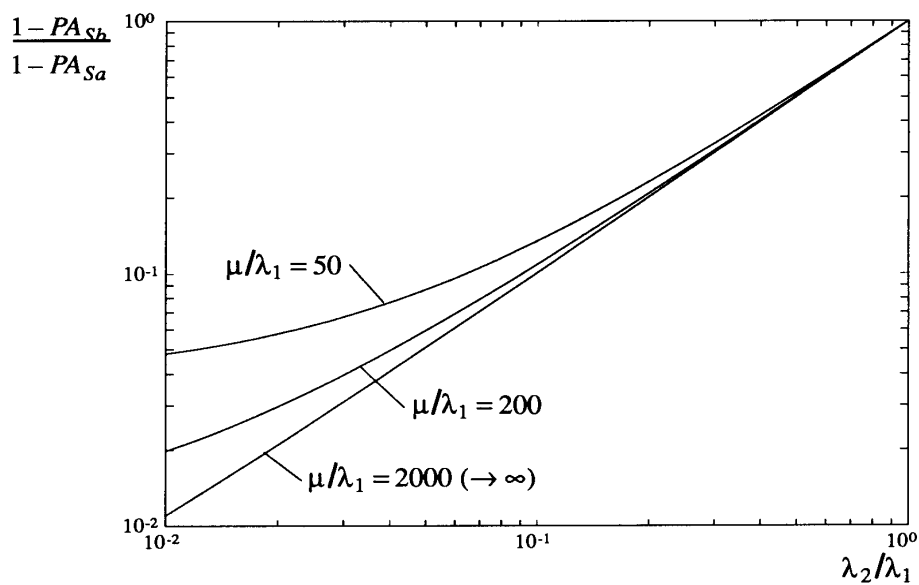
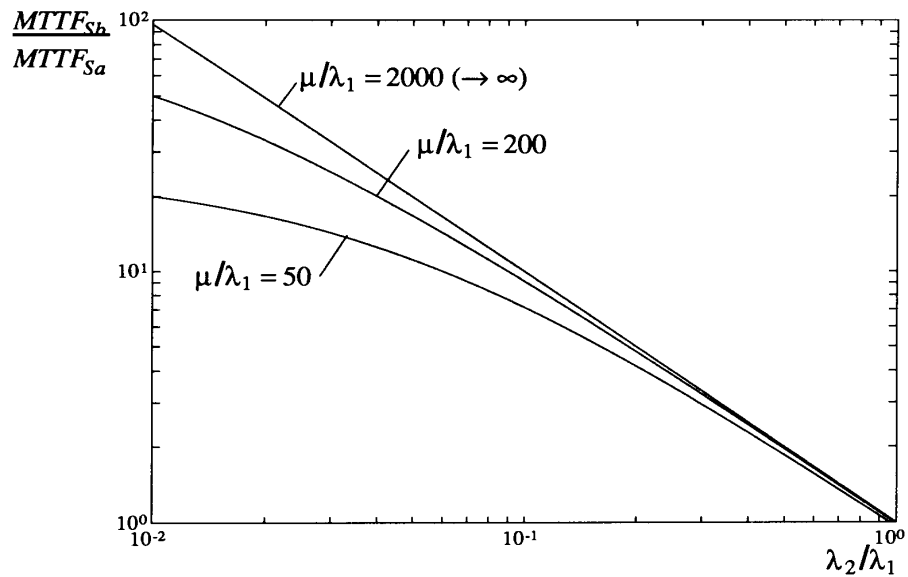
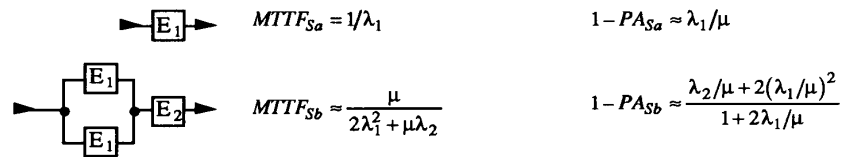


Bild 3.12 Vergleich bezüglich ausfallfreier Arbeitszeit $MTTF_{S0}$ und Unverfügbarkeit $1 - PA_S$ zwischen Grundstrukturen gemäss Tab. 3.5 ($\mu_1 = \mu_2 = \mu$), vgl. Bild 2.5 für nichtreparierbare Systeme ($\mu = 0$)

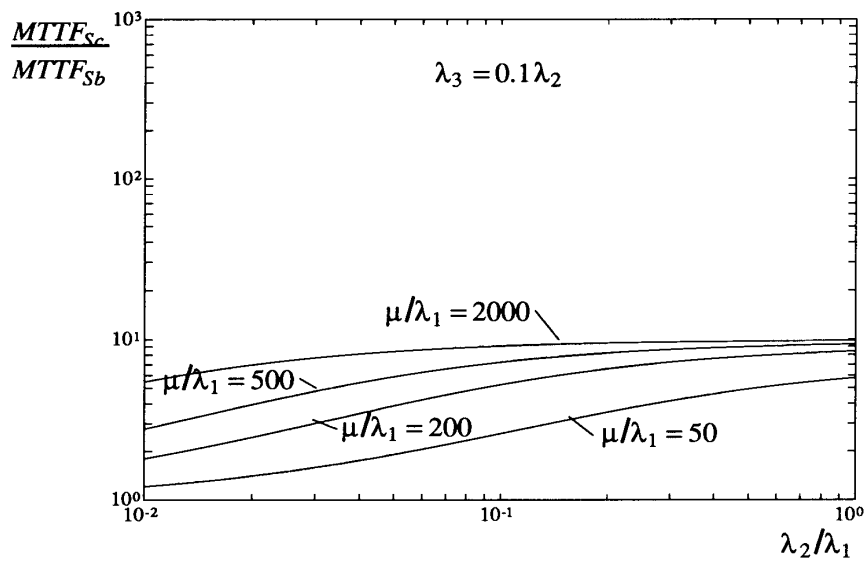
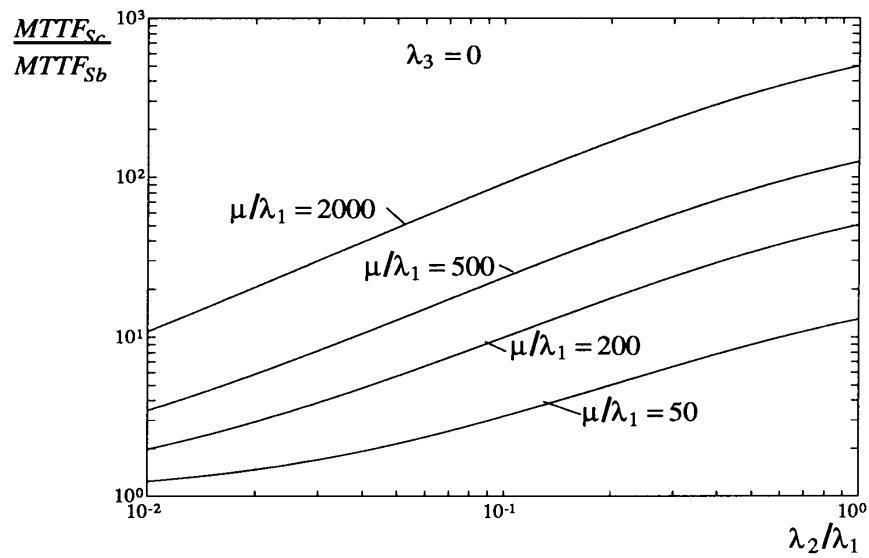
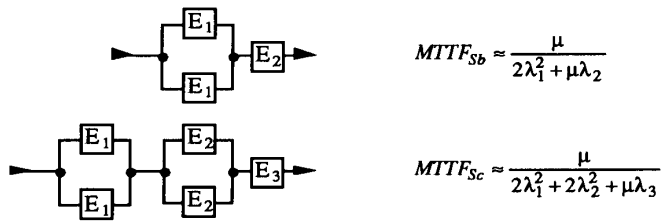


Bild 3.13 Vergleich bezüglich ausfallfreier Arbeitszeit $MTTF_{S0}$ zwischen Grundstrukturen gemäss Tab. 3.5 ($\mu_1 = \mu_2 = \mu_3 = \mu$), oben $\lambda_3 = 0$, unter $\lambda_3 = 0.1\lambda_2$, vgl. Bild 2.6 für nichtreparierbare Systeme ($\mu = 0$)

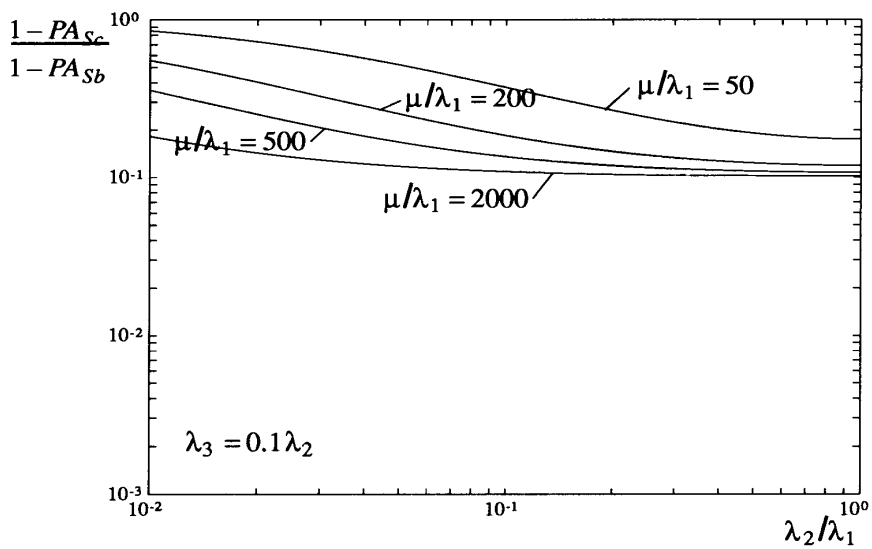
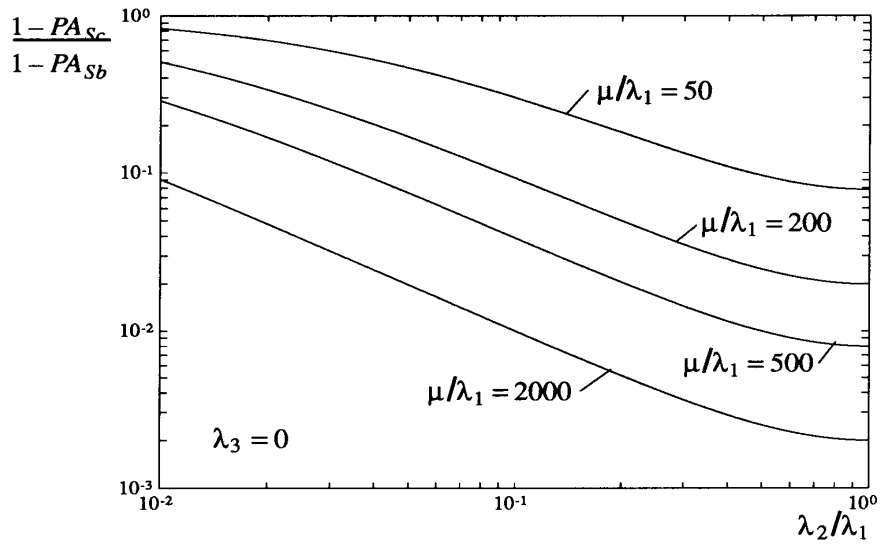
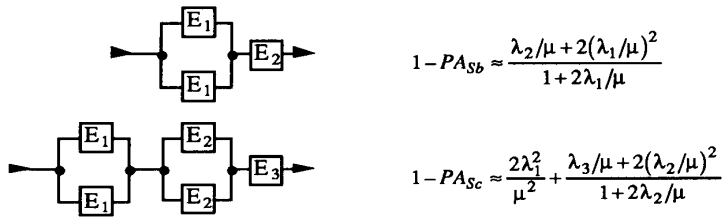


Bild 3.14 Vergleich bezüglich Unverfügbarkeit $1 - PA_S$ zwischen Grundstrukturen gemäss Tab. 3.5 ($\mu_1 = \mu_2 = \mu_3 = \mu$), oben $\lambda_3 = 0$, unter $\lambda_3 = 0.1\lambda_2$, vgl. Bild 2.6 für nichtreparierbare Systeme ($\mu = 0$)

4 Einfluss von Redundanz und Ein-Ausschalten auf Zuverlässigkeit und Energieverbrauch von Geräten und Systemen

Hohe Zuverlässigkeits- oder Verfügbarkeitsforderungen erfüllt man oft nur mit Hilfe von Redundanz. Diese kann prinzipiell als heisse ($\lambda_r = \lambda$), warme ($\lambda_r < \lambda$) oder kalte ($\lambda_r \equiv 0$) Redundanz realisiert werden ($\lambda = 1/\text{MTBF}$ ist die Ausfallrate im Arbeitszustand und λ_r stellt die Ausfallrate im Reservezustand dar). Wie die Resultate in Tabelle 3.3 zeigen, gilt für eine Redundanz 1 aus 2:

a) Mittelwert der ausfallfreien Arbeitszeit im reparierbaren Fall ($\mu = 1/\text{MTTR} \gg \lambda$)

$$MTTF_{S0} = \frac{2\lambda + \lambda_r + \mu}{\lambda(\lambda + \lambda_r)} \begin{cases} \approx \mu/2\lambda^2 & \text{heisse Redundanz} \\ \approx \mu/\lambda^2 & \text{kalte Redundanz} \end{cases} \quad (4.1)$$

b) $MTTF_{S0}$ im nicht reparierbaren Fall ($\mu = 0$)

$$MTTF_{S0} = \frac{2\lambda + \lambda_r}{\lambda(\lambda + \lambda_r)} \begin{cases} \approx 3/2\lambda^2 & \text{heisse Redundanz} \\ \approx 2/\lambda^2 & \text{kalte Redundanz} \end{cases} \quad (4.2)$$

c) stationärer Wert der Punktverfügbarkeit

$$PA_S = AA_S \approx 1 - \frac{\lambda(\lambda + \lambda_r)}{\mu^2} \begin{cases} \approx 1 - 2(\lambda/\mu)^2 & \text{heisse Redundanz} \\ \approx 1 - (\lambda/\mu)^2 & \text{kalte Redundanz} \end{cases} \quad (4.3)$$

Mit einer kalten Redundanz ist der Gewinn an Zuverlässigkeit (gemessen an $MTTF_{S0}$) bzw. die Reduktion an Unverfügbarkeit ($1 - PA_S$) rund um etwa ein Faktor 2 grösser als im Falle einer kalten Redundanz (1.5 im nicht-reparierbaren Fall). Wenn man ferner in Betracht zieht, dass bei einer heissen Redundanz 1 aus 2 der Leistungsverbrauch doppelt so gross ist wie bei einer kalten, stellt die kalte Redundanz (wenn sie realisiert werden kann) die beste Lösung dar. Ähnliche Überlegungen können für Geräte oder Systeme ohne Redundanz betreffend dem Ein- und Ausschalten gemacht werden.

Im Hinblick auf eine Minimierung des Leistungsverbrauchs stellt sich damit die Frage nach der Möglichkeit, Geräte oder Systeme zu realisieren, die ohne Einbusse auf die Zuverlässigkeit beliebig oft ein- und ausgeschaltet werden können, oder für welche kalte Redundanz verwendet werden kann. Es gibt Geräte, die solche Forderungen erfüllen müssen und zu diesem Zweck speziell konzipiert worden sind, zu diesen gehören z.B. die meisten Telefonendgeräte (andere Geräte sind auch unproblematisch. Wie eine Studie an der Professur für Zuverlässigkeitstechnik der ETH Zürich gezeigt hat [7], können für verschiedene Geräte Modifikationen in der Schaltungsauslegung vorgeschlagen werden, die es erlauben (eventuell mit geringen

4.1 Allgemeine Betrachtungen

Einbussen an Komfort) leicht belastete Zustände (tiefer Standby) mit kleinem Leistungsverbrauch zu realisieren. Bei anderen Geräten, insbesondere bei komplexen digitalen Anlagen der Computertechnik, ist sowohl das Ein- und Ausschalten wie auch die Verwendung einer kalten Redundanz oft problematisch (Latch-up Gefahr, transiente Vorgänge, Datensicherung).

Zusammenfassend kann gesagt werden, dass es durch geeignete Massnahmen in der Entwicklungsphase (Wahl der Technologie sowie die Auslegung der Schaltung, vgl. z.B. [2, 1991] für konkrete Entwicklungsrichtlinien bezüglich Zuverlässigkeit und Instandhaltbarkeit) sowie eventuell durch eine Redimensionierung des Komfortes, für viele elektronische Geräte und Systeme prinzipiell möglich ist, den Energieverbrauch zu minimieren. Wichtig dabei ist, dass solche Betrachtungen nicht erst bei der Installation einer Anlage gemacht werden, sondern eben in der Entwicklungsphase, wo Redundanz oft auf Modul- oder Geräteebene eingesetzt werden kann. In gewissen Fällen kann die Reduktion des Energieverbrauchs durch die Verwendung von kalter Redundanz, das Einführen eines tiefen Standby-Zustands oder einfach durch Ein- und Ausschalten noch verstärkt werden (die letzten Massnahmen können aber nicht bei allen Geräten oder Systemen realisiert werden und sind in der Regel mit Zuverlässigkeitsspezialisten zu besprechen).

Um einen besseren Einblick in die Wechselwirkung zwischen Zuverlässigkeit und Energieverbrauch zu geben, werden in den folgenden Abschnitten zwei konkrete Beispiele untersucht. Das erste Beispiel betrifft ein Multicomputer-System (zur Vereinfachung nicht reparierbar bis zum Systemausfall) und stützt sich auf die Resultate von Abschnitt 2.1.7. Das zweite Beispiel behandelt eine unterbrechungsfreie Stromversorgung (reparierbar) und stützt sich auf die Grundlagen vom Kapitel 3. Das Rohmaterial für diese Abschnitte wurde von den Herren G.A. Zardini und L. Strozzi geliefert.

4.2 Multicomputer-System

In diesem Beispiel sollen verschiedene System-Varianten für ein Rechenzentrum (als Beispiel für ein Multicomputer-System) unter Berücksichtigung von drei Hauptkomponenten

- Netz (allein oder mit unterbrechungsfreier Stromversorgung (USV))
- Rechenanlage (Computer oder Workstations)
- Klimaanlage

bezüglich Mittelwert der ausfallfreien Arbeitszeit $MTTF_s$ und Leistungsverbrauch verglichen werden. Für die einzelnen Teile werden folgende Annahmen getroffen (vgl. Tab. 5.2):

	$MTBF = 1/\lambda$	Leistungsverbrauch
Netz (N)	100 h	—
USV (U)	100'000 h	5 kW
Klimaanlage (K)	5'000 h	40 kW
Computer (C)	4'000 h	20 kW
Workstation (W)	40'000 h	0.5 kW

Alle Teile haben eine konstante Ausfallrate $\lambda = 1/MTBF$, redundante Teile sind in heisser Redundanz. Untersucht werden folgende Varianten:

1. Netz, Computer und Klimaanlage, ohne Redundanz.
2. Wie Fall 1 aber mit Redundanz auf Netz (USV).

3. Netz mit Redundanz (Annahmen für die USV: doppelte Leistung, gleiche Ausfallrate), 2 Computer in Redundanz (1 aus 2) und eine Klimaanlage (Annahmen für die Klimaanlage: doppelte Leistung, gleiche Ausfallrate).
4. Zwei (geographisch) getrennte Anlagen, jede gemäss Fall 1, auf 2 (getrennte) Netze.
5. Netz mit Redundanz (USV), Workstations (Annahmen: gleicher Energieverbrauch wie zwei Computer in Redundanz 1 aus 2, 50% der Workstations reichen für die Erfüllung der geforderten Funktion, das entspricht einer Redundanz 40 aus 80, keine Klimaanlage (die Workstations sind räumlich verteilt).
6. Wie Fall 5 aber die Workstations sind vernetzt mit Fileserver in Redundanz 1 aus 2.

Tabelle 4.1 zeigt die entsprechenden Zuverlässigkeitsblockdiagramme und fasst die Resultate zusammen.

Tab. 4.1 System-Varianten für ein Rechenzentrum

Variante	Zuverlässigkeits-Blockdiagramm	Zuverlässigkeits-Funktion $R_S = R_S(t)$	MTTF _S (Gl. 2.9) [h]	Verbrauchte Leistung [kW]
1		$R_N R_C R_K$	100	60
2		$(R_N + R_U - R_N R_U) \cdot R_C R_K$	2'200	65
3		$(R_N + R_U - R_N R_U) \cdot (2R_C - R_C^2) R_K$	3'000	130
4		$2R_N R_C R_K - (R_N R_C R_K)^2$	140	120
5		$(R_N + R_U - R_N R_U) \cdot \sum_{i=40}^{80} \binom{80}{i} R_W^i (1 - R_W)^{80-i}$	24'700	50
6		$(R_N + R_U - R_N R_U) \cdot \sum_{i=39}^{78} \binom{78}{i} R_W^i (1 - R_W)^{78-i} \cdot (2R_W - R_W^2)$	22'300	50

Bild 4.1 gibt den Verlauf der Zuverlässigkeitsfunktion $R_S(t)$ für die 6 Varianten wieder. Aus Tab. 4.1 und Bild 4.1 geht hervor, dass die Varianten 5 und 6 mit verteilter Intelligenz eindeu-

tig besser als die anderen Varianten abschneiden. Dies ist um so mehr zu betonen, weil heute die Leistungsfähigkeit einer Workstation heute schon jene eines mittelgrossen Computer erreicht [3,9].

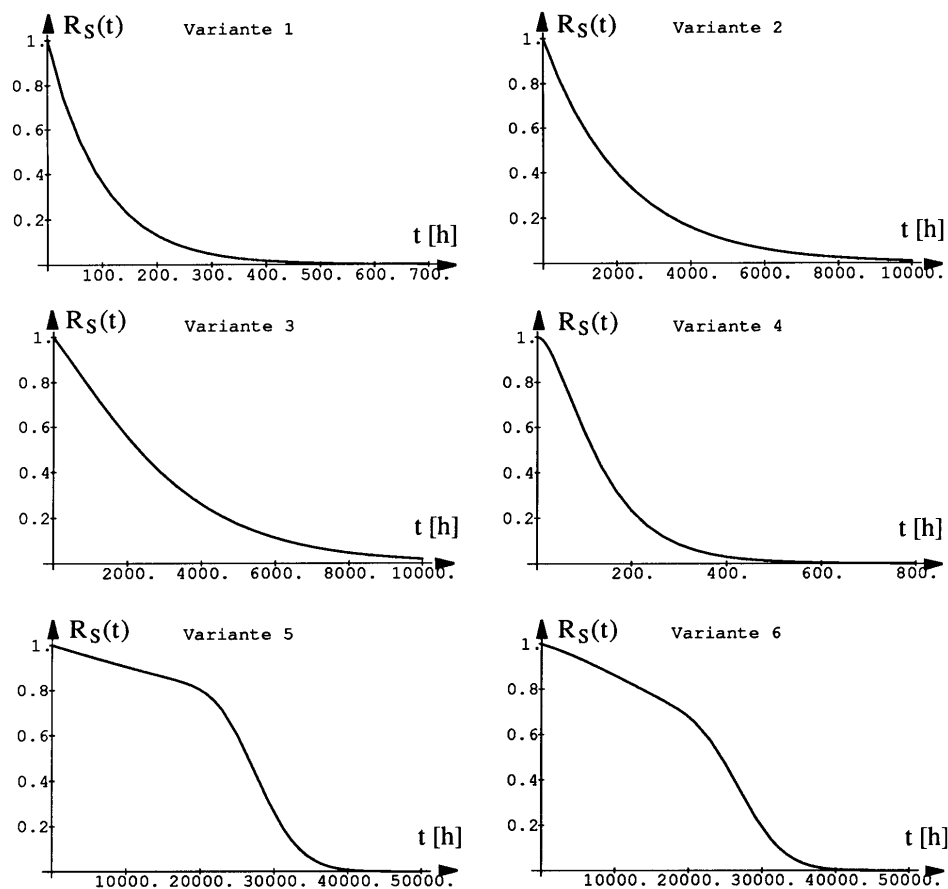


Bild 4.1 Verlauf der Zuverlässigkeitsfunktion $R_S(t)$ für die Varianten gemäss Tab. 4.1

4.3 Unterbrechungsfreie Stromversorgung

Eine unterbrechungsfreie Stromversorgung (USV) hat die Sicherstellung der Stromversorgung bei Netzstörungen und kurzen Netzunterbrüchen (bis etwa 10 min) sowie bei der Durchführung der Abschaltprozedur im Fall längerer Netzunterbrüche zur Aufgabe. Die meisten Computeranlagen können Netzunterbrüche bis zu etwa 100ms selbst überbrücken. Eingehende Untersuchungen haben aber gezeigt, dass selbst in "stabilen" Netzen im Mittel etwa alle 100 h Netzunterbrüche zwischen 130 bis 700 ms auftreten. Prinzipiell wird zwischen on-line USV-

Anlagen (die Last wird ständig durch einen Wechselrichter gespeisen) und off-line USV-Anlagen (die USV-Anlage wird nur beim Netzausfall eingeschaltet) unterschieden. Off-line Anlagen benötigen eine Anlaufzeit und können deshalb nicht immer eingesetzt werden. Bei den on-line USV-Anlagen wird zwischen Geradeaus-Anlagen (bei welchen der Wechselrichter die volle Last speisen kann) und reversiblen Anlagen (bei welchen der Wechselrichter als Netzstabilisator und Ladegerät wirkt) nicht unterschieden. Bild 4.2 zeigt das Blockschaltbild einer Geradeaus-Anlage.

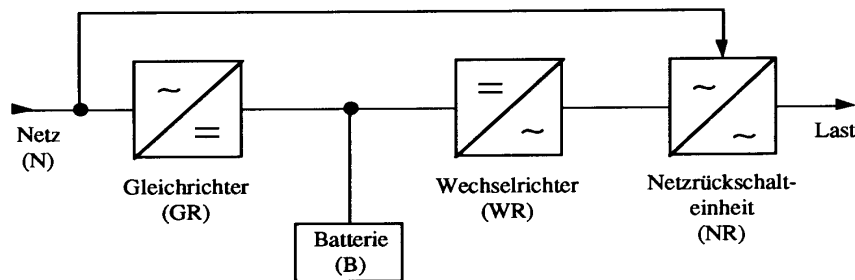


Bild 4.2 Blockschaltbild einer Geradeaus-USV-Anlage

Bild 4.3 zeigt das Zuverlässigkeitsblockdiagramm einer Geradeaus-USV-Anlage, a) ohne Redundanz und b) mit heisser Redundanz 1 aus 2. Das Zuverlässigkeitsblockdiagramm wurde mit Hilfe einer FMEA/FMECA (Abschnitt 2.2) zur korrekten Identifizierung der Schnittstellen-Elemente S und B in Bild 4.3b, aufgestellt.

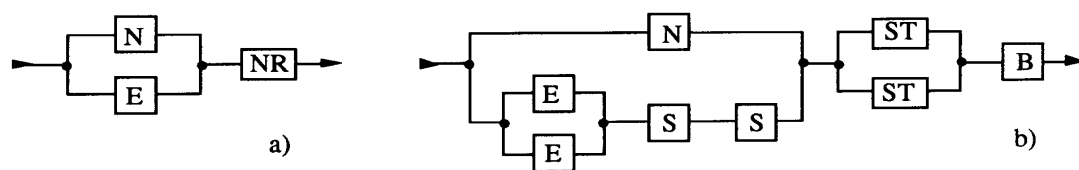


Bild 4.3 Zuverlässigkeitsblockdiagramme einer Geradeaus-USV-Anlage, a) ohne Redundanz und b) mit 2 USV-Anlagen in heisser Redundanz (N = Netz, E = Gleichrichter + Batterie + Wechselrichter), NR = Netzurückschalt-einheit, S = Synchronisation + Lastaufteilung, ST = Steuerung, B = Steuersignale

Tabelle 4.2 gibt die Ausfall- und die Reparaturraten der einzelnen Elemente sowie den Mittelwert der ausfallfreien Arbeitszeit auf Systemebene $MTTF_S = MTTF_{S0}$ für die beiden Varianten in Bild 4.3 an (die Berechnung erfolgt nach der Methode der Makrostrukturen, Tab. 3.5).

Tab. 4.2 Annahmen und Resultate für die Geradeaus-USV-Anlage gemäss Bild 4.3

	Bild 4.3a			Bild 4.3b				
	N	E	NR	N	E	S	ST	B
Ausfallrate (λ in h^{-1})	0.01	$40 \cdot 10^{-6}$	$6 \cdot 10^{-6}$	0.01	$40 \cdot 10^{-6}$	$4 \cdot 10^{-6}$	$6 \cdot 10^{-6}$	$0.5 \cdot 10^{-6}$
Reparaturrate (μ in h^{-1})	10	0.16	0.16	10	0.16	0.16	0.16	0.16
$MTTF_S = MTTF_{S0}$ in h	120'000			1'000'000				

Für die Stromversorgung grosser Rechenanlagen wird aus betrieblichen Gründen oft eine mittlere ausfallfreie Arbeitszeit grösser als 300'000h gefordert. Der relative Energieverlust einer USV-Anlage wird durch ihren Wirkungsgrad η bestimmt. Der Verlauf von η als Funktion der Last ist (zusammen mit der Identifizierung der Schnittstellen-Elemente im Zuverlässigkeitsblockdiagramm (Bild 4.3) und der Verifizierung der eingebauten Prüfmöglichkeiten der Anlage) ein weiterer wichtiger Aspekt bei der Beurteilung einer USV-Anlage. Der typische Verlauf von η für grosse USV-Anlagen ist etwa 85% bei 1/4 Last und zwischen 92% und 95% oberhalb der halben Last. Tabelle 4.3 zeigt den Energieverbrauch einer USV-Anlage gemäss Bild 4.3a unter Berücksichtigung einer typischen zeitlichen Verteilung der Last in einem Rechenzentrum.

Tab. 4.3 Energieverbrauch (Verlust mal Zeit) einer Geradeaus-USV-Anlage gemäss Bild 4.3a, bezogen auf dem Energieverbrauch der Computeranlage.

Zeit	Last	η	Verlust (1- η)/ η	Verlust × Zeit
60%	100%	94%	6.4%	3.8%
40%	50%	92%	8.7%	3.5%
				7.3%

Der Preis für eine sichere Stromversorgung mit einer Geradeaus-USV-Anlage liegt gemäss Tab. 4.3 bei etwa 7.3% Mehrenergieverbrauch (6.4% wenn die Belastung ständig 100% wäre). Im Falle einer USV-Anlage mit Redundanz 1 aus 2 wird in der Regel mit zwei kleineren Anlagen (etwa 80%) operiert, so dass der Mehrenergieverbrauch bei etwa 12% liegt; auch ist es möglich, bei tieferen Belastungen (z.B. unterhalb 50%) eine der redundanten Anlagen abzuschalten.

5 Festlegung und Durchsetzung von Zuverlässigkeitsforderungen

Zuverlässigkeitsforderungen sind wichtig, um spezifizierte operationelle Bedingungen zu erfüllen, sowie um Betriebs- und Instandhaltungskosten unter Kontrolle zu halten. Sie schaffen klare Voraussetzungen für den Entwicklungsingenieur, müssen aber nach dem Grundsatz *so gut wie nötig* formuliert werden, denn übertriebene Forderungen können die Marktchancen eines Geräts oder Systems negativ beeinflussen. Nach einer kurzen Darlegung der allgemeinen Kundenforderung bezüglich Qualitäts- und Zuverlässigkeitssicherungen von Geräten und Systemen sowie des Zusammenhangs zwischen Lebenslaufkosten und Qualitäts- bzw. Zuverlässigkeitsforderungen wird in diesem Kapitel ausführlich auf die Festlegung und Durchsetzung von Zuverlässigkeitsforderungen eingegangen.

5.1 Kundenforderungen

Kundenforderungen auf dem Gebiet der Qualitäts- und Zuverlässigkeitssicherung können qualitativen oder quantitativen Charakter haben. *Quantitative* Forderungen (z.B. MTBF, MTTR, Verfügbarkeit) werden wie technische Eigenschaften in Pflichtenheften und Verträgen festgelegt. Sie spezifizieren die Werte für die Zuverlässigkeit, Instandhaltbarkeit oder Verfügbarkeit mit den entsprechenden Angaben bezüglich geforderter Funktion, Umwelt- bzw. Arbeitsbedingungen, personeller und materieller Bedingungen, logistischer Unterstützung und Nachweiskriterien. *Qualitative* Forderungen sind in Normen, Standards oder Empfehlungen enthalten. Diese Normen fordern grundsätzlich ein Qualitätssicherungssystem. Ihre Hauptziele sind:

- *Standardisierung* bezüglich Konfiguration, Umweltbedingungen, Prüfbedingungen, Wahl von Bauteilen und Stoffen, logistischer Unterstützung usw.
- Vereinheitlichung der Qualitäts- und Zuverlässigkeitssicherungssysteme
- Festlegung einer gemeinsamen Sprache.

Die wichtigsten Normen und Standards auf dem Gebiet der Qualitäts- und Zuverlässigkeitssicherung sind in Tabelle 5.1 zusammengefasst.

Die ISO 9001 bzw. die äquivalente EN90001 sollen in den neunziger Jahren die nationalen Qualitätsnormen für Gebrauchsgüter ablösen. Sie sind für alle Arten von Produkten gedacht und damit eher allgemein gehalten und in der Struktur noch nicht ganz ausgereift. Die ISO 9001 fordert ein System zur Sicherstellung der *Qualität* in der Entwicklungs- und Fertigungsphase. Das System soll wirksam und wirtschaftlich sein und die Gewähr bieten, dass zu einer Abnahmeprüfung nur abnahmetaugliche Geräte vorgelegt werden. Die Forderungen erstrecken sich auf Organisation, Planung, Konfigurationsmanagement, Beschaffung, Steuerung der Fertigung, Prüf- und Messeinrichtungen, Qualitätsprüfung und Korrekturmaßnahmen.

Tabelle 5.1 Wichtigste Normen und Standards auf dem Gebiet der Qualitäts- und Zuverlässigkeitssicherung von Geräten und Anlagen

<i>Allgemeine Gebrauchsgüter</i>			
1987	Int.	ISO 9000	: Quality management and quality assurance standards - Guidelines for selection and use
		9001	: Quality systems - Model for quality assurance in design/development, production, installation and servicing see also 9000-3 (1991): Guidelines for the application of ISO 9001 to the development, supply and maintenance of software
		9002	: Quality systems - Model for quality assurance in production and installation
		9003	: Quality systems - Model for quality assurance in final inspection and test
		9004	: Quality management and quality system elements - Guidelines see also 9004-2 (1991): Guidelines for services
1989	CH	SN EN 45011	: Allgemeine Kriterien für Stellen, die Produkte zertifizieren
		EN 45012	: Allgemeine Kriterien für Stellen, die Qualitätssicherungs-systeme zertifizieren
<i>Militär</i>			
1959	USA	MIL-Q-9858	: Quality Program Requirements (Ed.A, 1963)
1965	USA	MIL-STD-785	: Reliability Program for Systems and Equipment Development and Production (Ed.B, 1980)
1965	USA	MIL-STD-781	: Reliability Testing for Engineering Development, Qualification and Production (Ed.D, 1986)
1966	USA	MIL-STD-470	: Maintainability Program for Systems and Equipment (Ed. A, 1983)
1968	NATO	AQAP-1	: NATO Requirements for an Industrial Quality Control System (3th. Ed. 1984)

Die Erfüllung der ISO 9001 wird zunehmend als Voraussetzung für die Erteilung von Beschaffungsaufträgen betrachtet. Darüber besteht aber noch keine internationale juristisch gestützte Abmachung, so dass die damit verbundene *Zertifizierung* nur bedingt international anerkannt wird. Zur Zertifizierung gehört die Genehmigung des Qualitätssicherungssystems des Auftragnehmers und insbesondere des Qualitätssicherungshandbuchs. Für die Fertigungsphase allein ist die *ISO 9002* massgebend.

Für die Aspekte der Zuverlässigkeit und der Instandhaltbarkeit sind die internationalen Normen noch im Aufbau begriffen. Als Referenz können z.Z. die MIL-STD-785 für die *Zuverlässigkeit* und die MIL-STD-470 für die *Instandhaltbarkeit* abschnittsweise herangezogen werden. Die *MIL-STD-785* fordert die Sicherstellung der Zuverlässigkeit und insbesondere die Realisierung eines *Zuverlässigkeitssicherungsprogramms*. Die Forderungen sind der Komplexität des Geräts bzw. Systems anzupassen. Sie sind sehr detailliert angegeben und umfassen Organisation, Planung, Analysen, Wahl und Qualifikation von Bauteilen und Stoffen, Entwurfsüberprüfungen, FMEA/FMECA, Qualifikationsprüfungen, Erfassung und Analyse der Ausfälle, Korrekturmassnahmen, Hebung der Zuverlässigkeit in der Fertigungsphase und Nachweisprüfungen. Zum *Nachweis einer MTBF* sind die *MIL-STD-781* und das *MIL-HDBK-781* massgebend. Die Prüfung wird grundsätzlich durch das Prüfniveau (geforderte Funktion, Temperaturzyklen, Ein- und Ausschaltung, Vibrationen usw.) und durch den Prüfplan bestimmt. Als Prüfpläne werden einfache Prüfpläne und Sequentialtests angeführt. Diese Dokumente sind nur für den Fall einer *konstanten Ausfallrate* λ bzw. einer $MTBF=1/\lambda$ anwendbar. Die *MIL-STD-470* fordert die Realisierung eines *Instandhaltbarkeitssicherungsprogramms* und insbesondere die Durchführung von Analysen, Entwurfsüberprüfungen, FMEA/FMECA und Nachweisprüfungen. Gefordert wird auch die Erarbeitung eines *Instandhaltungskonzepts* (Ausfallerkennung/Ausfalllokalisierung, Strukturierung des Geräts bzw. Systems, Kundendokumentation, Logi-

stik) mit den nötigen Verbindungen zur logistischen Unterstützung und zu den konstruktiven Massnahmen. Der *Nachweis der Instandhaltbarkeit* wird in der MIL-STD-471 behandelt.

5.2 Wirtschaftlichkeitsbetrachtungen

Bei der Beschaffung einer grossen Anzahl gleicher Geräte oder von komplexen Anlagen oder Systemen mit langer Anwendungsdauer spielen neben den reinen Anschaffungskosten auch die Betriebs-, Instandhaltungs- und Ausscheidungskosten eine wichtige Rolle. Alle diese Kosten hängen in der Regel stark von den Qualitäts- und Zuverlässigkeitsforderungen ab. Ihre Summe stellt die *Lebenslaufkosten* dar. Auch wenn genaue Modellvorstellungen zur Berechnung und Optimierung der Lebenslaufkosten noch wenig bekannt sind, ist bei der Festlegung von Qualitäts- und Zuverlässigkeitsforderungen der Zusammenhang zwischen diesen Forderungen und den Lebenslaufkosten so weit wie möglich zu berücksichtigen. Bild 5.1 zeigt als Beispiel den Einfluss des Qualitätsniveaus g_Q oder des Zuverlässigkeitsniveaus g_R auf die Summe C der Kosten für die Sicherstellung der Qualität, Zuverlässigkeit, Instandhaltbarkeit und Logistik von zwei komplexen Systemen mit verschiedenen Einsatzprofilen [2,1986].

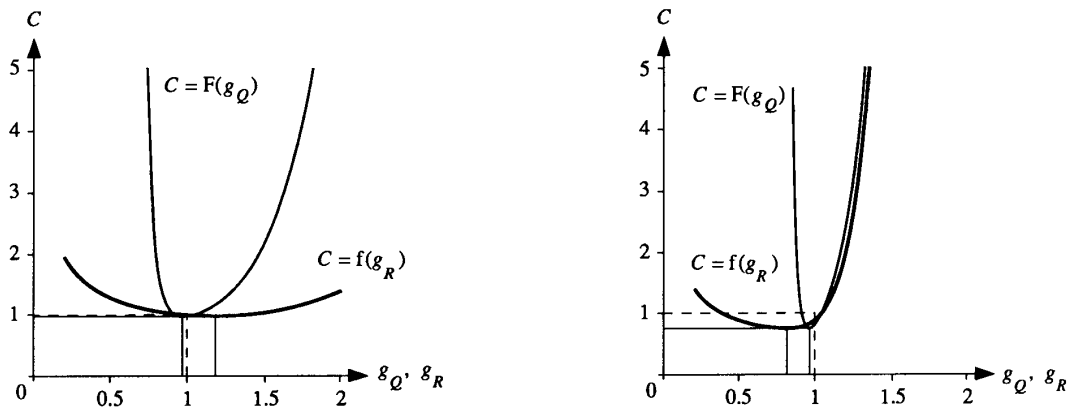


Bild 5.1 Summe der relativen Kosten C zur Qualitätssicherung und für die Sicherstellung der Zuverlässigkeit, Instandhaltbarkeit und Logistik von zwei komplexen Systemen als Funktion des Erfüllungsgrads des Qualitätsniveaus (g_Q = erreichtes Qualitätsniveau / vorgegebenes Ziel) oder der *MTBF* (g_R = erreichte *MTBF* / vorgegebenes Ziel), gestrichelt ist die Zielvorgabe (C ist als Erwartungswert zu betrachten)

Der Einstieg in solche Modellbetrachtungen soll anhand von Beispiel 5.1 geschaffen werden.

Beispiel 5.1

Gegeben sei eine Baugruppe mit n statistisch unabhängigen Bauteilen. Die Defektquote der Bauteile sei p , die Kosten für i Reparaturen (Lokalisierung und Ersatz von i defekten Bauteilen auf der Baugruppe) sei k_i . Gesucht wird: a) der Mittelwert C_a der gesamten Reparaturkosten (die Baugruppe darf kein defektes Bauteil enthalten) und b) der mögliche Gewinn, wenn die Bauteile einer Eingangsprüfung unterzogen werden, Prüfkosten C pro Bauteil und Herabsetzung der Defektquote von p auf p_0 .

Lösung

Die Wahrscheinlichkeit, dass von den n Bauteilen genau i defekt sind, ist durch die Binomialverteilung gegeben

$$p_i = \binom{n}{i} p^i (1-p)^{n-i}.$$

Für die Frage a) folgt damit

$$C_a = \sum_{i=1}^n k_i p_i = \sum_{i=1}^n k_i \binom{n}{i} p^i (1-p)^{n-i}. \quad (5.1)$$

Für die Frage b) gilt ferner

$$C_b = nk + \sum_{i=1}^n k_i \binom{n}{i} p_0^i (1-p_0)^{n-i}.$$

Der Gewinn (positiv oder negativ) folgt dann aus der Differenz $C_a - C_b$.

Mit ähnlichen Überlegungen wie jene im Beispiel 5.1 folgt z.B. für den *Mittelwert* der gesamten Reparaturkosten C_{CM} während der Anwendungsdauer T_{UL} einer Baugruppe mit Ausfallrate λ und mittleren Reparaturkosten k_{CM} per Ausfall, den Wert

$$C_{CM} = \lambda T_{UL} k_{CM} \quad (5.2)$$

Das Modell, welches auf die Darstellung im Bild 5.1 geführt hat, stützt sich auf folgender Gleichung für den Erwartungswert C der Kosten für die Qualitätssicherung und zur Sicherstellung der Zuverlässigkeit, Instandhaltbarkeit und Logistik von komplexen Systemen [2, 1986].

$$C = C_Q g_Q^{m_Q} + C_R g_R^{m_R} + C_{CM} g_{CM}^{m_{CM}} + C_{PM} g_{PM}^{m_{PM}} + C_L g_L^{m_L} + \frac{T_{UL} k_{CM}}{MTBF_S g_R} + (1 - AA_S) T_{UL} k_D + \frac{k_E}{g_Q^{m_Q} m_E}. \quad (5.3)$$

In der Gleichung (5.3) haben die Indices folgende Bedeutung

CM = Reparatur (Corrective Maintenance)

L = Logistik

PM = Wartung (Preventive Maintenance)

Q = Qualität

R = Zuverlässigkeit (Reliability)

$MTBF$ und AA_S sind die Zielvorgaben für den Mittelwert der ausfallfreien Arbeitszeit und für die Verfügbarkeit (stationärer Wert der durchschnittlichen Verfügbarkeit); T_{UL} ist die Anwendungsdauer des Gerätes bzw. Systems (Useful Life); n ist die Anzahl versteckter Defekte, die beim Kunden entdeckt werden; k_{CM} , k_D und k_E sind die Kosten für eine Reparatur, für eine Stunde Stillstandzeit und. pro versteckter Defekt; C_Q , C_R , C_{CM} , C_{PM} und C_L sind die Kosten zur Sicherstellung der vorgegebenen Ziele für die Qualität, Zuverlässigkeit, Instandsetzbarkeit, Wartbarkeit und. Logistik (vgl. Bild 1.3); m_i sind Parameter und g_i sind die Variablen nach welchen die Kosten minimiert werden sollen (g_i = erreichtes Niveau durch vorgegebenes Ziel für die betrachtete Grösse i). Die ersten fünf Glieder der Gl. (5.3) stellen einen Teil der *Anschaffungskosten* dar, die übrigen Glieder tragen zu den *Folgekosten* bei. Die Untersuchung hat gezeigt, dass das Modell gemäss Gl. (5.3) wenig empfindlich auf die Parameter m_i ist.

5.3 Energieverbrauch

Bei der Entwicklung und Herstellung von Produkten soll zunehmend die Umweltverträglichkeit berücksichtigt werden (von der Rohstoffaufbereitung bis zur Produktentsorgung). Ein wichtiger Aspekt ist hier auch der Energieverbrauch. Dieser hängt für elektronische Geräte stark von den verwendeten Technologien aber speziell vom Vorhandensein von Redundanz auf hoher Integrationsebene (Gerät, Anlage) ab. Die Verdoppelung einer Computeranlage, z.B. infolge einer übertriebenen Zuverlässigkeits- oder Verfügbarkeitsforderung, verlangt praktisch auch eine Verdoppelung der Infrastruktur, insbesondere der Klimaanlage, was auf eine Verdoppelung des gesamten Energieverbrauchs führen kann (die Klimaanlage braucht in der Regel doppelt so viel Energie wie der Computer selbst). Redundanz kann aber oft gezielt auf Modul- oder Geräteebe- eingesetzt werden, und in manchen Fällen kann man durch geeignete Ein- und Abschaltstrategien sogar auf sie verzichten. Bei der Festlegung von Zuverlässigkeitsforderungen ist es deshalb wichtig, sich ein Bild über ihre Tragweite zu machen. Wie im Falle der Lebenslaufkosten (vgl. Bild 1.3 und Anhang A1) und der Umweltverträglichkeit, ist die Schätzung der Tragweite von Zuverlässigkeitsforderungen in einer frühen Projektphase oft schwierig. Die Kapitel 2 und 3 dieser Monographie geben die Werkzeuge für eine solche Schätzung. Für konkretere Beispiele kann auf Kapitel 4 verwiesen werden.

5.4 Festlegung von Zuverlässigkeitsforderungen

Bei der Festlegung der quantitativen, projektspezifischen Qualitäts- und Zuverlässigkeitsforderungen müssen von Anfang an auch deren Realisierungs- und Nachweismöglichkeiten berücksichtigt werden. Die Forderungen ergeben sich aus den Kundenwünschen, bzw. aus den Marktbedürfnissen, unter Berücksichtigung der technischen und finanziellen Grenzen sowie, wie im Abschnitt 5.3 erwähnt, der Umweltprobleme, inkl. Energieverbrauch. Als Beispiel sollen in diesem Abschnitt die wichtigsten Aspekte bei der Formulierung der Forderungen für die *MTBF*, die *MTTR* und den stationären Wert der Verfügbarkeit $PA = MTBF / (MTBF + MTTR)$ betrachtet werden ($MTBF = 1/\lambda$ = Mittelwert der ausfallfreien Arbeitszeit im Falle konstanter Ausfallrate λ , *MTTR* = Mittelwert der Reparaturzeit).

Provisorische Ziele für die *MTBF*, die *MTTR* und die Verfügbarkeit *PA* werden anhand folgender Überlegungen aufgestellt:

- Operationelle Anforderungen bezüglich Zuverlässigkeit, Instandhaltbarkeit, Verfügbarkeit
- Vorgesehene logistische Unterstützung
- Geforderte Funktion und vorgesehene Umweltbedingungen
- Erfahrungen mit ähnlichen Geräten und Systemen
- Vorhandensein von Redundanz
- Bedingungen für die Lebenslaufkosten
- Tragweite der Forderungen bezüglich Kosten, Energieverbrauch, usw. (vgl. die Bemerkungen der Abschnitte 5.2 und 5.3)

Typische Ausfallrate für elektronische Geräte liegen zwischen etwa 500 und $10'000 \cdot 10^{-9} \text{h}^{-1}$ für eine Umgebungstemperatur θ_A von 35°C und einen *Einsatzfaktor* $d = 0.3$ und sind etwa 50 bis

100 mal grösser für Systeme, vgl. Tab. 5.2 für konkretere Richtwerte. Der Einsatzfaktor ($0 < d \leq 1$) berücksichtigt, dass das Gerät im Mittel

$$d \cdot 100 \%$$

der Zeit in Betrieb ist. Genauere Angaben über Ausfallrate von Geräten der Nachrichtentechnik können aus Tabelle 2.2 entnommen werden. Anstelle der Ausfallrate λ kann auch die *mittlere Anzahl Ausfälle per 100 Geräte (oder in %) und Jahr* m angegeben werden

$$m = \lambda \cdot 8'600 \text{ h} \cdot 100\% \approx \lambda \cdot 10^6 \text{ h}. \quad (5.2)$$

Tabelle 5.2 Richtwerte für Ausfallraten von Geräten und Systemen $\lambda = 1/\text{MTBF}$ in 10^{-9} h^{-1} und mittlere Anzahl m Ausfälle in % und pro Jahr (d = Einsatzfaktor, $\theta_A = 35^\circ\text{C}$, Bauteile in Industriequalität)

Gerät	d=100%		d=30%	
	λ	m	λ	m
Haustelephonzentrale	6000	6	2000	2
Komfort-Telephongerät	600	0.6	200	0.2
Kopiergerät (inkl. Mechanik)	900'000	900	300'000	300
PC	9000	9	3000	3
Radaranlage (Boden)	600'000	600	200'000	200
Steuerung (SPS)	3'000	3	1000	1
Grosscomputer	150'000	150	—	—

Werte von $m \leq 1\%$ stellen hohe Anforderungen für elektronische Geräte und können sich merklich auf den Gerätepreis auswirken.

Mit Hilfe von *Grobanalysen* und *Vergleichsstudien* werden die Ziele Schritt für Schritt verfeinert. Dabei ist es oft von Vorteil, wenn die Ziele bis zum Niveau Baugruppen aufgeteilt werden. Für den *Nachweis einer MTBF* $= 1/\lambda$ sind folgende Angaben wichtig:

- quantitative Forderungen (*MTBF*=minimal akzeptierbare *MTBF* (bzw. oberste Grenze der akzeptierbaren *MTBF*) oder $MTBF_0$ =spezifizierte *MTBF* und $MTBF_1$ =minimal akzeptierbare *MTBF*)
- geforderte Funktion (inkl. Zeitverlauf)
- Umweltbedingungen (Temperaturzyklen, Ein-/Ausschaltzyklen, Vibrationen usw.)
- Fehler 1. und/oder 2. Art (α und/oder β)
- Annahmebedingungen (Dauer der Prüfung, Anzahl zugelassener Ausfälle)
- Anzahl der zur Verfügung stehenden Geräte oder Systeme
- Parameter, die überprüft werden müssen, und Frequenz der Überprüfung
- Ausfälle, die nicht berücksichtigt werden
- Wartung und Vorbehandlung vor der Prüfung
- Instandhaltungsprozedur
- Protokolle und Berichte

- Bestimmungen im Falle eines negativen Prüfergebnisses.

Für den *Nachweis einer MTTR* sind folgende Angaben wichtig:

- quantitative Forderungen (*MTTR*, Varianz, Quantil)
- Prüfbedingungen (Personal, Werkzeuge, externe Prüfmittel, Ersatzteile)
- Anzahl und Art der durchzuführenden Reparaturen (simulierte Ausfälle)
- Einteilung der gesamten Reparaturzeit (Ausfallerkennung, -lokalisierung, -behebung, Funktionsprüfung, Logistikzeit)
- Annahmebedingungen
- Protokolle und Berichte
- Bestimmungen im Falle eines negativen Prüfergebnisses.

Der *Nachweis der Verfügbarkeit* erfolgt in der Regel indirekt über die Beziehung $PA = MTBF / (MTBF + MTTR)$. Für einen direkten Nachweis kann man sich auf die Prozedur für den Nachweis einer unbekannten Wahrscheinlichkeit p stützen.

5.5 Durchsetzung von Zuverlässigkeitsforderungen

Für einfache Geräte und Baugruppen genügt es oft, dass die Zuverlässigkeitsforderungen im Pflichtenheft aufgenommen werden und vom Entwicklungsingenieur in Zusammenarbeit mit den Zuverlässigkeitsfachleuten berücksichtigt werden. Für komplexe Geräte sowie für Anlagen und Systeme müssen die Aktivitäten der Qualitäts- und Zuverlässigkeitssicherung geplant und in jenen der Linienstellen als Arbeitspaket im gesamten Projektplan integriert werden. Für grosse Vorhaben erfolgt eine solche Integration im Rahmen eines Qualitäts- und Zuverlässigkeitssicherungsprogramms.

Das Qualitäts- und Zuverlässigkeitssicherungsprogramm fasst die projektbezogenen Aktivitäten zur Sicherstellung der Qualität und Zuverlässigkeit einer Betrachtungseinheit zusammen und legt die Abgrenzungen in Bezug auf Umfang und Zuständigkeiten fest. Es ist zweckmässig, zwei getrennte Qualitäts- und Zuverlässigkeitssicherungsprogramme für die Entwicklungs- und Fertigungsphase zu erstellen. Diese Programme müssen projektspezifisch formuliert und in die Arbeitspakete der Linienstellen integriert werden. Im folgenden werden die wichtigsten Aspekte bei der Erstellung und Beurteilung eines Qualitäts- und Zuverlässigkeitssicherungsprogramms für komplexe Geräte und Systeme dargelegt. Für einfachere Betrachtungseinheiten sind entsprechende Anpassungen bezüglich Umfang und Tiefe der einzelnen Aktivitäten notwendig.

Als Checkliste zur Erstellung eines Qualitäts- und Zuverlässigkeitssicherungsprogramms für komplexe Geräte und Systeme kann die Tabelle 5.3 verwendet werden. Diese Tabelle fasst die Aufgaben zur Sicherstellung der Qualität und Zuverlässigkeit während der Entwicklungs- und Fertigungsphase zusammen und gibt die prinzipielle Zuteilung der Aufgaben zu den Linienstellen an. Die wichtigsten Elemente eines Qualitäts- und Zuverlässigkeitssicherungsprogramms sind:

1. Projektorganisation, Projektplanung und Projektablauf
2. Zuverlässigkeits- und Sicherheitsanalysen
3. Wahl und Qualifikation von Bauteilen, Stoffen, Fertigungsprozessen und -abläufen

4. Konfigurationsmanagement
5. Qualitätsprüfungen
6. Qualitätsdatensystem

Im folgenden wird auf den Inhalt bzw. auf die Beurteilung eines Qualitäts- und Zuverlässigkeitssicherungsprogramms für die Entwicklung komplexer Geräte oder Systeme eingegangen. Wie aus Tab. 5.3 zu entnehmen ist, liegt die Kompetenz für die Erstellung des Qualitäts- und Zuverlässigkeitssicherungsprogramms bei der Qualitäts- und Zuverlässigkeitssicherungsstelle. Die Durchführung bzw. Nachführung dieses Programms ist hingegen Aufgabe des Projektleiters.

5.5.1 Projektorganisation, Projektplanung, Projektablauf

Die Durchführung eines Qualitäts- und Zuverlässigkeitssicherungsprogramms setzt einen klaren Projektablauf sowie eine transparente Projektorganisation und Projektplanung voraus. Ausgangspunkt dabei ist das Pflichtenheft. Folgendes Beispiel gibt das Inhaltsverzeichnis eines Pflichtenhefts für komplexe Geräte und Systeme an:

- 1 . Veranlassung
- 2 . Ziele
- 3 . Kosten, Termine
- 4 . Marktchance (Umsatz, Preis, Konkurrenzsituation)
- 5 . Technische Eigenschaften
- 6 . Umweltbedingungen
- 7 . Operationelle Eigenschaften
- 8 . Qualitäts- und Zuverlässigkeitssicherung
- 9 . Spezielle Aspekte (Technologieengpässe, Patentwesen, Wertanalyse usw.)
10. Beilagen

Für komplexe Systeme oder im Fall grosser Abnehmer kann der Kunde auf den Inhalt der Abschnitte 5 bis 8 des Pflichtenhefts einwirken. Ausgehend vom Pflichtenheft wird das Projekt in Arbeitspakete aufgeteilt. Aus der Strukturierung lässt sich die Projektorganisation ableiten, woraus wiederum die Tätigkeitslisten sowie der Netzplan, der Balkenplan und die Meilensteinliste aufgestellt werden können. Schliesslich bleibt die Quantifizierung der personellen, materiellen und finanziellen Mittel für die Durchführung des Projekts. Das Qualitäts- und Zuverlässigkeitssicherungsprogramm muss sicherstellen, dass das Projekt formell richtig geplant und organisiert ist, und dass die personellen Besetzungen vorgenommen worden sind.

Tabelle 5.3 Prinzipielle Zuteilung der Aufgaben zur Sicherstellung der Qualität und Zuverlässigkeit komplexer Geräte und Systeme (Checkliste zur Erstellung eines projektspezifischen Qualitäts- und Zuverlässigkeitssicherungsprogramms)

Zuteilung der Aufgaben zur Sicherstellung der Qualität und Zuverlässigkeit komplexer Geräte und Systeme (Checkliste für die Festlegung der projektspezifischen Aufgaben) Abkürzungen: F = Federführung, M = Mitwirkung, I = Information	Verkauf	Entwicklung	Fertigung	Zentrale Stelle
1. Ermittlung der Markt- bzw. Kundenforderungen				
1 Ermittlung der Beurteilung der ausgelieferten Geräte	F	I	I	M
2 Ermittlung der Wünsche und Bedürfnisse des Markts bzw. der Kunden	F	I	I	M
3 Kundenberatung	F			M
2. Durchführung von Grobanalysen				
1 Festlegung provisorischer Ziele für die Zuverlässigkeit, Instandhaltbarkeit, Verfügbarkeit, Sicherheit und für das Qualitätsniveau	M	M	M	F
2 Durchführung von Grobanalysen und Bestimmung potentieller Probleme	I	M		F
3 Durchführung von Vergleichsstudien	I	M		F
3. Erstellung bzw. Überprüfung von Pflichtenheften, Offerten usw.				
1 Definition der geforderten Funktion	I	F		M
2 Festlegung der Umweltbedingungen	M	F		M
3 Festlegung realistischer Forderungen für die Zuverlässigkeit, Instandhaltbarkeit, Verfügbarkeit, Sicherheit und für das Qualitätsniveau	M	M	M	F
4 Festlegung der Prüf- und Nachweiskriterien	M	M	M	F
5 Abklärung der Möglichkeit, Daten aus der Nutzungsphase zu erhalten	F			M
6 Schätzung der Kosten für die Sicherstellung der Qualität und der Zuverlässigkeit	M	M	M	F
4. Erstellung des Qualitäts- und Zuverlässigkeitssicherungsprogramms				
1 Erstellung	M	M	M	F
2 Nachführung				
– Entwicklungsphase	I	F	I	M
– Fertigungsphase	I	I	F	M
5. Analyse der Zuverlässigkeit und der Instandhaltbarkeit				
1 Spezifizierung der geforderten Funktion der einzelnen Elemente		F		M
2 Festlegung der umweltbedingten, funktionsbedingten und zeitbedingten Beanspruchungen (detaillierte Arbeitsbedingungen)		F		M
3 Festlegung der Belastungsfaktoren		M		F
4 Aufteilung der Zuverlässigkeit und der Instandhaltbarkeit		M		F
5 Aufstellung der Zuverlässigkeitsblockdiagramme				
– Niveau Baugruppe		F		M
– Niveau Gerät und System		M		F
6 Identifikation und Analyse von Schwachstellen (FMEA/FMECA, Worst-Case, Driftanalysen, Stress-Strength-Analysen)				
– Niveau Baugruppe		F		M
– Niveau Gerät und System		M		F
7 Durchführung von Vergleichsstudien				
– Niveau Baugruppe		F		M
– Niveau Gerät und System		M		F

Tabelle 5.3 (Fortsetzung)

9. Erstellung der projektabhängigen Spezifikationen				
1 Zuverlässigkeits-Richtlinien		M		F
2 Instandhaltbarkeits-Richtlinien	M	M	I	F
3 Sicherheits-Richtlinien	I	M	I	F
4 Andere Vorschriften, Spezifikationen usw. für die Linienstellen der Entwicklung		F	I	M
5 Arbeitsanweisungen, Vorschriften usw. für die Linienstellen der Fertigung		I	F	M
6 Überwachung der Einhaltung	M	M	M	F
10. Planung und Steuerung von Dokumentation und Bauzustand				
1 Planung und Überwachung des Konfigurations- managements	M	M	M	F
2 Durchführung des Konfigurationsmanagements				
– Identifikation und Überwachung der Konfiguration				
• während der Entwicklungsphase		F		M
• während der Fertigungsphase			F	M
• während der Nutzungsphase (Garantieperiode)	F	I	I	M
– Überprüfung der Konfiguration (Entwurfs- überprüfungen)				
• während der Definitionsphase	M	F	M	M
• während der Entwicklungsphase	I	F	I	M
• während der Vorserienphase	M	F	M	M
– Steuerung der Konfiguration (Evaluierung, Koordinierung und Annahme resp. Verwerfung der Änderungsvorschläge)	M	F	M	M
11. Qualifikation der Prototypen				
1 Planung	I	F	I	M
2 Durchführung und Berichterstattung		F	M	M
3 Analyse der Prüfergebnisse	M	F	I	M
4 Spezielle Zuverlässigkeits-, Instandhaltbarkeits- und Sicherheitsprüfungen	I	M	I	F
12. Sicherstellung der Fertigungsprozesse und -abläufe				
1 Wahl und Qualifikation von Prozessen und Verfahren		F	M	M
2 Fertigungsplanung und -steuerung		M	F	M
3 Laufende Überprüfung bzw. Überwachung von Betriebsmitteln, Produktionsprozessen, Arbeits- einflüssen, Handhabung, Lagerung und Verpackung		I	F	M
13. Durchführung von Zwischenprüfungen				
1 Planung		M	F	M
2 Durchführung und Berichterstattung		I	F	I
14. Durchführung der Endprüfung/Abnahmeprüfung				
1 Umweltprüfungen und/oder Vorbehandlung von Serieneinheiten				
– Planung	I	M	M	F
– Durchführung und Berichterstattung	I	I	M	F
– Analyse der Prüfergebnisse	I	M	M	F
2 Endprüfung/Abnahmeprüfung				
– Planung	M	M	M	F
– Durchführung und Berichterstattung	I	I	M	F
– Analyse der Prüfergebnisse	I	M	M	F
3 Beschaffung, Instandhaltung und Eichung von Prüfgeräten und Prüfeinrichtungen (auch für 7. 11 und 13)	I	M	M	F

Tabelle 5.3 (Fortsetzung)

9. Erstellung der projektabhängigen Spezifikationen				
1 Zuverlässigkeits-Richtlinien		M		F
2 Instandhaltbarkeits-Richtlinien	M	M	I	F
3 Sicherheits-Richtlinien	I	M	I	F
4 Andere Vorschriften, Spezifikationen usw. für die Linienstellen der Entwicklung		F	I	M
5 Arbeitsanweisungen, Vorschriften usw. für die Linienstellen der Fertigung		I	F	M
6 Überwachung der Einhaltung	M	M	M	F
10. Planung und Steuerung von Dokumentation und Bauzustand				
1 Planung und Überwachung des Konfigurations- managements	M	M	M	F
2 Durchführung des Konfigurationsmanagements				
– Identifikation und Überwachung der Konfiguration				
• während der Entwicklungsphase		F		M
• während der Fertigungsphase			F	M
• während der Nutzungsphase (Garantieperiode)	F	I	I	M
– Überprüfung der Konfiguration (Entwurfs- überprüfungen)				
• während der Definitionsphase	M	F	M	M
• während der Entwicklungsphase	I	F	I	M
• während der Vorserienphase	M	F	M	M
– Steuerung der Konfiguration (Evaluierung, Koordinierung und Annahme resp. Verwerfung der Änderungsvorschläge)	M	F	M	M
11. Qualifikation der Prototypen				
1 Planung	I	F	I	M
2 Durchführung und Berichterstattung		F	M	M
3 Analyse der Prüfergebnisse	M	F	I	M
4 Spezielle Zuverlässigkeits-, Instandhaltbarkeits- und Sicherheitsprüfungen	I	M	I	F
12. Sicherstellung der Fertigungsprozesse und -abläufe				
1 Wahl und Qualifikation von Prozessen und Verfahren		F	M	M
2 Fertigungsplanung und steuerung		M	F	M
3 Laufende Überprüfung bzw. Überwachung von Betriebsmitteln, Produktionsprozessen, Arbeits- einflüssen, Handhabung, Lagerung und Verpackung		I	F	M
13. Durchführung von Zwischenprüfungen				
1 Planung		M	F	M
2 Durchführung und Berichterstattung		I	F	I
14. Durchführung der Endprüfung/Abnahmeprüfung				
1 Umweltprüfungen und/oder Vorbehandlung von Serieneinheiten				
– Planung	I	M	M	F
– Durchführung und Berichterstattung	I	I	M	F
– Analyse der Prüfergebnisse	I	M	M	F
2 Endprüfung/Abnahmeprüfung				
– Planung	M	M	M	F
– Durchführung und Berichterstattung	I	I	M	F
– Analyse der Prüfergebnisse	I	M	M	F
3 Beschaffung, Instandhaltung und Eichung von Prüfgeräten und Prüfeinrichtungen (auch für 7. 11 und 13)	I	M	M	F

Tabelle 5.3 (Fortsetzung)

<i>15. Erfassung, Analyse und Korrektur der Defekte und Ausfälle</i>				
1 Erfassung	M	M	M	F
2 Analyse und Entscheid über Abhilfen während der				
– Qualifikationsphase		F		M
– Zwischenprüfung		M	F	M
– Endprüfung/Abnahmeprüfung	M	M	M	F
– Nutzung (Garantieperiode)	F	M	M	M
3 Durchführung der Korrekturmaßnahmen an der Hardware bzw. der Computer-Software (Änderung, Reparatur, Nacharbeit)	I	M	M	F
4 Durchführung der Korrekturmaßnahmen an der Dokumentation	M	M	M	F
5 Verdichtung, Speicherung, Auswertung und Rückkopplung der Information	I	I	I	F
6 Steuerung des Qualitätsdatensystems	I	I	I	F
<i>16. Logistische Unterstützung</i>				
1 Bereitstellung der speziellen Meß- und Prüfgeräte für die Instandhaltung	M	F	I	M
2 Erstellung der Kundendokumentation	F	I	I	I
3 Schulung des Bedienungs- und Instandhaltungspersonals	F	I		I
4 Ermittlung des Bedarfs an Ersatzteilen, Instandhaltungspersonal usw.	F	M		M
5 Kundendienst	F	I	I	M
<i>17. Durchführung von Koordinations- und Überwachungsaufgaben</i>				
1 Projektspezifisch	M	M	M	F
2 Projektunabhängig	I	I	I	F
3 Durchführung unabhängiger Überprüfungen (Audits)				
– projektspezifisch	M	M	M	F
– projektunabhängig	I	I	I	F
4 Informationsrückkopplung	I	I	I	F
<i>18. Erfassung und Analyse der Qualitätskosten</i>				
1 Erfassung der Kosten	M	M	M	F
2 Analyse der Kosten und nötigenfalls Einleitung von Maßnahmen	M	M	M	F
3 Erstellung periodischer und spezieller Berichte	I	I	I	F
4 Prüfung der Wirtschaftlichkeit	I	I	I	F
<i>19. Erarbeitung von Konzepten, Methoden und Verfahren</i>				
1 Erarbeitung von Konzepten und Verfahren	M	M	M	F
2 Erstellung und Nachführung des Qualitäts- und Zuverlässigkeitssicherungshandbuchs	M	M	M	F
3 Untersuchung von Methoden	I	I	I	F
4 Bereitstellung von Computer-Programmen für Analysen und Auswertungen	I	I	I	F
5 Erfassung, Auswertung, Speicherung und Verteilung von Daten, Erfahrung und Know-how	I	I	I	F
<i>20. Motivation und Schulung der Linienstellen</i>				
1 Aufstellung des internen Aus- und Weiterbildungsprogramms	M	M	M	F
2 Erstellung der Lehrgänge	M	M	M	F
3 Durchführung bzw. Teilnahme	M	M	M	F

5.5.2 Zuverlässigkeits- und Sicherheitsanalysen

Zuverlässigkeits- und Sicherheitsanalysen umfassen Ausfallrateanalysen, Ausfallartanalysen (FMEA/FMECA, FTA, vgl. Abschnitt 2.2), Untersuchung der konkreten Möglichkeiten zur Verbesserung der Zuverlässigkeit und der Sicherheit (Unterlastung, Vorbehandlung, Redundanz) und Vergleichsstudien. Die Methoden dazu sind in Kapitel 2 (in Kapitel 3 für reparierbare Systeme) dargelegt. Das Qualitäts- und Zuverlässigkeitssicherungsprogramm soll angeben, was für das betreffende Projekt konkret getan wird. Insbesondere soll es folgende Fragen beantworten können:

1. Wie werden die Belastungsfaktoren festgelegt?
2. Welche Ausfallratenkataloge werden verwendet? Wie sind die verschiedenen Faktoren bestimmt worden?
3. Wie werden die Arbeitsbedingungen auf Bauteilebene bestimmt?
4. Wie lautet die Prozedur für die Ausfallartanalyse? Für welche Elemente wird sie durchgeführt?
5. Welche Arten von Vergleichsstudien werden durchgeführt?

Zudem sollen die Schnittstellen zur Wahl und Qualifikation von Bauteilen und Stoffen, zu den Entwurfsüberprüfungen, zum Prüfkonzept, zum Vorbehandlungskonzept, zu den Zuverlässigkeitsprüfungen, zum Qualitätsdatensystem und zu den Aktivitäten bei den Unterlieferanten aufgezeigt werden. Die Unterlagen für die Berechnung der Ausfallraten der Bauteile müssen kritisch beurteilt werden (Quelle, Aktualität, Festlegung der Faktoren).

5.5.3 Wahl und Qualifikation von Bauteilen, Stoffen, Fertigungsprozessen und -abläufen

Bauteile und Stoffe sowie Fertigungsprozesse und -abläufe haben einen direkten Einfluss auf die Qualität und Zuverlässigkeit der fertiggestellten Geräte oder Systeme. Sie müssen deshalb sorgfältig ausgewählt und qualifiziert werden. Die Kriterien und Methoden zur Wahl und Qualifikation elektronischer Bauteile werden in [2, 1991] ausführlich dargelegt. Aus dem Qualitäts- und Zuverlässigkeitssicherungsprogramm soll hervorgehen, wie Bauteile, Stoffe und Prozesse ausgewählt und qualifiziert werden. Insbesondere soll es folgende Fragen beantworten können:

1. Existiert eine Liste der bevorzugten Bauteile und Stoffe? Wie wurde diese Liste aufgestellt? Sind die Annahmen für die Erstellung dieser Liste vereinbar mit den Umwelt- und Belastungsbedingungen in dem betreffenden Projekt? Sind die Aspekte des Zweitlieferanten und der langfristigen Beschaffung berücksichtigt worden?
2. Wie werden neue Bauteile und Stoffe qualifiziert? Wie sieht das Freigabeverfahren aus.?
3. Unter welchen Bedingungen kann ein Entwicklungsingenieur oder ein Konstrukteur nicht qualifizierte Bauteile und Stoffe verwenden?
4. Wie sind die Standard-Fertigungsprozesse qualifiziert worden?
5. Wie werden die speziellen Fertigungsprozesse qualifiziert?

Spezielle Fertigungsprozesse sind solche, deren Qualität nicht direkt am Gerät oder System überprüft werden kann (z. B. chemisch-metallurgische Verfahren) und solche, welche hohe An-

Forderungen bezüglich Reproduzierbarkeit erfüllen müssen oder deren Anwendung negative Einflüsse auf die Qualität oder Zuverlässigkeit des Geräts oder Systems haben können.

5.5.4 Konfigurationsmanagement

Das Konfigurationsmanagement ist ein wichtiges Werkzeug der Qualitätssicherung. Im Rahmen eines Projekts wird es in Beschreibung, Überprüfung, Steuerung und Überwachung der Konfiguration aufgeteilt. Die Beschreibung der Konfiguration wird in der Dokumentation festgehalten. Für komplexe Geräte und Systeme wird in der Regel die Dokumentation in Projektdokumentation, technische Dokumentation, Fertigungsdokumentation und Kundendokumentation eingeteilt v. 1. Bild 5.2.

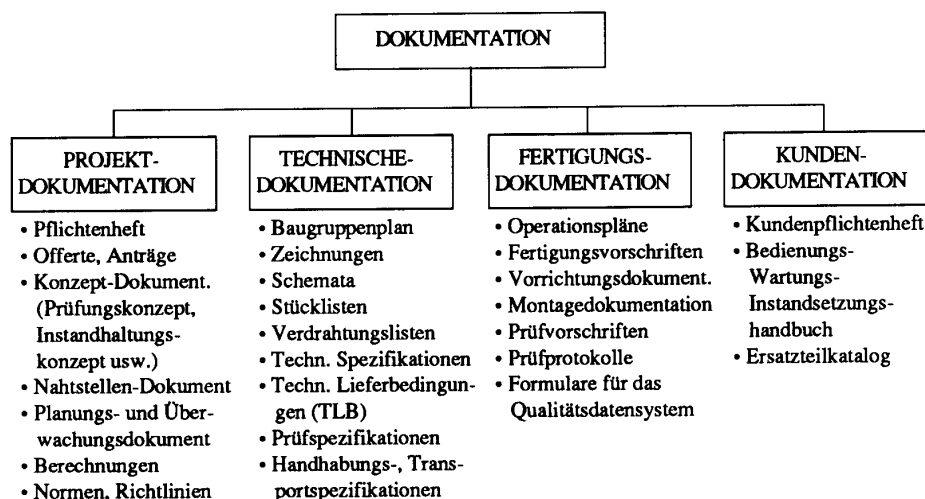


Bild 5.2 Einteilung der Dokumentation im Fall komplexer Geräte und Systeme

Zur Überprüfung der Konfiguration werden Entwurfsüberprüfungen (Design Reviews) durchgeführt. Ihr Ziel ist es, sicherzustellen, dass das Gerät bzw. System den gestellten Anforderungen genügt. Dabei werden mit Hilfe von Checklisten die wichtigsten Aspekte der Entwicklung und der Konstruktion (Wahl von Bauteilen und Stoffen, Dimensionierung, Schnittstellenprobleme, Konstruktionsprobleme usw.), der Fabrikation (Herstellbarkeit, Prüfbarkeit, Reproduzierbarkeit), der Zuverlässigkeit, der Instandhaltbarkeit, der Sicherheit, des Patentwesens, des Value-Engineerings und der Wertanalyse kritisch beurteilt. Die wichtigsten Entwurfsüberprüfungen sind in Tab. 5.4 angegeben. Die Entwurfsüberprüfungen stehen unter der Führung des Projektleiters. Die Mitarbeit des Projekt-Verantwortlichen für die Qualitäts- und Zuverlässigkeitssicherung ist notwendig. An den wichtigen Entwurfsüberprüfungen sollen sich auch folgende Personen beteiligen:

- * die betroffenen Teilprojektleiter
- * Sachbearbeiter, die über die verschiedenen Aspekte berichten sollen
- * je ein Vertreter aus Fertigung und Verkauf
- * ein bis zwei unvoreingenommene erfahrene Entwickler oder externe Fachleute
- * ein bis zwei Vertreter des Auftraggebers (falls gefordert).

Tabelle 5.4 Wichtige Entwurfsüberprüfungen (Design Reviews) während der Definitions- und Entwicklungsphase

	System-Entwurfsüberprüfung (System Design Review, SDR)	Vorläufige Entwurfsüberprüfungen (Preliminary Design Reviews, PDR)	Kritische Entwurfsüberprüfung (Critical Design Review, CDR)
Durchführung	Am Ende der Definitionsphase	Während der Entwicklungsphase (laufend, jedesmal wenn eine Baugruppe fertig entwickelt ist)	Am Ende der Qualifikation der Prototypen
Ziel(e)	Kritische Überprüfung des Pflichtenhefts anhand der Erkenntnisse aus der Evaluation von Marktforschungen, Grobanalysen, Vergleichsstudien, Patentlage usw.	- Kritische Überprüfung der erstellten Unterlagen (Berechnungen, Schaltungen, Zeichnungen, Prüfspezifikation, usw.) - Vergleich der erhaltenen Resultate mit den Pflichtenheft-Forderungen - Überprüfung der Schnittstellen mit anderen Baugruppen	- Kritischer Vergleich der Ergebnisse aus der Qualifikation der Prototypen mit den Pflichtenheft-Forderungen - Formale Überprüfung der Übereinstimmung der hergestellten Prototypen mit der technischen Dokumentation - Überprüfung der Herstellbarkeit, Prüfbarkeit, Reproduzierbarkeit
Eingang (Input)	- Traktandenliste - Pflichtenheft (Entwurf) - Dokumentation (Berechnungen, Begründungen usw.) - Checklisten	- Traktandenliste - Dokumentation (Berechnungen, Schemata, Zeichnungen, Stücklisten, Prüfspezifikationen, Baugruppenplan, Schnittstellenspezifikationen usw.) - Protokolle der früheren vorläufigen Entwurfsüberprüfungen - Checklisten	- Traktandenliste - Technische Dokumentation - Prüfplan und Prüfvorschriften für die Qualifikation der Prototypen - Ergebnisse der Qualifikationsprüfung der Prototypen - Liste der Abweichungen zu den Pflichtenheft-Forderungen - Instandhaltungskonzept - Checklisten
Ausgang (Output)	- Pflichtenheft - Projektantrag für die Entwicklungsphase - Schnittstellenspezifikation - grobes Instandhaltungs- und Logistikkonzept - Protokoll	- Referenzen-Konfiguration der betroffenen Baugruppe(n) - Liste der Abweichungen zu den Pflichtenheft-Forderungen - Protokoll	- Liste definitiver Abweichungen zu den Pflichtenheft-Forderungen - Bereinigtes Pflichtenheft - Qualifizierte und freigegebene Prototypen - Eingefrorene techn. Dokumentation - Bereinigtes Instandhaltungskonzept - Vorserie-Antrag - Protokoll

Für jede Entwurfsüberprüfung sollen von den Teilnehmern projektbezogene Checklisten aufgestellt werden. Fragenkataloge als Grundlage zur Erstellung solcher Listen sind in der Tabelle 5.5 gegeben.

Die Steuerung der Konfiguration umfasst die systematische Evaluierung, Koordinierung und Annahme bzw. Verwerfung der vorgeschlagenen Änderungen und Modifikationen der Konfiguration. Dazu gehört insbesondere das Änderungswesen. Änderungen entstehen als Folge von Defekten oder Ausfällen; Modifikationen haben ihren Ursprung in einer Änderung im Pflichtenheft.

Mit der Überwachung der Konfiguration wird erreicht, dass alle genehmigten Änderungen und Modifikationen der Konfiguration realisiert und protokolliert werden. Ein definiertes Verfahren für die Überwachung der Konfiguration ist notwendig, wenn die Änderungen und Modifikationen sowohl in der Dokumentation als auch in der Hardware bzw. Software durchgeführt werden müssen. Dabei ist auf die ständige Übereinstimmung zwischen Dokumentation und Hardware bzw. Software (Bauzustandsüberwachung) sowie auf allfällige Kompatibilitätsprobleme zu achten. Eine lückenlose Verfolgung und Protokollierung aller Lebenslaufschritte einer bestimmten Betrachtungseinheit ist notwendig, um die Forderung nach der Verfolgbarkeit zu erfüllen.

Die Verfolgbarkeit wird selten für ein ganzes System gefordert, oft aber für jene Teile, welche die Sicherheit des Systems bestimmen können.

Dem Konfigurationsmanagement soll im Qualitäts- und Zuverlässigkeitssicherungsprogramm genügend Platz eingeräumt werden. Dabei stehen folgende Beurteilungsfragen im Vordergrund:

1. Welche Dokumente müssen wann, mit welchem Inhalt und durch welche Stelle erstellt werden?
2. Entspricht der Inhalt der Dokumente den Qualitätsforderungen?
3. Ist das Freigabeverfahren der Dokumente vereinbar mit den Qualitätsforderungen?
4. Sind die Änderungsverfahren der Projektdokumentation, der technischen Dokumentation, der Fertigungsdokumentation und der Kundendokumentation festgelegt? Ist die Aufwärtskompatibilität sichergestellt?
5. Wie wird die Bauzustandsüberwachung gewährleistet? Welche Betrachtungseinheiten sind der Verfolgbarkeit unterzogen?

5.5.5 Qualitätsprüfungen

Qualitätsprüfungen umfassen alle Aktivitäten zur Feststellung, inwieweit das Gerät bzw. System den gestellten Anforderungen genügt. Dazu gehören Eingangsprüfungen, Qualifikationsprüfungen, Zwischenprüfungen und Endprüfungen (inkl. Zuverlässigkeits-, Instandhaltbarkeits- und Sicherheitsprüfungen). Im Hinblick auf eine Optimierung des Prüfaufwands ist es wichtig, dass alle Prüfungen im Rahmen eines Prüf- und Vorbehandlungskonzepts auf Geräte- bzw. Systemebene geplant und durchgeführt werden. Das Qualitäts- und Zuverlässigkeitssicherungsprogramm muss das ganze Prüfkonzzept behandeln und folgende Fragen beantworten:

1. Wie sieht das Prüf- und Vorbehandlungskonzept auf Geräte- bzw. Systemebene aus? Wie wurde es freigegeben?
2. Wie werden die Unterlieferanten ausgewählt, qualifiziert und überwacht?
3. Was wird in den Beschaffungsunterlagen festgelegt? Wie wurden sie freigegeben?

Tabelle 5.5 Fragenkatalog als Hilfe bei der Aufstellung projektspezifischer Checklisten für die Entwurfsüberprüfungen bei der Entwicklung komplexer Geräte und Systeme

System-Entwurfsüberprüfung

1. Welche Erfahrungen mit ähnlichen Geräten und Systemen liegen vor?
2. Welches sind die Ziele für die Leistungsfähigkeit, Zuverlässigkeit, Instandhaltbarkeit, Verfügbarkeit und Sicherheit? Auf welches Anforderungsprofil (geforderte Funktion und Umweltbedingungen) beziehen sie sich?
3. Welche provisorische Aufteilung der Leistungsparameter, der Zuverlässigkeit und ggf. der Instandhaltbarkeit wurde vorgenommen?
4. Welches sind die kritischen Stellen? Liegen potentielle Probleme vor (neue Technologien, Schnittstellen)?
5. Sind Vergleichsstudien durchgeführt worden? Welche Resultate liegen vor?
6. Sind Stör- oder Entstörprobleme zu erwarten?
7. Liegen potentielle Probleme bezüglich Sicherheit vor?
8. Sind die Forderungen realistisch? Entsprechen sie den Marktbedürfnissen?
9. Liegt ein grobes Instandhaltungskonzept vor? Sind spezielle Forderungen bezüglich logistischer Unterstützung oder Ergonomie zu erwarten?
10. Sind spezielle Forderungen an die (Computer-) Software zu richten?
11. Ist die Patentlage überprüft worden? Sind Lizenzen notwendig? Sollen Patente angemeldet werden?
12. Liegen Schätzungen für die Lebenslaufkosten vor? Hat man diese unter Einbezug der Zuverlässigkeits- und Instandhaltbarkeitsforderungen überschlagsmäßig optimiert?
13. Liegt eine Rentabilitätsbetrachtung vor? Wie ist die Konkurrenzlage? Ist das Entwicklungsrisiko abgeschätzt worden?
14. Sind die geschätzten Termine realistisch? Kann das Gerät bzw. das System zum richtigen Zeitpunkt auf den Markt gebracht werden?
15. Sind Beschaffungsschwierigkeiten bei der Serienproduktion zu erwarten?

Vorläufige Entwurfsüberprüfung auf Baugruppenebene

- a) Allgemeines
1. Handelt es sich bei der betrachteten Baugruppe um eine Neuentwicklung, eine Weiterentwicklung oder eine Änderung bzw. Modifikation?
2. Können bestehende Elemente verwendet werden?
3. Wird die Baugruppe auch in anderen Geräten verwendet?
4. Liegen Erfahrungswerte aus ähnlichen Baugruppen vor? Welches waren die Probleme?
5. Liegt ein Blockschema bzw. ein Funktionsschema vor?
6. Sind die Pflichtenheft-Forderungen erfüllt? Wie wurde dies überprüft? Wo bestehen Abweichungen? Können einzelne Forderungen reduziert werden?
7. Gibt es Probleme mit dem Patentwesen? Müssen Lizenzen gekauft werden?
8. Werden die vorgesehenen Kosten und Termine eingehalten?
9. Wurden die Methoden des Value-Engineering angewandt?
10. Haben sich seit Beginn der Entwicklungsphase die Kunden- bzw. Marktbedürfnisse geändert?

Tabelle 5.5 (Fortsetzung)

b) Leistungsparameter

1. Welches sind die maßgebenden Leistungsparameter der betrachteten Baugruppe?
2. Wie wurde die Erfüllung der Forderungen bezüglich Leistungsparameter sichergestellt (Annahmen, Modelle, Berechnungen)?
3. Hat man bei den Berechnungen auch den Worst Case betrachtet?
4. Sind die Probleme der elektromagnetischen Verträglichkeit gelöst?
5. Wurden bei der Dimensionierung die vorgeschriebenen Richtlinien und Normen beachtet?
6. Sind die Schnittstellenprobleme mit den anderen Baugruppen gelöst?
7. Sind Versuchsaufbau und Funktionsmuster im Labor ausreichend erprobt worden?

c) (Computer-) Software

1. Sind die wichtigsten Qualitätsmerkmale und Entwicklungsrichtlinien schriftlich festgelegt worden? Wurden sie beachtet?
2. Sind die Detailspezifikationen genügend umfassend und klar?
3. Ist das Programm Modular, Top-down entwickelt und Bottom-up integriert und geprüft worden?
4. Sind die Schnittstellen zwischen den Programm-Modulen sauber definiert worden?
5. Sind das Programm und die Daten gegen Fehlbedienung, äußere Störungen und Überlast geschützt?
6. Ist die Software in ihrer endgültigen Hardware- und Software-Umgebung genügend erprobt worden? Wie?
7. Ist die Software ausreichend klar und korrekt dokumentiert?

d) Umwelteinflüsse

1. Sind die Umweltbedingungen, auch als Funktion der Zeit, definitiv festgelegt worden?
2. Wurden diese mit den internen Belastungen kombiniert und damit die Arbeitsbedingungen der einzelnen Elemente der Baugruppe ermittelt?
3. Sind die äußeren Störungen ermittelt oder angenommen worden? Ist der Einfluß dieser Störungen in den Berechnungen berücksichtigt worden (Worst Case)?

e) Bauteile und Stoffe

1. Bei welchen Bauteilen und Stoffen hat man sich nicht an die Liste der bevorzugten Bauteile und Stoffe halten können? Aus welchen Gründen? Wie wurde die Qualifikation dieser Bauteile und Stoffe durchgeführt?
2. Welche Bauteile und Stoffe werden vorbehandelt?
3. Sind für bestimmte Bauteile oder Stoffe spezielle Qualitätsklassen (z. B. Raumfahrt-Klasse) notwendig? Warum? Kann man diese Teile rechtzeitig beschaffen?
4. Ist die Beschaffung auch für die Serienproduktion sichergestellt? Ist mindestens ein Zweitlieferant vorhanden? Sind die Forderungen bezüglich Qualität, Zuverlässigkeit und Sicherheit für die zugelieferten Bauteile und Stoffe erfüllt?
5. Sind Probleme bei den Eingangsprüfungen zu erwarten?

f) Zuverlässigkeit

1. Ist es eine Neuentwicklung, eine Weiterentwicklung oder eine Änderung bzw. Modifikation?
2. Können bestehende Elemente verwendet werden?
3. Wird die Betrachtungseinheit auch in anderen Geräten verwendet?
4. Liegen Erfahrungsdaten aus ähnlichen Betrachtungseinheiten vor? Welches waren die Probleme?

. Ist eine Liste der bevorzugten Bauteile und Stoffe aufgestellt worden?

Tabelle 5.5 (Fortsetzung)

6. Ist die Wahl/Qualifikation der nicht standardisierten Bauteile und Stoffe geregelt? Wie?
7. Sind die maßgebenden Parameter jedes Elements der Betrachtungseinheit ausreichend definiert?
8. Sind die gegenseitigen Beziehungen und Beeinflussungen der verschiedenen Elemente berücksichtigt worden?
9. Sind die Pflichtenheft-Forderungen erfüllt? Können einzelne Forderungen reduziert werden?
10. Ist das Anforderungsprofil definiert und als Grundlage berücksichtigt worden?
11. Ist ein Zuverlässigkeitsblockdiagramm aufgestellt worden?
12. Sind die Umweltbedingungen ausreichend definiert? Wie wurden die Arbeitsbedingungen der einzelnen Elemente bestimmt?
13. Ist die Unterlastung konsequent durchgesetzt worden?
14. Wurde die Temperatur der elektronischen Bauteile so niedrig wie möglich gehalten?
15. Wurde der Einfluß von Ein- und Ausschaltvorgängen und von externen Störungen berücksichtigt?
16. Ist die Ausfallrate jedes Elements bekannt? Stimmt sie überein mit dem bei der Aufteilung der Zuverlässigkeit festgelegten Wert?
17. Sind Elemente mit beschränkter Lebensdauer vorhanden?
18. Ist die vorausgesagte Zuverlässigkeit berechnet worden?
19. Ist eine FMEA/ FMECA durchgeführt worden? Für welche Teile? Sind Sicherheitsprobleme vorhanden? Sind die Bedingungen für eine elektrische Zulassung erfüllt?
20. Sind Drift- und Worst-Case-Analysen durchgeführt worden?
21. Ist es notwendig, die Zuverlässigkeit durch Redundanz zu verbessern?
22. Sind Elemente vorhanden, die eine Vorbehandlung benötigen?
23. Sind Vereinfachungen des Entwurfs oder der Konstruktion möglich?
24. Wurde für eine rasche und einfache Erkennung, Lokalisierung und Behebung der Ausfälle gesorgt? Können verborgene Ausfälle auftreten? Ist der Aufbau modular?
25. Sind die Aspekte der Bedienungsfreundlichkeit berücksichtigt worden?
26. Liegt ein Programm für die Zuverlässigkeitsprüfungen vor? Was umfaßt es?
27. Sind die Fragen der Herstellbarkeit, Prüfbarkeit und Reproduzierbarkeit beachtet worden?
28. Sind die Beschaffungsprobleme (Zweitlieferanten, langfristige Lieferung) gelöst?
29. Sind die Probleme von Transport und Lagerung berücksichtigt worden?
30. Sind die logistischen Aspekte (Standardisierung der Werkzeuge und Prüfmittel, ersatzteilkonforme Strukturierung der Betrachtungseinheit, leichte Zugänglichkeit usw.) berücksichtigt worden?

g) Instandhaltbarkeit

1. Wurde ein Modularaufbau angestrebt? Sind die Module unabhängig? Ist das Gerät bzw. System in Ersatzteilen strukturiert?
2. Ist ein Konzept für die Ausfallerkennung und die Ausfall-Lokalisierung aufgestellt und realisiert worden? Werden verborgene Ausfälle ebenfalls erkannt und lokalisiert? Können redundante Elemente ohne Betriebsunterbrechung auf Geräte- bzw. Systemebene repariert werden?
3. Ist eine automatische Betriebsüberwachung vorgesehen? Welche Funktionen werden überwacht? Ist eine Zustandsprüfung vorhanden?

Tabelle 5.5 (Fortsetzung)

4. Sind genügend Prüfpunkte vorgesehen? Sind diese leicht zugänglich? Sind sie richtig bezeichnet? Sind sie nach Funktionen geordnet? Haben sie Pull-up/Pull-down Widerstände?
 5. Wurde der Abgleich auf ein Minimum beschränkt? Sind die entsprechenden Elemente leicht zugänglich und gut bezeichnet? Sind Wechselwirkungen vermieden worden? Ist die Justierung unkritisch?
 6. Ist die Anzahl der externen Prüfmittel auf ein Minimum beschränkt worden?
 7. Hat man einen größtmöglichen Standardisierungsgrad angestrebt?
 8. Sind die Ersatzteile definiert? Sind sie bezeichnet? Sind sie elektrisch und mechanisch leicht auftrennbar? Sind sie leicht auswechselbar?
 9. Sind die Elemente mit beschränkter Lebensdauer leicht zugänglich? Sind die Zugangstüren selbsthaltend? Lassen sie sich leicht und rasch öffnen bzw. lösen?
 10. Sind die Einschübe mit Gleitschienen und Arretierungen versehen? Sind Griffe zum Heben der Module vorhanden?
 11. Sind die Befestigungen leicht und mit standardisierten Werkzeugen lösbar? Sind sie auf ein Minimum reduziert worden?
 12. Sind die Stecker schnell lösbar? Ist ihre Bezeichnung klar? Sind Verwechslungen möglich? Sind Reservekontakte vorhanden?
 13. Sind Drähte und Kabel mit der richtigen Isolation versehen? Ist ihre Identifikation klar? Sind die Zuführungen ausreichend geschützt? Sind die Drähte nicht zu stark gespannt?
 14. Sind empfindliche Elemente gegen äußere und innere Störungen (elektrische, elektromagnetische, thermische, klimatische, mechanische) geschützt?
 15. Ist die Kühlung ausreichend und zweckmäßig?
 16. Ist die Wartung einfach? Kann sie während des Betriebs durchgeführt werden?
 17. Ist die Bedienungskonsole optimal ausgelegt (Bezeichnung, Beleuchtung, Ergonomie)? Sind die Abläufe sequentiell angeordnet?
 18. Sind die Aspekte der Unfallverhütung berücksichtigt worden (Schutzdeckel, Erdung, Sicherungen, Warnschilder usw.)?
 19. Ist die technische Sicherheit beachtet worden (FMEA/FNIECA, FTA)?
 20. Ist die Herstellbarkeit genügend berücksichtigt worden?
- h) Sicherheit
1. Sind die Vorschriften bezüglich Unfallverhütung beachtet worden?
 2. Ist die Sicherheit auch in Zusammenhang mit äußeren Einwirkungen (Naturkatastrophen, Sabotage usw.) untersucht worden?
 3. Ist eine FMEA/ FMECA durchgeführt worden? Sind "Single-Point Failures" aufgetreten? Lassen sich diese vermeiden?
 4. Sind "Fail-Safe"-Überlegungen und -Untersuchungen gemacht worden? Welches sind die Resultate?
 5. Welche Sicherheitsprüfungen sind vorgesehen? Sind sie ausreichend?
 6. Sind die Aspekte der Sicherheit in der Dokumentation genügend beachtet worden?

Tabelle 5.5 (Fortsetzung)

i) Menschliche Faktoren, Ergonomie

1. Hat man bei der Festlegung der Sequenzen für den Betrieb und für die Instandhaltung an das Ausbildungsni-
2. Ist bei der Auslegung der Bedienungs- und Signalisierungselemente den ergonomischen Aspekten genügend Beachtung geschenkt worden?
3. Hat man versucht, die Baugruppe optimal an den Menschen anzupassen?

j) Standardisierung

1. Sind soweit wie möglich standardisierte Bauteile und Stoffe verwendet worden?
2. Hat man während der Dimensionierung und der Konstruktion an die Austauschbarkeit der Bauteile gedacht?
3. Entsprechen die verwendeten Bezeichnungen in der Dokumentation und in der Hardware den vorgeschriebenen Richtlinien und Normen?

k) Konfiguration

1. Ist die technische Dokumentation vollständig und einwandfrei? Entspricht sie dem aktuellen Stand des Projekts?
2. Sind die Schnittstellenprobleme mit den anderen Baugruppen gelöst?
3. Kann man den aktuellen Stand der Dokumentation einfrieren und als Referenz-Dokumentation betrachten?
4. Kann man den aktuellen Stand der Konfiguration einfrieren und als Referenz-Konfiguration betrachten?

l) Fabrikation, Prüfung

1. Sind sämtliche Fragen der Herstellbarkeit, Prüfbarkeit und Reproduzierbarkeit beachtet worden?
2. Sind spezielle Fertigungsprozesse notwendig? Wurden sie untersucht, qualifiziert, erprobt? Welche Erfahrungen liegen vor?
3. Welche Qualifikationsprüfungen sind für die Prototypen vorgesehen? Werden in diese Prüfungen Zuverlässigkeits-, Instandhaltbarkeits- und Sicherheitsprüfungen einbezogen? Lassen sich diese Prüfungen in das Prüfkonzept für das ganze Gerät bzw. System integrieren?
4. Sind besondere Verpackungs-, Transport- oder Lagerungsprobleme zu erwarten?

Kritische Entwurfsüberprüfung auf Geräte-bzw. Systemebene

a) Formelles

1. Ist die technische Dokumentation vollständig?
2. Ist die technische Dokumentation auf die Richtigkeit überprüft worden? Welche Stellen waren daran beteiligt?
3. Sind die verschiedenen Dokumente der technischen Dokumentation auf Kohärenz überprüft worden? Durch wen?
4. Ist die Eindeutigkeit bei der Numerierung auch bezüglich Änderungsindex gewährleistet? Kann damit die Bauzustandsüberwachung sichergestellt werden?
5. Ist die Beschriftung der Hardware zweckdienlich? Genügt sie den Erfordernissen der Fabrikation und der Instandhaltung?
6. Ist die Übereinstimmung zwischen den Prototypen und der technischen Dokumentation überprüft worden?
7. Ist die Software ausreichend dokumentiert?

Tabelle 5.5 (Fortsetzung)

- b) Technisches
1. Ist die technische Dokumentation ausreichend, um eine eindeutige Interpretation der Prüfergebnisse und der Prüfvorschriften zuzulassen? Ist diese Dokumentation auf dem letzten Stand? Ist die Übereinstimmung mit der Hardware überprüft worden?
 2. Ist das Anforderungsprofil für die Zuverlässigkeitsprüfungen definiert worden?
 3. Ist das realisierte Instandhaltungskonzept ausreichend? Sind Ersatzteile mit verschiedenem Entwicklungsstand austauschbar?
 4. Sind die Kriterien für die Prüfung der Instandhaltbarkeit definiert worden? Welche Ausfälle wurden simuliert? Wie sind die personellen und materiellen Bedingungen für die Wartung und die Reparatur festgelegt worden? Wurden die Bedienungsfreundlichkeit und andere ergonomische Aspekte geprüft? Wie?
 7. Sind die Kriterien für die Prüfung der Sicherheit (Unfallverhütung und technische Sicherheit) definiert worden? Sind die Fehler- und Ausfallkriterien für die kritischen Parameter definiert worden? Ist für diejenigen deterministischen Parameter, die während der Prüfungen nicht genügend genau gemessen werden können, eine indirekte Messung vorgesehen worden?
 8. Sind die vorgesehenen Zwischen- und Endprüfungen in der Fertigung aus heutiger Sicht ausreichend? Lassen sich schon jetzt Probleme der Prüfbarkeit oder der Herstellbarkeit erkennen?
 9. Sind die Prüfvorschriften, die Prüfspezifikationen und die Prüfplanung vollständig? Sind die Bedingungen für die Funktionsprüfungen, die Umweltprüfungen, die Zuverlässigkeitsprüfungen, die Instandhaltbarkeitsprüfungen und die Sicherheitsprüfungen festgelegt worden?
 10. Sind die Verpackungs-, Transport- und Lagerungsprobleme untersucht worden?
 11. Sind Defekte und Ausfälle systematisch analysiert worden (Art, Ursache, Auswirkung)? Sind die entsprechenden Korrekturmaßnahmen überprüft worden? Wie? Durch wen? Auch kostenmäßig?
 12. Ist die Software ausreichend, klar und korrekt dokumentiert?
 13. Sind das Programm und die Daten robust gegen Fehlbedienung, äußere Störungen und Überlast?
 14. Ist die Software in ihrer endgültigen Hardware- und Software-Umgebung genügend erprobt worden?
 15. Ist die Liste der Abweichungen vollständig? Sind diese Abweichungen annehmbar? Erfüllt das Gerät bzw. das System die Kunden- bzw. Marktbedürfnisse noch? Kann die Vorserie freigegeben werden?
 4. Wie werden die Eingangsprüfungen durchgeführt? Mit welchen Stichprobenplänen? Mit welchen Prüfspezifikationen und Prüfvorschriften? Wie sieht das Freigabeverfahren dieser Dokumente aus?
 5. Welche Bauteile bzw. Stoffe werden hundertprozentig geprüft? Welche vorbehandelt? Wie lauten die Prozeduren?
 6. Wie werden die Prototypen qualifiziert? Wie sieht das Freigabeverfahren der entsprechenden Dokumentation aus? Wer entscheidet über die Prüfergebnisse?
 7. Wie werden die Zwischen- und Endprüfungen durchgeführt? Wie ist das Freigabeverfahren der entsprechenden Dokumentation? Wer beurteilt die Prüfergebnisse?
 8. Wie werden die Umweltprüfungen durchgeführt? Wie sieht das Freigabeverfahren der entsprechenden Dokumentation aus? Wer beurteilt die Prüfergebnisse?
- Ist eine Vorbehandlung für die Serieneinheiten vorgesehen?

10. Wie werden fehlerhafte Prüflinge behandelt?
- 1 1. Wie lauten die Anweisungen für Handhabung, Transport, Lagerung und Verpackung der geprüften Betrachtungseinheiten? Wie sieht das Freigabeverfahren der entsprechenden Dokumentation aus?

5.5.6 Qualitätsdatensystem

Vom Beginn der Qualifikation der Prototypen an (Qualifikationsphase) müssen alle Defekte und Ausfälle systematisch erfasst, analysiert und korrigiert werden. Dabei soll nach Möglichkeit, die Analyse bis hin zur Ursache geführt werden, um das Wiederauftreten derselben Probleme zu verhindern. Für die Beurteilung des Qualitäts- und Zuverlässigkeitssicherungsprogramms sind folgende Fragen wichtig:

- 1 . Wie erfolgt die Erfassung der Defekte und Ausfälle? Ab welcher Projektphase wirkt das Qualitätsdatensystem?
- 2 . Wie werden Defekte und Ausfälle analysiert? Welche Aufteilung wird vorgenommen? Wie lauten die Prozeduren?
- 3 . Wer führt die Korrekturmaßnahmen durch? Wer überwacht ihre Durchführung? Wer überprüft die Endkonfiguration?
- 4 . Wie erfolgt die Verdichtung, Auswertung und Rückkopplung der Informationen (Daten) zu den verschiedenen Stellen?
- 5 . Welche Stelle ist für das Qualitätsdatensystem als ganzes zuständig? Hat die Fertigung ein eigenes Qualitätsdatensystem mit beschränkter Reichweite? Wie sind die Schnittstellen definiert?

Anhang A1: Definitionen und Begriffserklärungen

In diesem Anhang werden die wichtigsten Begriffe auf dem Gebiet der Qualitäts- und Zuverlässigkeitssicherung technischer Systeme eingeführt und diskutiert. Die angegebenen Definitionen berücksichtigen so weit wie möglich die einschlägigen Normen. Auf noch nicht allgemein anerkannte Begriffe wie Bum-in, Einlaufen, MTBF, Verlässlichkeit und Vorbehandlung wird speziell eingegangen. Die prinzipielle hierarchische Einstufung der behandelten Begriffe ist im Bild A1.1 gegeben.

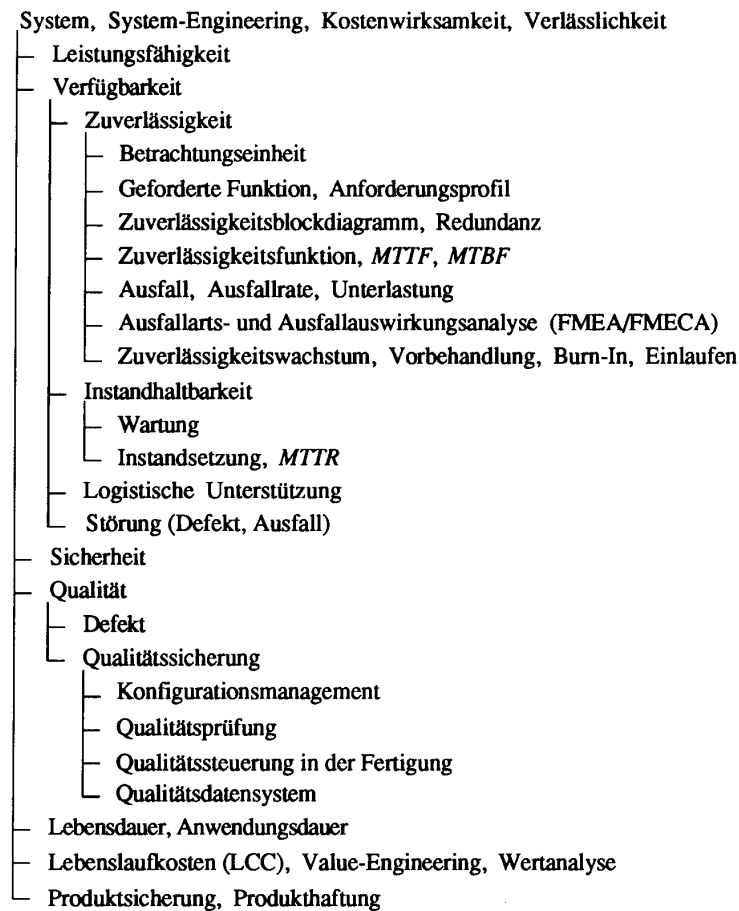


Bild A1.1 Prinzipielle Einstufung der wichtigsten Begriffe auf dem Gebiet der Qualitäts- und Zuverlässigkeitssicherung

Anforderungsprofil (Mission Profile, profile de la mission requise, profilo della missione richiesta)

Spezifische Aufgabe, die eine Betrachtungseinheit während einer bestimmten Zeit unter vorgegebenen Bedingungen ausführen muss.

Das Anforderungsprofil legt die geforderte Funktion und die Umweltbedingungen sowie deren Zeitverlauf fest. Ein repräsentatives Anforderungsprofil und die entsprechenden Zuverlässigkeitsziele sollen bereits im Pflichtenheft festgelegt werden.

Anwendungsdauer (Useful Life, *vie utile*, *vita utile*)

Zeitspanne der Anwendung, nach deren Ablauf eine Betrachtungseinheit definitiv aus dem Betrieb genommen wird.

Typische Werte für die Anwendungsdauer sind 3 bis 6 Jahre für normale Investitionsgüter, 5 bis 15 Jahre für Militäreinrichtungen und 10 bis 30 Jahre für Vermittlungs- und Energieanlagen. Anstelle von Anwendungsdauer wird oft auch von Brauchbarkeitsdauer gesprochen.

Ausfall (Failure, *défaillance*, *guasto*)

Beendigung der Fähigkeit einer Betrachtungseinheit, ihre geforderte Funktion auszuführen.

Ausfälle werden bezüglich AM Ursache und Auswirkung (ev. auch Ausfallmechanismus) eingeteilt. Bei der Bewertung der Auswirkungen muss berücksichtigt werden, ob man sich auf die direkt betroffene oder auf eine übergeordnete Betrachtungseinheit bezieht. Da der Ausfall nicht die einzige Ursache einer Betriebsunterbrechung ist, wird oft für Betriebsunterbrechungen (von der Wartung abgesehen) der Begriff Störung (Betriebsstörung) verwendet.

Ausfallarts- und -folgen-Analyse (Failure Modes and Effects Analysis, *analyse des modes de défaillance et de leur effets*, *analisi dei modi di guasto e delle relative conseguenze*)

Systematische Untersuchung aller möglichen Ausfallarten einer Betrachtungseinheit bezüglich ihrer Auswirkung auf die Funktionstüchtigkeit und die Sicherheit der betreffenden Betrachtungseinheit und den von dieser beeinflussten Betrachtungseinheiten.

Ziel einer FMEA/ FMECA (Failure Modes and Effects Analysis/Failure Modes, Effects and Criticality Analysis) ist das Auffinden aller potentiellen Gefahren (Hazards) und die Analyse der Vorkehrungen zur Beseitigung bzw. Milderung ihrer Auswirkung oder zur Reduktion ihrer Auftretswahrscheinlichkeit. Die verschiedenen Ausfallarten und Ausfallursachen werden systematisch untersucht, ausgehend von der tiefsten Integrationsebene (Bottom-up). Neben der FMEA/FMECA wird oft auch die Fault Tree Analysis (FTA) verwendet (Top-down). Die Prozedur der FMEA/FMECA lässt sich unmittelbar auf die Untersuchung der Auswirkung von Störungen erweitern, so dass die Abkürzung FMEA/ FMECA sinngemäss für Fault Modes and Effects Analysis/Fault Modes, Effects and Criticality Analysis verwendet werden kann.

A 1.2

Ausfallrate (Failure Rate, taux de défaillance, tasso di guasto)

Wahrscheinlichkeit bezogen auf δt , dass eine Betrachtungseinheit im Intervall $(t, t+\delta t]$ ausfallen wird, unter der Bedingung, dass sie zur Zeit $t=0$ eingeschaltet wurde und im Intervall $(0, t]$ nicht ausgefallen ist.

Die Ausfallrate wird mit $\lambda(t)$ bezeichnet. Wenn τ die ausfallfreie Arbeitszeit einer Betrachtungseinheit ist, so gilt

$$\lambda(t) = \lim_{\delta t \downarrow 0} \frac{1}{\delta t} \Pr\{t < \tau \leq t + \delta t \mid \tau > t\}.$$

Zwischen $\lambda(t)$ und der Zuverlässigkeitsfunktion $R(t)$ gilt (für $R(0)=1$) $R(t) = e^{-\int_0^t \lambda(x) dx}$. Für $\lambda(t) = \lambda$ folgt $R(t) = e^{-\lambda t}$. In diesem und nur in diesem Fall kann für die Schätzung von λ der Wert $\hat{\lambda} = k/T$ verwendet werden, wobei T die kumulative Betriebszeit (über beliebig viele, statistisch identische Betrachtungseinheiten) und k die totale Anzahl Ausfälle in T sind. Der typische Verlauf der Ausfallrate einer Gesamtheit statistisch identischer Betrachtungseinheiten setzt sich in der Regel aus der Phase der Frühausfälle, der Phase der Ausfälle mit konstanter (oder nahezu konstanter) Ausfallrate und der Phase der Verschleissausfälle zusammen. (Diese Phasen werden oft auch als Frühausfallphase, Phase der Zufallsausfälle und Spätausfallphase bezeichnet; der Begriff Zufallsausfall sollte allerdings vermieden werden, denn Ausfälle treten meist zufällig auf.)

Betrachtungseinheit (Item, unité, unità)

Beliebige Anordnung, wie z. B. Stoff, Bauteil, Unterbaugruppe, Baugruppe, Gerät, Anlage, System, welche für Untersuchungen oder Analysen als Einheit betrachtet wird.

Bei der Betrachtungseinheit handelt es sich in der Regel um eine Funktions- oder Konstruktionseinheit. Anstelle von Betrachtungseinheit wird oft der Begriff Einheit verwendet. Betrachtungseinheit wird auch als materieller oder immaterieller Gegenstand der Betrachtung definiert.

Burn-in

Betrieb einer Betrachtungseinheit unter erhöhten Beanspruchungen, um Ausfallmechanismen zu beschleunigen.

Für elektronische Betrachtungseinheiten sind die Beanspruchungen in der Regel eine hohe konstante Umgebungstemperatur (125°C für ICs, 85°C für Baugruppen) und eine erhöhte elektrische Belastung. Das Burn-in kann als Teil einer Vorbehandlungsequenz (hundertprozentig ausgeführt um Frühausfälle zu provozieren) oder als zeitraffende Zuverlässigkeitsprüfung (Untersuchung der Ausfallrate) betrachtet werden. Es darf nicht mit Einlaufen verwechselt werden. Die Beanspruchungen liegen höher, als sie in der Nutzungsphase zu erwarten sind, sie sollen aber nicht Ausfallmechanismen auslösen, die im normalen Betrieb nicht auftreten würden. Bei reparierbaren Betrachtungseinheiten werden die während des Burn-in auftretenden Ausfälle durch geeignete Korrekturmassnahmen behoben.

Defekt (Defect, défaut, difetto)

Abweichung einer Eigenschaft einer Betrachtungseinheit von den festgelegten Anforderungen.

A1.3

Defekte müssen die Funktionstüchtigkeit einer Betrachtungseinheit nicht unbedingt beeinträchtigen. Sie werden in der Regel durch Fehler in der Entwicklungs-, Fertigungs- oder Nutzungsphase verursacht.

Einlaufen (Run-in)

Betrieb einer Betrachtungseinheit unter normalen Beanspruchungen, um Fertigungsfehler zu erkennen und zu eliminieren.

Die Fehler, welche beim Einlaufen erkannt werden, sind oft regelmässig und damit in allen Betrachtungseinheiten eines gegebenen Loses vorhanden. In der Regel treten sie ohne Erhöhung der Beanspruchungen auf. Einlaufen wird deshalb oft nur an einer beschränkten Anzahl Betrachtungseinheiten vorgenommen (Vorserien). Einlaufen nicht mit Burn-in verwechselt werden.

Geforderte Funktion (Required Function, fonction requise, funzione richiesta)

Anforderungen an eine Betrachtungseinheit, gegeben in der Regel als Sollwerte und Toleranzen.

Die Festlegung der geforderten Funktion bildet den Ausgangspunkt jeder Zuverlässigkeitsanalyse, weil damit auch der Ausfall definiert wird. Dabei ist es aus praktischen Gesichtspunkten von Vorteil, wenn für alle Grössen Toleranzbereiche und nicht nur feste Werte vorgeschrieben werden.

Instandhaltbarkeit (Maintainability, maintenabilité, manutenibilità)

Mass für die Fähigkeit einer Betrachtungseinheit, funktionstüchtig gehalten werden zu können, ausgedrückt durch die Wahrscheinlichkeit, dass unter festgelegten materiellen und personellen Bedingungen der Zeitaufwand für eine Wartung bzw. für eine Instandsetzung kleiner als ein vorgegebenes Zeitintervall ist.

Es ist üblich, die Instandhaltbarkeit als eine (operationelle) Eigenschaft eines technischen Systems zu betrachten und in Wartbarkeit (Wartung) und Instandsetzbarkeit (Reparatur) einzuteilen. Bei der Festlegung der Instandhaltbarkeit müssen die entsprechenden personellen (Anzahl Leute, Ausbildung) und materiellen (Werkzeuge, Ersatzteile) Bedingungen angegeben werden. Eine Wartung erfolgt in der Regel planmässig und off-line, eine Reparatur hingegen zufällig und on-line.

Instandsetzung (Corrective Maintenance, maintenance corective, manutenzione correttiva)

Alle Aktivitäten zur Wiederherstellung des Sollzustands einer ausgefallenen Betrachtungseinheit.

A 1.4

Instandsetzung wird auch als Reparatur bezeichnet. Die Hauptschritte bei einer Reparatur sind: Ausfall-Lokalisierung, Ausfallbehebung, Abgleich und Funktionsprüfung. Für die Untersuchungen wird in der Regel angenommen, das ausgefallene Element (im Zuverlässigkeitsblockdiagramm) sei nach jeder Reparatur wieder neuwertig. Diese Annahme vereinfacht die Analysen. Sie gilt für die ganze Betrachtungseinheit (bis hin zur Systemebene), wenn jedes Element eine konstante Ausfallrate hat.

Konfigurationsmanagement (Configuration Management, gestion de la configuration, gestione della configurazione)

Verfahren zur Festlegung, Beschreibung, Prüfung und Genehmigung der Konfiguration einer Betrachtungseinheit sowie zu ihrer Steuerung und Überwachung bei Änderungen und Modifikationen.

Unter dem Begriff Konfiguration versteht man die Gesamtheit der funktionellen und physikalischen Eigenschaften einer Betrachtungseinheit, wie sie in der Dokumentation festgelegt und in der Hardware bzw. (Computer-) Software vorhanden ist. Das Konfigurationsmanagement wird eingeteilt in Beschreibung, Überprüfung, Steuerung und Überwachung der Konfiguration. Das Konfigurationsmanagement ist ein Teil der Qualitätssicherung.

Kostenwirksamkeit (Cost Effectiveness, efficacité des coûts, efficacia dei costi)

Mass für die Fähigkeit einer Betrachtungseinheit, die geforderte Funktion mit dem bestmöglichen Verhältnis von Nutzen zu Lebenslaufkosten zu erfüllen.

Anstelle von Kostenwirksamkeit wird oft von Systemwirksamkeit gesprochen.

Lebensdauer (Life Time, durée de vie, durata di vita)

Zeitintervall zwischen Beanspruchungsbeginn und Ausfallzeitpunkt einer nicht reparierbaren Betrachtungseinheit.

Lebenslaufkosten (Life Cycle Costs, coût global de durée de vie, costo globale sul ciclo di vita)

Summe der Kosten für die Anschaffung, den Betrieb, die Instandhaltung und die Ausscheidung einer Betrachtungseinheit.

Die Optimierung der Lebenslaufkosten wird im Rahmen der Kosten- bzw. Systemwirksamkeit und des SystemEngineerings vorgenommen.

Leistungsfähigkeit (Capability/Performance, capacité/performance, capacità/prestazione)

Mass für die Fähigkeit einer Betrachtungseinheit, den vorgegebenen funktionellen Anforderungen zu genügen.

AI.5

Alle Massnahmen und Aktivitäten, die ausgeführt werden, um eine wirksame und wirtschaftliche Verwendung einer Betrachtungseinheit während der Nutzungsphase zu ermöglichen.

Die logistische Unterstützung muss früh in der Entwicklungsphase konzipiert (Aufstellung des Instandhaltungskonzepts) und während der Entwicklung, Fertigung und Anwendung einer Betrachtungseinheit realisiert werden.

MTBF (Mean Time Between Failures, moyenne du temps de bon fonctionnement, tempo medio fra asti)

$$MTBF = 1 / \lambda .$$

MTBF ist nur für Betrachtungseinheiten mit konstanter Ausfallrate $\lambda(t)=\lambda$ zu verwenden. In diesem Falle gilt $R(t) = e^{-\lambda t}$ und $MTBF = 1 / \lambda$ ist gleich dem Mittelwert der ausfallfreien Arbeitszeit der Betrachtungseinheit (vgl. *MTTF*). *Anmerkung:* Die hier aufgeführte Definition steht im Einklang mit den verbreiteten Methoden zur statistischen Schätzung bzw. zum statistischen Nachweis einer *MTBF*, insbesondere der Punktschätzung $\hat{MTBF} = 1 / \hat{\lambda} = T / k$ mit T als kumulativer Betriebszeit und k als Anzahl Ausfälle während T (die Schätzung T / k darf nur in Zusammenhang mit einem Poisson-Prozess verwendet werden). Die aufgeführte Definition weicht von jener in verschiedenen Normen ab, wo *MTBF* für reparierbare und *MTTF* für nicht reparierbare Betrachtungseinheiten definiert werden. Mit der Verwendung der *MTBF* als Mittelwert der Zeit zwischen aufeinanderfolgenden Ausfällen einer reparierbaren Betrachtungseinheit stellt aber die *MTBF* die Summe der Mittelwerte der effektiven ausfallfreien Arbeits- und der Reparaturzeit dar, eine Auffassung, welche mit der Schätzung T / k nicht kompatibel ist.

MTTF (Mean Time To Failure, durée moyenne de fonctionnement avant défaillance, tempo medio fino al guasto)

Mittelwert der ausfallfreien Arbeitszeit einer Betrachtungseinheit.

Die *MTTF* wird aus der Zuverlässigkeitsfunktion $R(t)$ via $MTTF = \int_0^{\infty} R(t) dt$ berechnet. Was nach dem Ausfall mit der Betrachtungseinheit geschieht, ist für die *MTTF* nicht relevant. Ist sie reparierbar, so wird (damit die *MTTF* ihre Bedeutung nicht verliert) implizit angenommen, dass die Betrachtungseinheit nach der Reparatur neuwertig ist; der Mittelwert der nächsten ausfallfreien Arbeitszeit (ab Ende der Reparatur) ist dann gleich jenem des vorhergehenden (*MTTF*). Als Schätzung der *MTTF* kann $\hat{MTTF} = (t_1 + \dots + t_n) / n$ verwendet werden, wobei $t_1, \dots, t_n$ unabhängige Realisierungen (Beobachtungen) von ausfallfreien Arbeitszeiten statistisch identischer Betrachtungseinheiten sind.

MTTR (Mean Time To Repair, moyenne du temps de réparation, tempo medio di riparazione)

Mittelwert der Reparaturzeit einer Betrachtungseinheit.

Mit der Angabe der *MTTR* müssen auch die entsprechenden personellen (Anzahl Leute, Ausbildung) und materiellen (Werkzeuge, Ersatzteile) Bedingungen spezifiziert werden. Als Schätzung der *MTTR* kann $\hat{MTTR} = (t_1 + \dots + t_n) / n$ verwendet werden, wobei $t_1, \dots, t_n$ unabhängige Realisierungen (Beobachtungen) von Reparaturzeiten sta-

tistisch identischer Betrachtungseinheiten sind. Analytisch wird die *MTTR* aus der Verteilungsfunktion $G(t)$ der Reparaturzeiten via $MTTR = \int_0^{\infty} (1 - G(t)) dt$ berechnet.

Produkthaftung (Product Liability, responsabilité du fait du produit, responsabilità da prodotto)

Rechtliche Verantwortung des Herstellers für Personen-, Sach- oder Vermögensschäden, die durch den Gebrauch defekter oder ausgefallener Betrachtungseinheiten entstehen.

Dabei wird grundsätzlich eine sachgemäße, vom Hersteller vorgeschriebene Anwendung der Betrachtungseinheit vorausgesetzt. Es wird unterschieden zwischen verschuldensabhängiger Haftung (Tort Liability) und verschuldensunabhängiger Haftung (Kausalhaftung oder Strict Liability). Verschuldensunabhängige Haftung ist üblich in den USA und wird vermehrt auch in Europa verwendet.

Produktsicherung (Product Assurance, assurance produit, assicurazione prodotto)

Alle geplanten und systematischen Massnahmen, die ausgeführt werden, um das geforderte Qualitätsniveau sicherzustellen sowie die festgelegte Zuverlässigkeit, Instandhaltbarkeit, Verfügbarkeit und Sicherheit einer Betrachtungseinheit zu erreichen.

Im Rahmen der Produktsicherung versteht man unter Betrachtungseinheit die Gesamtheit der vom Hersteller gelieferten Leistungen, d. h. die Hardware, (Computer-)Software, Dokumentation und logistische Unterstützung. Produktsicherung wird in der Raumfahrttechnik und zunehmend auch im Militär- und Zivilbereich angewendet.

Qualität (Quality, qualité, qualità)

Gesamtheit der Eigenschaften und Merkmale eines Produkts oder einer Tätigkeit, die sich auf deren Eignung zur Erfüllung gegebener Erfordernisse beziehen.

Diese Definition ist sehr allgemein. Sie hat den Vorteil, dass sie alle (objektiven und subjektiven) Eigenschaften eines Produkts oder einer Tätigkeit berücksichtigen kann. Der Nachteil liegt im Verlust der Aussagekraft. Eine frühere, für technische Systeme möglicherweise besser geeignete Definition lautet: Mass für den Grad, zu welchem eine Betrachtungseinheit, den durch den Verwendungszweck gestellten funktionellen, operationellen und physikalischen Eigenschaften und Anforderungen, genügt.

Qualitätsdatensystem (Quality Data Reporting System, système des données de qualité, sistema dei dati di qualità)
System zur Erfassung, Analyse und Korrektur aller Defekte und Ausfälle, die während der Herstellung und Prüfung einer Betrachtungseinheit auftreten, sowie zur Verdichtung, Speicherung, Auswertung und Rückkopplung der entsprechenden Qualitäts- und Zuverlässigkeitsdaten.

Das Qualitätsdatensystem kann rechnerunterstützt sein. Die Analyse der Defekte und Ausfälle soll deren Ursachen aufzeichnen, damit geeignete Korrektur- und/oder Präventivmassnahmen zur Verhinderung einer Wiederholung der gleichen Probleme durchgeführt werden können. Wenn möglich, soll das Qualitätsdatensystem auch während der Nutzungsphase wirksam bleiben. Das Qualitätsdatensystem ist ein Teil der Qualitätssicherung.

Qualitätsprüfung (Quality Test, contrôle de la qualité, controllo qualità)

Prüfung, inwieweit eine Betrachtungseinheit den gestellten Anforderungen genügt.

Qualitätsprüfung umfasst Eingangsprüfungen, Qualifikationsprüfungen, Zwischenprüfungen und Endprüfungen (samt Zuverlässigkeits-, Instandhaltbarkeits- und Sicherheitsprüfungen). Aus Gründen der Kosten und Effizienz sollten alle Prüfungen in einem Prüf- bzw. Prüf- und Vorbehandlungskonzept integriert werden. Qualitätsprüfung ist ein Teil der Qualitätssicherung.

Qualitätssicherung (Quality Assurance, assurance de la qualité, assicurazione qualità)

Alle geplanten und systematischen Massnahmen, die ausgeführt werden, um das geforderte Qualitätsniveau einer Betrachtungseinheit sicherzustellen.

Zur Qualitätssicherung gehören das Konfigurationsmanagement, die Qualitätsprüfungen, die Qualitätssteuerung in der Fertigung und das Qualitätsdatensystem. Für komplexe Betrachtungseinheiten werden die Aktivitäten der Qualitätssicherung durch ein Qualitätssicherungsprogramm koordiniert und gesteuert. Ein wichtiges Ziel der Qualitätssicherung ist die Erreichung der geforderten Qualität mit minimalem Zeit- und Kostenaufwand (Integrale Qualitätssicherung oder Total Quality Management (TQM)).

Qualitätssteuerung in der Fertigung (Process Quality Control, contrôle de la qualité du processus, controllo della qualità del processo)

Steuerung der Fertigungsprozesse und -abläufe mit dem Ziel, das geforderte Qualitätsniveau einer Betrachtungseinheit sicherzustellen.

Qualitätssteuerung in der Fertigung ist ein Teil der Qualitätssicherung und wird auch Qualitätslenkung bezeichnet.

Redundanz (Redundancy, redondance, ridondanza)

Vorhandensein von mehr funktionsfähigen Mitteln in einer Betrachtungseinheit, als für die Erfüllung der geforderten Funktion notwendig sind.

Für die Hardware wird zwischen heisser (aktiver, paralleler), warmer (leicht belasteter) und kalter (Standby-) Redundanz unterschieden. Redundanz bedeutet nicht unbedingt eine Verfielfachung der Hardware, sie kann z. B. auch durch Kodierung, Software oder sequentiell realisiert werden. Erfüllen die redundanten Elemente (Reserve-Elemente) nur noch einen Teil der geforderten Funktion, so spricht man von einer Pseudoredundanz.

Sicherheit (Safety, sécurité, sicurezza)

Mass für die Fähigkeit einer Betrachtungseinheit, weder Menschen, Sachen noch Umwelt zu gefährden.

Es ist üblich, die Sicherheit als (operationelle) Eigenschaft eines technischen Systems zu betrachten und unter den Aspekten der Unfallverhütung (Sicherheit, wenn die Betrachtungseinheit korrekt funktioniert und betrieben wird) und der technischen Sicherheit (Sicherheit, wenn die Betrachtungseinheit oder ein Teil davon ausgefallen ist) zu untersuchen.

Störung (Fault, panne/dérangement, avaria)

Fehlende, fehlerhafte oder unvollständige Erfüllung der geforderten Funktion durch eine Betrachtungseinheit, von der Wartung abgesehen.

Eine Störung kann sich als Ausfall oder Defekt manifestieren (Symptom) und wird auch Versagen genannt. Die Ursachen von Ausfällen sind Ausfallmechanismen, jene der Defekte sind Fehler.

System (System, système, sistema)

Zusammenfassung technischer und organisatorischer Mittel zur autonomen Erfüllung eines Aufgabenkomplexes.

Ein System besteht im allgemeinen aus Hardware, (Computer-)Software, Menschen (Bedienungs- und Instandhaltungspersonal) und logistischer Unterstützung. Zur Vereinfachung der Berechnungen werden oft ideale Bedingungen für die menschlichen Aspekte und die logistische Unterstützung angenommen. In solchen Fällen spricht man in der Regel von technischen Systemen.

System-Engineering

Anwendung der Natur- und Ingenieur-Ressourcen zur Transformation eines operationellen Bedürfnisses in ein System, unter Berücksichtigung der Lebenslaufkosten sowie aller funktionellen und operationellen Eigenschaften.

Unterlastung (Derating, réduction de la charge, riduzione del carico)

Nichtausnützung der maximalen Belastbarkeit einer Betrachtungseinheit, um die Ausfallrate zu verringern.

Der Belastungsfaktor gibt das Verhältnis der tatsächlichen Belastung zur maximalen Belastbarkeit bei normalen Arbeitsbedingungen an (20 bis 40°C für die Umgebungstemperatur).

A 1.9

Anwendung der Methoden der Wertanalyse in der Entwicklungsphase zur präventiven Kostenbeeinflussung.

Verfügbarkeit/Punkt-Verfügbarkeit (Availability/Point Availability, disponibilité/disponibilità instantanée, disponibilità/disponibilità istantanea)

Mass für die Fähigkeit einer Betrachtungseinheit, zu einem gegebenen Zeitpunkt funktionstüchtig zu sein, ausgedrückt durch die Wahrscheinlichkeit, dass die Betrachtungseinheit zu einem gegebenen Zeitpunkt t die geforderte Funktion unter vorgegebenen Arbeitsbedingungen ausführt.

Es ist üblich, die Punkt-Verfügbarkeit als eine (operationelle) Eigenschaft eines technischen Systems zu betrachten und mit $PA(t)$ zu bezeichnen. Für die Untersuchungen wird in der Regel ein Dauerbetrieb angenommen (ständiges Wechseln vom Arbeits- zum Reparaturzustand und umgekehrt) und vorausgesetzt, dass die Betrachtungseinheit nach jeder Reparatur neuwertig ist. Falls die menschlichen Faktoren und die logistische Unterstützung nicht berücksichtigt werden, erhält man dann für den stationären Wert der Punkt-Verfügbarkeit den Ausdruck $PA(t) = PA = MTTF / (MTTF + MTTR)$. Je nach Anwendung können andere Verfügbarkeitsarten definiert werden (durchschnittliche Verfügbarkeit, Missions-Verfügbarkeit, Arbeitsmissions-Verfügbarkeit, welche logistische oder andere Faktoren berücksichtigen). Anstelle von Punkt-Verfügbarkeit kann auch der Begriff momentane Verfügbarkeit verwendet werden.

Verlässlichkeit (Dependability, sûreté de fonctionnement, fidatezza)

Mass für den Grad, in welchem eine Betrachtungseinheit ihre geforderte Funktion zu jedem Zeitpunkt während einer Mission mit gegebenem Anforderungsprofil erfüllt, unter der Annahme, zu Beginn der Mission funktionstüchtig zu sein.

Verlässlichkeit ist ein Mass für den Erfolg einer Mission (MIL-STD-721). Ihre Verwendung als Gesamtbegriff für Verfügbarkeit, Zuverlässigkeit, Instandhaltbarkeit und Sicherheit kann zu Verwirrungen führen. Alternativen für einen solchen Gesamtbegriff könnten operationelle Wirksamkeit, operationelle Verfügbarkeit (Bild 1.3) oder Funktionsfähigkeit sein.

Vorbehandlung (Screening/Environmental Stress Screening, déverminage, setacciatura)

Folge von Beanspruchungen, denen ein Los statistisch identischer Betrachtungseinheiten unterworfen wird, um Frühaussfälle zu provozieren.

Die Wirksamkeit eines bestimmten Vorbehandlungsschritts ist, je nach Integrationsebene, unterschiedlich. Für ICs ist z. B. das Burn-in wirksam, für Baugruppen sind dies die thermischen Zyklen und die Vibrationen. Die Erfahrung zeigt, dass die bezüglich Wirksamkeit und Kosten optimale Vorbehandlungssequenz fallweise zu bestimmen ist und flexibel (anpassungsfähig) gehandhabt werden soll. Für die Vorbehandlung auf Baugruppenebene wird oft der Begriff Environmental Stress Screening (ESS) verwendet.

Wartung (Preventive Maintenance, manutention préventive, manutenzione preventiva)

Alle Aktivitäten zur Erhaltung des Sollzustands einer funktionstüchtigen Betrachtungseinheit.

Ziel der Wartung ist die Kontrolle des Funktionszustands, die Entdeckung und Beseitigung von verborgenen Ausfällen und die Vermeidung von Drift- bzw. Verschleissausfällen. In den Untersuchungen wird oft angenommen, dass nach jeder Wartung die Betrachtungseinheit neuwertig ist. Diese Annahme vereinfacht die Analysen wesentlich. Sie trifft insbesondere dann zu, wenn alle Elemente der Betrachtungseinheit konstante Ausfallraten haben.

Wertanalyse (Value Analysis, analyse de la valeur, analisi del valore)

Optimierung der Konfiguration einer Betrachtungseinheit sowie der Herstellungsprozesse und -abläufe, damit die notwendigen Funktionen zu möglichst niedrigen Gesamtkosten erfüllt werden können, ohne die Leistungsfähigkeit, die Zuverlässigkeit, die Instandhaltbarkeit, die Sicherheit oder das Qualitätsniveau zu beeinträchtigen.

Zuverlässigkeit (Reliability, fiabilité, affidabilità)

Mass für die Fähigkeit einer Betrachtungseinheit, funktionstüchtig zu bleiben, ausgedrückt durch die Wahrscheinlichkeit, dass die geforderte Funktion unter vorgegebenen Arbeitsbedingungen während einer festgelegten Zeitdauer ausfallfrei ausgeführt wird.

Es ist üblich, die Zuverlässigkeit als eine (operationelle) Eigenschaft eines technischen Systems zu betrachten und mit R zu bezeichnen. Die Zuverlässigkeit gibt die Wahrscheinlichkeit an, dass in der festgelegten Zeitspanne (T) keine Betriebsunterbrechungen auftreten werden. Dies bedeutet aber nicht, dass redundante Teile nicht ausfallen dürfen; solche Teile können ausfallen und (ohne Betriebsunterbrechung auf Ebene Betrachtungseinheit) instandgesetzt werden. Der Begriff Zuverlässigkeit kann für nichtreparierbare wie auch für reparierbare Betrachtungseinheiten verwendet werden. Wird Tals Parameter (t) betrachtet, so ergibt sich die Zuverlässigkeitsfunktion $R(t)$.

Zuverlässigkeitsblockdiagramm (Reliability Block Diagram, diagramme de fiabilité, schema a blocchi dell'affidabilità)

Zusammenschaltung sämtlicher Elemente einer Betrachtungseinheit, die an der Erfüllung der geforderten Funktion beteiligt sind, in Form eines Flussdiagramms; darin erscheinen die für die Funktionserfüllung notwendigen Elemente in Serie und die Redundanten parallel.

Das Zuverlässigkeitsblockdiagramm ist ein Ereignisdiagramm. Es gibt Antwort auf die Frage: Welche Elemente müssen für die Erfüllung der geforderten Funktion funktionieren, und welche Elemente können ausfallen? Da es sich um ein Ereignisdiagramm handelt, dürfen für jedes Element nur eine Ausfallart (dominierende, z. B. Kurzschluss oder Unterbrechung) und nur zwei Zustände (gut/ausgefallen) angenommen werden.

AI.1 1

Zuverlässigkeit, ausgedrückt als Funktion der Zeit.

Die Zuverlässigkeitsfunktion wird mit $R(t)$ bezeichnet. $R(t)$ gibt die Wahrscheinlichkeit an, dass kein Ausfall im Intervall $(0, t]$ auftreten wird. In der Regel wird $R(0) = 1$ angenommen. In diesem Fall ist der Mittelwert der ausfallfreien Arbeitszeit gegeben durch $MTTF = \int_0^{\infty} R(t) dt$.

Zuverlässigkeitswachstum (Reliability Growth, accroissement de la fiabilité, incremento della affidabilità)

Erhöhung der Zuverlässigkeit einer Betrachtungseinheit durch erfolgreiche Behebung von Entwicklungs- oder Fertigungsmängeln.

Die im Laufe des Zuverlässigkeitswachstums erkannten Mängel (Schwachstellen) sind grösstenteils systematisch, d. h. bei allen Betrachtungseinheiten eines gegebenen Loses vorhanden. Zuverlässigkeitswachstum wird deshalb in der Regel bei der Qualifikation der Prototypen und der Serienreifmachung an einer beschränkten Anzahl Betrachtungseinheiten (selten hundertprozentig) durchgeführt. Im Gegensatz zum Einlaufen erfolgt das Zuverlässigkeitswachstum oft unter verschärften Umweltbedingungen (ähnlich wie bei der Vorbehandlung).

Anhang A2: Mathematische Grundlagen

Die Untersuchung der Zuverlässigkeit und Verfügbarkeit von Geräten und Systemen erfolgt mit den Werkzeugen der Wahrscheinlichkeitstheorie. In diesem Anhang werden die Grundlagen für solche Analysen kurz zusammengefasst, für eine ausführlichere Darlegung sowie für die Aspekte der mathematischen Statistik (im Zusammenhang mit Zuverlässigkeits- oder Verfügbarkeitsprüfungen) wird auf [2,1991] verwiesen.

A2.1 Auszug aus der Wahrscheinlichkeitsrechnung

A2.1.1 Der Begriff der Wahrscheinlichkeit

In der Wahrscheinlichkeitsrechnung werden Ausdrücke wie “das Gerät arbeitet ausfallfrei im Intervall $(0,t]$ ” oder “die Reparaturzeit ist kürzer als eine vorgegebene Zeitspanne t ” als *Ereignisse* betrachtet, für welche Wahrscheinlichkeiten definiert werden können (es wird impliziert bzw. stillschweigend angenommen, dass die Ereignisse zu einem Borelschen Ereignisfeld gehören). Die Wahrscheinlichkeit eines Ereignisses A wird mit

$$\Pr\{A\}$$

bezeichnet. Für sie gilt

$$0 \leq \Pr\{A\} \leq 1. \quad (\text{A2.1})$$

Dabei ist 0 die Wahrscheinlichkeit des unmöglichen Ereignisses Φ und 1 jene des sicheren Ereignisses Ω (nach dem Kolmogoroffschen axiomatischen Aufbau der Wahrscheinlichkeitsrechnung wird ferner postuliert, dass für unvereinbare Ereignisse $A_1, A_2, \dots$ gilt $\Pr\{A_1 \cup A_2 \cup \dots\} = \Pr\{A_1\} + \Pr\{A_2\} + \dots$). Die Wahrscheinlichkeit des komplementären Ereignisses $\bar{A}$ (nicht Eintreten von A) ist gegeben durch

$$\Pr\{\bar{A}\} = 1 - \Pr\{A\}. \quad (\text{A2.2})$$

Enthält das sichere Ereignis Ω nur endlich viele Elemente, so kann $\Pr\{A\}$ direkt mit der Formel der klassischen Wahrscheinlichkeit bestimmt werden

$$\Pr\{A\} = \frac{\text{Anzahl Elemente in } A}{\text{Anzahl Elemente in } \Omega} = \frac{\text{Anzahl günstige Fälle}}{\text{Anzahl mögliche Fälle}}. \quad (\text{A2.3})$$

Im allgemeinen Fall, wenn man eine unbeschränkte Folge von statistisch identischen Versuchen betrachtet, bei welchen ein bestimmtes Ereignis A entweder eintritt oder nicht eintritt, so zeigt die Erfahrung, dass die relative Häufigkeit k/n für das Eintreten von A mit grösser werdendem n immer weniger von einem festen Wert $p(A)$ abweicht. Es scheint daher sinnvoll für $n \rightarrow \infty$, den Grenzwert $p(A)$ als die Wahrscheinlichkeit $\Pr\{A\}$ mit $k/n = \hat{p}(A)$ als Schätzung zu bezeichnen.

A2.1.2 Bedingte Wahrscheinlichkeit, Unabhängigkeit

Der Begriff der bedingten Wahrscheinlichkeit ist für die praktische Anwendung der Wahrscheinlichkeitsrechnung von fundamentaler Bedeutung. Man kann sich leicht vorstellen, dass die Information "bei einem Versuch ist das Ereignis A eingetreten" zu einer Umbewertung der Wahrscheinlichkeiten anderer Ereignisse führen kann. Diese neuen Wahrscheinlichkeiten nennt man bedingte Wahrscheinlichkeiten und bezeichnet sie mit $\Pr\{B|A\}$. Ist z. B. A ein Element von B , so sollte vernünftigerweise $\Pr\{B|A\} = 1$, und somit in der Regel von der ursprünglichen unbedingten Wahrscheinlichkeit $\Pr\{B\}$ verschieden sein. Bei der Einführung des Begriffs der bedingten Wahrscheinlichkeit $\Pr\{B|A\}$, von B unter der Bedingung " A ist eingetreten", kann man sich wieder von Eigenschaften relativer Häufigkeiten leiten lassen. Dies führt zu folgender *Definition der bedingten Wahrscheinlichkeit*

$$\Pr\{B|A\} = \frac{\Pr\{A \cap B\}}{\Pr\{A\}}, \quad (\text{A2.4})$$

und damit zu

$$\Pr\{A \cap B\} = \Pr\{A\} \Pr\{B|A\} = \Pr\{B\} \Pr\{A|B\}. \quad (\text{A2.5})$$

Unter Beachtung der Gl. (A2.5) folgt damit, dass zwei Ereignisse A und B dann und nur dann unabhängig sind, wenn

$$\Pr\{A \cap B\} = \Pr\{A\} \Pr\{B\} \quad (\text{A2.6})$$

gilt. Die Ereignisse $A_1, \dots, A_n$ heißen (vollständig oder stochastisch) *unabhängig*, wenn für jedes k ($1 < k \leq n$) und jede Auswahl $i_1, \dots, i_k \in \{1, \dots, n\}$ gilt

$$\Pr\{A_{i_1} \cap \dots \cap A_{i_k}\} = \Pr\{A_{i_1}\} \dots \Pr\{A_{i_k}\}. \quad (\text{A2.7})$$

A2.1.3 Grundregeln der Wahrscheinlichkeitsrechnung

Die Berechnung der Wahrscheinlichkeit zusammengesetzter Ereignisse stützt sich auf die in diesem Abschnitt zusammengestellten Grundregeln.

A2.1.3.1 Additionssatz für zwei unvereinbare Ereignisse

Die Ereignisse A und B sind *unvereinbar*, wenn das Eintreten des einen Ereignisses das Eintreten des anderen *ausschliesst*. Wenn z. B. ein Halbleiterbauelement betrachtet wird, das entweder wegen Kurzschluss oder Unterbrechung ausfallen kann, so sind die Ereignisse {der Ausfall tritt wegen eines Kurzschlusses auf} und {der Ausfall tritt wegen einer Unterbrechung auf} unvereinbar. Für unvereinbare Ereignisse gilt

$$\Pr\{A \cup B\} = \Pr\{A\} + \Pr\{B\}. \quad (\text{A2.8})$$

Beispiel A2.1

Eine Lieferung enthält 95 gute Dioden, 3 Dioden mit Kurzschluss und 2 Dioden mit Unterbrechung. Es wird eine Diode herausgezogen. Gesucht ist die Wahrscheinlichkeit, dass diese defekt ist.

Lösung

Aus den Gln. (A2.8) und (A2.3) folgt

$$\Pr\{\text{Diode defekt}\} = \frac{3}{100} + \frac{2}{100} = \frac{5}{100}.$$

Falls die Ereignisse $A_1, A_2, \dots$ paarweise unvereinbar sind so sind sie auch vollständig unvereinbar und man hat (Axiom)

$$\Pr\{A_1 \cup A_2 \cup \dots\} = \sum_i \Pr\{A_i\}. \quad (\text{A2.9})$$

A2.1.3.2 Multiplikationssatz für zwei unabhängige Ereignisse

Zwei Ereignisse sind *unabhängig*, wenn die Information über das Eintreten (oder Nichteintreten) des einen Ereignisses keinen Einfluss auf die *Wahrscheinlichkeit* für das Eintreten des anderen Ereignisses hat. Für unabhängige Ereignisse gilt Gl. (A2.6)

$$\Pr\{A \cap B\} = \Pr\{A\} \Pr\{B\}.$$

Beispiel A2.2

Ein System besteht aus den zwei Elementen E_1 und E_2 , die für die Erfüllung der geforderten Funktion notwendig sind. Ein Ausfall von einem Element hat keinen Einfluss auf das andere. $R_1 = 0.8$ sei die Zuverlässigkeit von E_1 und $R_2 = 0.9$ jene von E_2 . Gesucht ist die Zuverlässigkeit R_S des Systems.

Lösung

Laut Definition gilt $R_1 = \Pr\{E_1 \text{ erfüllt die geforderte Funktion}\}$ und $R_2 = \Pr\{E_2 \text{ erfüllt die geforderte Funktion}\}$. Für das System hat man $R_S = \Pr\{E_1 \text{ erfüllt die geforderte Funktion} \cap E_2 \text{ erfüllt die geforderte Funktion}\}$, woraus mit Gl. (A2.6) folgt $R_S = R_1 R_2 = 0.72$.

A2.1.3.3 Multiplikationssatz für beliebige Ereignisse

Für beliebige Ereignisse A und B mit $\Pr\{A\} > 0$ gilt Gl. (A2.5)

$$\Pr\{A \cap B\} = \Pr\{A\} \Pr\{B \mid A\}.$$

Beispiel A2.3

In einer Sendung mit 95 guten und 5 defekten ICs werden 2 ICs herausgezogen. Gesucht ist die Wahrscheinlichkeit a) kein defektes IC und b) genau ein defektes IC zu bekommen.

Lösung

a) Aus den Gln. (A2.5) und (A2.3)) folgt

$$\Pr\{\text{erstes IC gut} \cap \text{zweites IC gut}\} = \frac{95}{100} \cdot \frac{94}{99} = 0.902.$$

b) Man hat $\Pr\{\text{genau ein IC defekt}\} = \Pr\{(\text{erstes IC gut} \cap \text{zweites IC defekt}) \cup (\text{erstes IC defekt} \cap \text{zweites IC gut})\}$; aus den Gln. (A2.2) und (A2.3) folgt dann

$$\Pr\{\text{ein IC defekt}\} = \frac{95}{100} \cdot \frac{5}{99} + \frac{5}{100} \cdot \frac{95}{99} = 0.096.$$

Die Verallgemeinerung von Gl. (A2.5) führt zum *Multiplikationssatz*

$$\Pr\{A_1 \cap \dots \cap A_n\} = \Pr\{A_1\} \Pr\{A_2 \mid A_1\} \Pr\{A_3 \mid (A_1 \cap A_2)\} \dots \Pr\{A_n \mid (A_1 \cap \dots \cap A_{n-1})\}. \quad (\text{A2.10})$$

Dabei wird $\Pr\{A_1 \cap \dots \cap A_{n-1}\} > 0$ vorausgesetzt. Ein wichtiger Spezialfall tritt auf, wenn die Ereignisse $A_1, \dots, A_n$ (vollständig) *unabhängig* sind, dann gilt

$$\Pr\{A_1 \cap \dots \cap A_n\} = \Pr\{A_1\} \dots \Pr\{A_n\} = \prod_{i=1}^n \Pr\{A_i\}. \quad (\text{A2.11})$$

A2.1.3.4 Additionssatz für zwei beliebige Ereignisse

Die Wahrscheinlichkeit für das Eintreten von *mindestens einem* der (beliebig verknüpften) Ereignisse A und B ist gegeben durch

$$\Pr\{A \cup B\} = \Pr\{A\} + \Pr\{B\} - \Pr\{A \cap B\}. \quad (\text{A2.12})$$

Beispiel A2.4

Zur Erhöhung der Zuverlässigkeit einer Anlage werden 2 Maschinen in Redundanz verwendet. Die Zuverlässigkeit jeder Maschine ist 0.9; die Maschinen werden voneinander unabhängig arbeiten und ausfallen. Gesucht ist die Zuverlässigkeit der Anlage.

Lösung

Aus den Gln. (A2.12) und (A2.6) folgt $\Pr\{\text{die erste Maschine erfüllt die geforderte Funktion} \cup \text{die zweite Maschine erfüllt die geforderte Funktion}\} = 0.9 + 0.9 - 0.9 \cdot 0.9 = 0.99$.

A2.1.3.5 Satz der totalen Wahrscheinlichkeit

Es seien $A_1, A_2, \dots$ paarweise unvereinbare Ereignisse ($A_i \cap A_j = \emptyset$ für alle $i \neq j$) und es gelte $\Omega = A_1 \cup A_2 \cup \dots$, sowie $\Pr\{A_i\} > 0$, $i = 1, 2, \dots$. Für ein beliebiges Ereignis B gilt der Satz der *totalen Wahrscheinlichkeit*

$$\Pr\{B\} = \sum_i \Pr\{B \cap A_i\} = \sum_i \Pr\{A_i\} \Pr\{B \mid A_i\}. \quad (\text{A2.13})$$

Beispiel A2.5

Es werden ICs von 3 Lieferanten A_1 , A_2 und A_3 in den Mengen 1'000, 600 und 400 gekauft. Die Wahrscheinlichkeit, ein defektes IC zu bekommen, sei 0.006 für A_1 , 0.02 für A_2 , und 0.03 für A_3 . Die ICs werden im Lager zusammengelegt. Gesucht ist die Wahrscheinlichkeit, dass ein aus dem Lager herausgegriffenes IC defekt ist.

Lösung

Aus den Gln. (A2.13) und (A2.3) folgt

$$\Pr\{\text{IC aus } A_1 \mid \text{IC defekt}\} = \frac{\frac{1'000}{2'000} \cdot 0.006}{0.015} = 0.2.$$

A2.1.4 Zufallsgrössen, Verteilungsfunktion

In den Anwendungen treten oft Grössen auf, die Zufallscharakter aufweisen; d.h. Grössen, die bei der Wiederholung des gleichen Versuchs verschiedene Werte (aus der Menge der möglichen Werte) mit bestimmten Wahrscheinlichkeiten annehmen können. Beispiele dafür sind die ausfallfreie Arbeitszeit einer Betrachtungseinheit, die Dauer einer Reparatur, usw. Solche Grössen werden Zufallsgrössen genannt und hier mit griechischen Buchstaben τ , ξ , ζ , usw. bezeichnet. Eine Zufallsgrösse ξ wird in der Regel durch die Verteilungsfunktion

$$F(x) = \Pr\{\xi \leq x\} \quad (\text{A2.14})$$

charakterisiert. Der Wert $F(x)$ gibt die Wahrscheinlichkeit an, mit der die Zufallsgrösse ξ einen Wert kleiner oder gleich x annimmt. $F(x)$ ist eine rechtseitig stetige, nicht fallende Funktion mit $F(-\infty) = 0$ und $F(+\infty) = 1$. Die Wahrscheinlichkeit, dass τ einen Wert aus dem Intervall $(a, b]$ annimmt ist gegeben durch

$$\Pr\{a < \tau \leq b\} = F(b) - F(a). \quad (\text{A2.15})$$

In vielen Anwendungen ist $F(x)$ differenzierbar. Ihre Ableitung

$$f(x) = \frac{dF(x)}{dx} \quad (\text{A2.16})$$

wird *Verteilungsdichte* von x genannt. Für $f(x)$ gilt

$$F(x) = \int_{-\infty}^x f(y) dy \quad \text{und} \quad \int_{-\infty}^{\infty} f(x) dx = 1. \quad (\text{A2.17})$$

In der Zuverlässigkeitstheorie bezeichnet man die ausfallfreie Arbeitszeit einer Betrachtungseinheit oft mit τ . τ ist eine positive Zufallsgrösse und es gilt somit $F(0) = 0$. Die *Zuverlässigkeitsfunktion* $R(t)$ gibt die Wahrscheinlichkeit an, dass die Betrachtungseinheit im Intervall $(0, t]$ ausfallfrei arbeitet, d.h. man hat

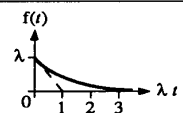
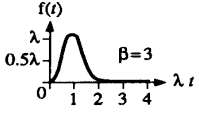
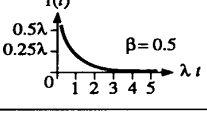
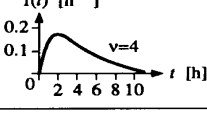
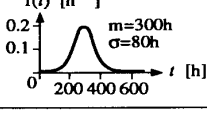
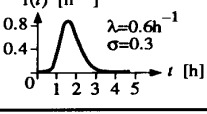
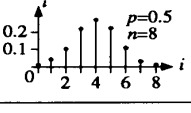
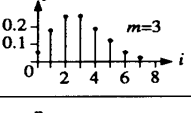
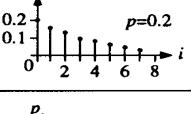
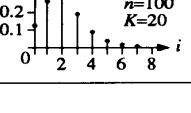
$$R(t) = \Pr\{t > \tau\} \quad (\text{A2.18})$$

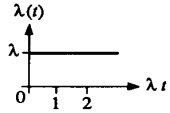
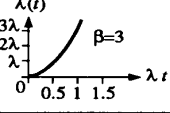
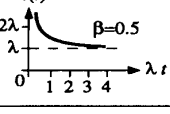
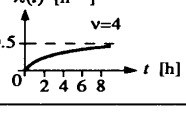
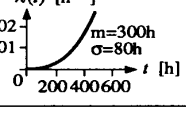
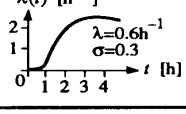
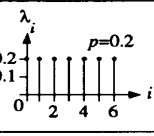
Wenn $F(t)$ die Verteilungsfunktion von t ist, dann gilt nach Gl. (A2.2)

$$R(t) = 1 - F(t) \quad (\text{A2.19})$$

Tab. A2.1 gibt die in der Zuverlässigkeitstheorie auftretenden Verteilungsfunktionen an.

Tabelle A2.1 Wichtige Verteilungsfunktionen in der Zuverlässigkeitstheorie [2, 1991]

Name	Verteilungsfunktion $F(t) = \Pr\{\tau \leq t\}$	Verteilungsdichte $f(t) = dF(t) / dt$	Wertebereich
Exponential	$1 - e^{-\lambda t}$		$t \geq 0$ $\lambda > 0$
Weibull	$1 - e^{-(\lambda t)^\beta}$		$t \geq 0$ $\lambda, \beta > 0$
Gamma	$\frac{1}{\Gamma(\beta)} \int_0^{\lambda t} x^{\beta-1} e^{-x} dx$		$t \geq 0$ $\lambda, \beta > 0$
Chi-Quadrat (χ^2)	$\frac{\int_0^t x^{(\nu/2)-1} e^{-x/2} dx}{2^{\nu/2} \Gamma(\nu/2)}$		$t \geq 0$ $\nu = 1, 2, \dots$ (Freiheitsgrade)
Normal	$\frac{1}{\sigma \sqrt{2\pi}} \int_{-\infty}^t e^{-\frac{(x-m)^2}{2\sigma^2}} dx$		$-\infty < t, m < \infty$ $\sigma > 0$
Logarithmische Normal- verteilung (Lognormal)	$\frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\frac{\ln(\lambda t)}{\sigma}} e^{-x^2/2} dx$		$t \geq 0$ $\lambda, \sigma > 0$
Binomial	$\Pr\{\zeta \leq k\} = \sum_{i=0}^k p_i$ $p_i = \binom{n}{i} p^i (1-p)^{n-i}$		$k = 0, \dots, n$ $0 < p < 1$
Poisson	$\Pr\{\zeta \leq k\} = \sum_{i=0}^k p_i$ $p_i = \frac{m^i}{i!} e^{-m}$		$k = 0, 1, \dots$ $m > 0$
Geometrisch	$\Pr\{\zeta \leq k\} = \sum_{i=1}^k p_i = 1 - (1-p)^k$ $p_i = p(1-p)^{i-1}$		$k = 1, 2, \dots$ $0 < p < 1$
Hyper- geometrisch	$\Pr\{\zeta = i\} = p_i = \frac{\binom{K}{i} \binom{N-K}{n-i}}{\binom{N}{n}}$		$i = 0, 1, \dots$ $\dots, \min(K, n)$

Ausfallrate $\lambda(t) = f(t) / (1 - F(t))$	$E[\tau]$	$\text{Var}[\tau]$	Eigenschaften
	$\frac{1}{\lambda}$	$\frac{1}{\lambda^2}$	Gedächtnislos $\Pr\{\tau > t + x_0 \mid \tau > x_0\} = \Pr\{\tau > t\} = e^{-\lambda t}$
	$\frac{\Gamma(1 + \frac{1}{\beta})}{\lambda}$	$\frac{\Gamma(1 + \frac{2}{\beta}) - \Gamma^2(1 + \frac{1}{\beta})}{\lambda^2}$	Monotone Ausfallrate (wachsend für $\beta > 1$ ($\lambda(0) = 0, \lambda(\infty) = \infty$), fallend für $\beta < 1$ ($\lambda(0) = \infty, \lambda(\infty) = 0$))
	$\frac{\beta}{\lambda}$	$\frac{\beta}{\lambda^2}$	$\tilde{f}(s) = \lambda^\beta / (s + \lambda)^\beta$; monotone Ausfallrate ($\lambda(\infty) = \lambda$), Erlang-Verteilung für $\beta = n = 2, 3, \dots$
	v	$2v$	Gamma-Verteilung mit $\beta = v/2$ und $\lambda = 1/2$ $F(t) = 1 - \sum_{i=0}^{v/2-1} \frac{(t/2)^i}{i!} e^{-t/2}$
	m	σ^2	$F(t) = \Phi(\frac{t-m}{\sigma})$ $\Phi(t) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^t e^{-x^2/2} dx$
	$\frac{e\sigma^2/2}{\lambda}$	$\frac{e^{2\sigma^2} - e^{\sigma^2}}{\lambda^2}$	$\ln \tau$ ist normalverteilt
nicht von Bedeutung	np	$np(1-p)$	$p_k = \Pr\{\zeta = k\} = \Pr\{k \text{ Erfolge in } n \text{ Bernoullischen Versuchen}\}$ (unabh. Versuche mit $\Pr\{A\} = p = \text{konst.}$)
nicht von Bedeutung	m	m	$\frac{(np)^i}{i!} e^{-np} \approx \binom{n}{i} p^i (1-p)^{n-i}$ $p_i = \Pr\{i \text{ Ausfälle in } (0, t] \mid \text{exponential verteilte ausfallfrei Arbeitszeit mit Parameter } \lambda\}, m = \lambda t$
	$\frac{1}{p}$	$\frac{1-p}{p^2}$	Gedächtnislos $p_i = \Pr\{\zeta = i\} = \Pr\{\text{erster Erfolg erst beim } i\text{-ten Bernoullischen Versuch}\}$
nicht von Bedeutung	$n \frac{K}{N}$	$\frac{K n (N-K) (N-n)}{N^2 (N-1)}$	Stichprobe ohne Zurücklegung

Die Eigenschaften der Exponential- und der Weibull-Verteilung sollen im folgenden mit Hilfe des Begriffs der Ausfallrate kurz besprochen werden. Die *Ausfallrate* $\lambda(t)$, eingeführt im Abschnitt 1.2, kann bei bekannter Verteilungsfunktion der ausfallfreien Arbeitszeit τ folgendermassen definiert werden

$$\lambda(t) = \lim_{\delta t \downarrow 0} \frac{1}{\delta t} \Pr\{t < \tau \leq t + \delta t \mid \tau > t\}. \quad (\text{A2.20})$$

Damit ist folgendes gemeint: die Betrachtungseinheit (mit ausfallfreier Arbeitszeit τ) sei zur Zeit $t = 0$ in Betrieb genommen worden und sei zur Zeit t noch nicht ausgefallen, $\lambda(t)\delta t$ ist dann gleich der Wahrscheinlichkeit, dass sie im nächsten Intervall δt ausfallen wird (Bild A2.1 soll diese Überlegung verdeutlichen).



Bild A2.1 Zur Definition der Ausfallrate

Aus Gl. (A2.20) folgt

$$\lambda(t) = \frac{f(t)}{1 - F(t)} = -\frac{dR(t)/dt}{R(t)}. \quad (\text{A2.21})$$

Die *Exponential-Verteilung* wird folgendermassen definiert

$$F(t) = 1 - e^{-\lambda t}, \quad t \geq 0, \quad \lambda > 0. \quad (\text{A2.22})$$

Für sie gilt

$$\lambda(t) = \lambda. \quad (\text{A2.23})$$

Die Ausfallrate ist im Falle der Exponentialverteilung konstant (zeitunabhängig). Diese wichtige Eigenschaft *charakterisiert* die Exponentialverteilung. Sie erleichtert die Berechnungen wesentlich, denn im Falle einer konstanten Ausfallrate gilt (Gedächtnislosigkeit):

Ist bekannt, dass die Betrachtungseinheit im jetzigen Augenblick funktioniert, so wird ihr weiteres Zeitverhalten nicht davon abhängen, wie lange sie schon gearbeitet hat, insbesondere ist die Wahrscheinlichkeit, dass im nächsten Intervall δt ausfällt, konstant und gleich $\lambda \delta t$.

Beispiel A2.6

Die ausfallfreie Arbeitszeit τ einer Baugruppe sei exponentiell verteilt mit $\lambda = 10^{-5} \text{ h}^{-1}$. Mit welcher Wahrscheinlichkeit liegt τ a) über 2'000 h, b) über 20'000 h, c) über 100'000 h, d) zwischen 20'000 h und 100'000 h?

Lösung

Aus Gln. (A2.22), (A2.18) und (A2.15) folgt

- a) $\Pr\{\tau > 2'000 \text{ h}\} = e^{-0.02} \approx 0.98$
- b) $\Pr\{\tau > 20'000 \text{ h}\} = e^{-0.2} \approx 0.819$
- c) $\Pr\{\tau > 100'000 \text{ h}\} = e^{-1} \approx 0.368$

$$d) \Pr\{20'000 \text{ h} < \tau \leq 100'000 \text{ h}\} = e^{-0.2} - e^{-1} \approx 0.451.$$

Eine Verallgemeinerung der Exponential-Verteilung führt zur *Weibull-Verteilung*

$$F(t) = 1 - e^{-(\lambda t)^\beta}, \quad t \geq 0, \quad \lambda, \beta > 0. \quad (\text{A2.24})$$

In diesem Fall folgt für die Ausfallrate

$$\lambda(t) = \beta \lambda (\lambda t)^{\beta-1}. \quad (\text{A2.25})$$

Für $\beta = 1$ hat man die Exponential-Verteilung. Für $\beta > 1$ (bzw. für $\beta < 1$) ist die momentane Ausfallrate monoton wachsend (bzw. monoton fallend). Die Weibull-Verteilung tritt in den Anwendungen oft auf, vorallem mit $\beta > 1$ als Verteilung der ausfallfreien Arbeitszeit von Bauteilen, die Verschleiss und/oder Ermüdung unterworfen sind, wie z.B. Röhren, Relais, mechanische Bauteile usw.

Besonders wichtig in der Zuverlässigkeitstheorie ist die Verteilung der Summe von zwei unabhängigen, positiven Zufallsgrössen. Es seien τ_1 und τ_2 zwei unabhängige Zufallsgrössen mit Verteilungsfunktionen $F_1(t)$ und $F_2(t)$, bzw. Verteilungsdichten $f_1(t)$ und $f_2(t)$, mit $F_1(0) = F_2(0) = 0$, gesucht ist die Verteilung der Summe von $\eta = \tau_1 + \tau_2$. Mit Hilfe des Satzes der totalen Wahrscheinlichkeit kann man zeigen, dass

$$F_\eta(t) = \Pr\{\eta \leq t\} = \int_0^t f_1(x) F_2(t-x) dx \quad (\text{A2.26})$$

und damit

$$f_\eta(t) = \int_0^t f_1(x) f_2(t-x) dx. \quad (\text{A2.27})$$

A2.1.5 Momente

Für eine grobe Charakterisierung einer Zufallsgrösse genügt in vielen Anwendungen die Angabe von verschiedenen Zahlengrössen. Zu diesen gehören insbesondere der Erwartungswert und die Varianz (vgl. Tab. A2.1).

A2.1.5.1 Erwartungswert (Mittelwert)

Der Erwartungswert $E[\tau]$ einer Zufallsgrösse τ mit Dichte $f(t)$ wird aus

$$E[\tau] = \int_{-\infty}^{\infty} t f(t) dt \quad (\text{A2.28})$$

berechnet, was für positive stetige Zufallsgrössen auf

$$E[\tau] = \int_0^{\infty} t f(t) dt \quad (\text{A2.29})$$

reduziert wird. Für positive Zufallsgrössen gilt auch

$$E[\tau] = \int_0^{\infty} (1 - F(t)) dt = \int_0^{\infty} R(t) dt. \quad (\text{A2.30})$$

Beispiel A2.7

Die ausfallfreie Arbeitszeit einer Betrachtungseinheit sei exponentiell verteilt mit Parameter λ . Gesucht ist der Erwartungswert (Mittelwert)

Lösung

Aus den Gln. (A2.22) und (A2.30) folgt

$$E[\tau] = \int_0^{\infty} e^{-\lambda t} dt = \frac{1}{\lambda}.$$

Es ist üblich $1/\lambda$ mit MTBF zu bezeichnen (Gl. 1.8). Man merke, dass für $t = 1/\lambda = \text{MTBF}$, $R(\text{MRBF}) = e^{-1} \approx 0.37$ gilt. Dies bedeutet, dass im Falle der Exponential-Verteilung nur rund 40% der Betrachtungseinheiten eine ausfallfreie Arbeitszeit grösser als MTBF haben werden.

In der mathematischen Statistik wird als (statistische) Schätzung für den Erwartungswert der arithmetische Mittelwert der unabhängigen Beobachtungen (Realisierungen) $t_1, \dots, t_n$ der Zufallsgrösse τ

$$\hat{E}[\tau] = \frac{1}{n} \sum_{i=1}^n t_i \quad (\text{A2.31})$$

definiert. $\hat{E}[\tau]$ ist für grosse Werte von n normalverteilt mit Erwartungswert $E[t]$ und Varianz $\text{Var}[\tau]/n$. Diese Eigenschaft von $\hat{E}[\tau]$ erklärt die Dualität der Begriffe Erwartungswert und Mittelwert.

A2.5.2 Varianz

Die Varianz einer Zufallsgrösse τ ist ein Mass dafür, wie stark die Zufallsgrösse um ihren Mittelwert $E[\tau]$ streut. Sie wird als

$$\text{Var}[\tau] = E[(\tau - E[\tau])^2] = E[\tau^2] - (E[\tau])^2. \quad (\text{A2.32})$$

definiert. Für eine stetige Zufallsgrösse gilt

$$\text{Var}[\tau] = \int_{-\infty}^{\infty} (t - E[\tau])^2 f(t) dt. \quad (\text{A2.33})$$

Die Grösse

$$\sigma = \sqrt{\text{Var}[\tau]} \quad (\text{A2.34})$$

wird als *Streuung* oder *Standardabweichung* und die Grösse

$$\kappa = \frac{\sigma}{E[\tau]} \quad (\text{A2.35})$$

als *Variationskoeffizient* von τ bezeichnet. Die Zufallsgrösse $(\tau - E[\tau]) / \sigma$, hat den Mittelwert 0 und die Varianz 1 (Normierung der Zufallsgrösse τ).

A2.2 Auszug aus der Theorie der stochastischen Prozesse

A2.2.1 Einführung

Stochastische Prozesse sind mathematische Modelle von Zufallserscheinungen, die in der Zeit ablaufen, wie z. B. das Zeitverhalten eines reparierbaren Geräts oder Systems. Sie werden hier mit $\xi(t)$ bezeichnet. Betrachtet man das Zeitverhalten eines Systems, das zufälligen Einflüssen unterliegt, und es sei $\mathcal{T}$ der uns interessierende Zeitbereich, z. B. $\mathcal{T} = [0, \infty)$, und $Z_0, \dots, Z_m$ der Zustandsraum, d. h. die Menge der möglichen Zustände des Systems; dann ist der Zustand des Systems zu einer gegebenen Zeit t_0 eine gewöhnliche Zufallsgrösse $\xi(t_0)$. Die Zufallsgrössen $\xi(t)$, $t \in \mathcal{T}$, sind so miteinander gekoppelt, dass für $n = 1, 2, \dots$ und beliebige Werte $t_1, \dots, t_n \in \mathcal{T}$ des Zeitparameters eine n -dimensionale Verteilungsfunktion existiert, die als die Verteilungsfunktion des Zufallsvektors $\xi(t_1), \dots, \xi(t_n)$ gilt

$$F(x_1, \dots, x_n; t_1, \dots, t_n) = \Pr\{\xi(t_1) \leq x_1, \dots, \xi(t_n) \leq x_n\}. \quad (\text{A2.36})$$

Für viele praktische Anwendungen genügen oft einige spezifische Kenngrössen des zugrundeliegenden stochastischen Prozesses (dessen Existenz man voraussetzt), wie bestimmte Zustandswahrscheinlichkeiten oder Verweilzeiten, zu deren Ermittlung nicht die Kenntnis aller Verteilungsfunktionen (A2.36) erforderlich ist. Aus der Problemstellung und den getroffenen Modellannahmen erkennt man oft direkt

- die Gestalt des *Zeitbereichs* $\mathcal{T}$ (stetig oder diskret, endlich oder unendlich)
- die Gestalt des *Zustandsraumes* (stetig oder diskret)
- den *Nachwirkungsgrad* (Abhängigkeitsstruktur) des zu untersuchenden stochastischen Prozesses
- *Invarianzeigenschaften* des Prozesses in Bezug auf Zeitverschiebungen (stationäre bzw. zeit-homogene Prozesse).

In der Zuverlässigkeitstheorie treten praktisch nur Prozesse mit stetigem Zeitparameter und diskretem Zustandsraum auf. Als Zeitbereich wird in der Regel die positive Zeitachse ($t \geq 0$) genommen. Die Zustände werden mit $Z_0, \dots, Z_m$ bezeichnet. Bezüglich *Nachwirkungsgrad* sind folgende Prozess-Klassen von Bedeutung

- Erneuerungsprozesse
- regenerative Prozesse mit bis zu einem einzigen regenerativen Zustand
- Markoff Prozesse
- Semi-Markoff Prozesse

- Semi-regenerative Prozesse (Prozesse mit einem eingebetteten Semi-Markoff-Prozess).

Hinsichtlich der Nachwirkungsfreiheit stellen die Markoff-Prozesse eine direkte Verallgemeinerung der Folgen unabhängiger Zufallsgrößen dar. Grob gesagt ist $\xi(t)$ ein *Markoff-Prozess*, wenn sein Verlauf nach einem beliebigen Zeitpunkt t durch t und seinem Zustand in t , nicht aber von seinem Verlauf vor t , abhängig ist. Der Markoff-Prozess ist damit nachwirkungsfrei. Oft ist in einem Markoff-Prozess auch die Abhängigkeit von t nicht vorhanden (zeithomogene Markoff-Prozesse); in diesem Fall ist der Markoff-Prozess bis zur Kenntnis des Zustands zum Zeitpunkt der Betrachtung gedächtnislos. *Regenerative Prozesse* haben die Eigenschaft, dass es gewisse, in der Regel zufällige, Zeitpunkte gibt (Eintrittszeitpunkte bestimmter Zustände), in denen der Prozess seine Vergangenheit vergisst. Solche Punkte werden als *Erneuerungspunkte* (Regenerationspunkte) bezeichnet. Zwischen den Regenerationspunkten kann die Abhängigkeitsstruktur kompliziert sein. *Semi-Markoff-Prozesse* sind regenerative Prozesse, bei welchen *alle Zustände* regenerativ sind.

Zur Beschreibung des Zeitverhaltens von Systemen, die sich im statistischen Gleichgewicht (stationärem Zustand) befinden, eignen sich stationäre und zeithomogene Prozesse. Der Prozess $\xi(t)$ wird *stationär* (im engeren Sinne) genannt, falls für $n = 1, 2, \dots$, für beliebige Zeiten $t_1, \dots, t_n$ und a ($t_i, t_i + a \in T$)

$$F(x_1, \dots, x_n; t_1 + a, \dots, t_n + a) = F(x_1, \dots, x_n; t_1, \dots, t_n) \quad (\text{A2.37})$$

gilt. Für $n = 1$ besagt Gl. (A2.37), dass die Verteilungsfunktion der Zufallsgröße $\xi(t)$ unabhängig von t ist. Damit sind alle Momente, insbesondere $E[\xi(t)]$ und $\text{Var}[\xi(t)]$ *zeitunabhängig*. Ferner ist für $n = 2$ die Verteilungsfunktion der zweidimensionalen Zufallsgröße $(\xi(t), \xi(t + u))$ nur eine Funktion von u . Daraus folgt, dass auch der Korrelationskoeffizient zwischen $\xi(t)$ und $\xi(t + u)$ nur eine Funktion von u ist (sind nur der Erwartungswert, die Varianz und der Korrelationskoeffizient zeitunabhängig, so wird der Prozess als stationär im weiteren Sinne bezeichnet). Ein Prozess $\xi(t)$ wird *zeithomogen* genannt (Prozess mit stationären Zuwächsen), falls für $n = 1, 2, \dots$, beliebige Zeitintervalle (b_i, t_i) , beliebigen Zeitabschnitt a ($t_i, t_i + a, b_i, b_i + a \in T$) und beliebige Werte $x_1, \dots, x_n$ gilt

$$\begin{aligned} \Pr\{\xi(t_1 + a) - \xi(b_1 + a) \leq x_1, \dots, \xi(t_n + a) - \xi(b_n + a) \leq x_n\} \\ = \Pr\{\xi(t_1) - \xi(b_1) \leq x_1, \dots, \xi(t_n) - \xi(b_n) \leq x_n\}. \end{aligned} \quad (\text{A2.38})$$

Ist der Prozess $\xi(t)$ stationär, dann ist er auch zeithomogen (homogen).

Im folgenden wird man sich hier auf die Einführung der Markoff-Prozesse beschränken.

A2.2.2 Markoff-Prozesse mit endlich vielen Zuständen

A2.2.2.1 Definition

Ein stochastischer Prozess $\xi(t)$ mit endlich vielen Zuständen $Z_1, \dots, Z_m$ ist ein Markoff-Prozess, falls für $n = 1, 2, \dots$, für beliebige Zeitpunkte $t + a > t > t_n > \dots > t_1$ und beliebige $i, j, i_1, \dots, i_n \in \{0, \dots, m\}$ die Gleichung

$$\begin{aligned} \Pr\{\xi(t+a) = Z_j \mid (\xi(t) = Z_i \cap \xi(t_n) = Z_{i_n} \cap \dots \cap \xi(t_1) = Z_{i_1})\} \\ = \Pr\{\xi(t+a) = Z_j \mid \xi(t) = Z_i\} \end{aligned} \quad (\text{A2.39})$$

gilt. Die bedingten Zustandswahrscheinlichkeiten in Gl. (A2.39) werden *Übergangswahrscheinlichkeiten* genannt und mit $P_{ij}(t, t+a)$ bezeichnet

$$P_{ij}(t, t+a) = \Pr\{\xi(t+a) = Z_j \mid \xi(t) = Z_i\}. \quad (\text{A2.40})$$

Der Markoff-Prozess ist *zeithomogen*, wenn die Übergangswahrscheinlichkeiten $P_{ij}(t, t+a)$ zeitunabhängig sind

$$P_{ij}(t, t+a) = P_{ij}(a). \quad (\text{A2.41})$$

Im folgenden werden nur noch zeithomogene Markoff-Prozesse betrachtet. Die Übergangswahrscheinlichkeiten $P_{ij}(a)$ genügen den Bedingungen

$$P_{ij}(a) \geq 0 \quad \text{und} \quad \sum_{j=0}^m P_{ij}(a) = 1, \quad i = 0, \dots, m. \quad (\text{A2.42})$$

Die Elemente $P_{ij}(a)$ bilden damit eine *stochastische Matrix*. Zusammen mit der *Anfangsverteilung*

$$P_i(0) = \Pr\{\xi(0) = Z_i\}, \quad i = 0, \dots, m, \quad (\text{A2.43})$$

bestimmen sie das Verteilungsgesetz des Markoff-Prozesses, denn für $t > 0$ sind die *Zustandswahrscheinlichkeiten*

$$P_j(t) = \Pr\{\xi(t) = Z_j\}, \quad j = 0, \dots, m \quad (\text{A2.44})$$

durch

$$P_j(t) = \sum_{i=0}^m P_i(0) P_{ij}(t) \quad (\text{A2.45})$$

gegeben. Setzt man für $i, j = 0, \dots, m$ $P_{ij}(0) = 0$ für $i \neq j$ und $P_{ii}(0) = 1$ und nimmt man an, dass die Übergangswahrscheinlichkeiten $P_{ij}(t)$ in $t=0$ stetig sind, so kann man zeigen, dass die $P_{ij}(t)$ in $t=0$ auch differenzierbar sind. Die Grenzwerte

$$\lim_{\delta t \downarrow 0} \frac{P_{ij}(\delta t)}{\delta t} = \rho_{ij} \quad \text{für } i \neq j \quad \text{und} \quad \lim_{\delta t \downarrow 0} \frac{1 - P_{ii}(\delta t)}{\delta t} = \rho_i \quad (\text{A2.46})$$

existieren, und es gilt

$$\rho_i = \sum_{j=0}^m \rho_{ij}, \quad i = 0, \dots, m, \quad \rho_{ii} = 0. \quad (\text{A2.47})$$

Beachtet man, dass $P_{ij}(\delta t) = \Pr\{\xi(t + \delta t) = Z_j \mid \xi(t) = Z_i\}$ für jedes $t \geq 0$ gilt, so erhält man folgende nützliche Deutung für die Grössen p_{ij} und p_i (für $\delta t \downarrow 0$ und t beliebig, insbesondere $t = 0$)

$$p_{ij} \delta t = \Pr\{\text{Übergang von } Z_i \text{ nach } Z_j \text{ in } (t, t + \delta t)\} \quad (\text{A2.48})$$

$$p_i \delta t = \Pr\{Z_i \text{ wird verlassen in } (t, t + \delta t)\}. \quad (\text{A2.49})$$

Man bezeichnet p_{ij} und p_i als *Übergangsraten*. Sie spielen bei der Analyse von Markoff-Prozessen eine ähnliche Rolle wie die Übergangswahrscheinlichkeiten p_{ij} für die Markoff-Ketten.

Markoff-Prozesse können ganz allgemein zur Untersuchung des Zeitverhaltens reparierbarer Systeme dann eingesetzt werden, wenn *alle* auftretenden Zufallsgrössen *unabhängig und exponentiell verteilt sind*. Bei ihrer Modellierung ist es für die praktischen Anwendungen nützlich, die in einem Intervall $(t, t + \delta t]$, mit t beliebig (z. B. $t = 0$) und δt sehr klein, möglichen Übergänge und die zugehörigen Übergangsraten p_{ij} und p_i in einem *Diagramm der Übergangswahrscheinlichkeiten* in $(t, t + \delta t]$ festzuhalten. Dieses Diagramm ist eine direkte Erweiterung des *Diagramms der Zustandsübergänge*. Es ist ein gerichteter Graph mit den Zuständen Z_i , $i = 0, \dots, m$, als Knoten und den Übergangswahrscheinlichkeiten $P_{ij}(\delta t)$ als Verbindungen, *dabei werden die Glieder der Ordnung $o(\delta t)$ vernachlässigt*. Diese Übergangswahrscheinlichkeiten in $(t, t + \delta t]$ lassen sich leicht aus den konstanten Ausfall- und Reparaturraten der Elemente im Zuverlässigkeitsblockdiagramm bestimmen. Bild A2.2 illustriert dies anhand einer Struktur bestehend aus der Serienschaltung einer heissen Redundanz 1 aus 2 mit identischen Elementen $E_2 = E_3 = E$ und eines Schaltelements E_1 ; das System verfügt über eine einzige Reparaturmannschaft und es wird auf E_1 repariert sobald es ausfällt (Priorität der Reparatur), unabhängig vom Zustand der Elemente E_2 und E_3 . Im Bild A2.2b ist der Fall betrachtet, bei welchem nach einem Ausfall auf Ebene System kein weiterer Ausfall mehr auftreten kann (bis das System nicht wieder funktionstüchtig ist), durch diese Annahme wird der Zustandsraum wesentlich verkleinert. Die Redundanz 1 aus 2 allein wird in den Beispielen A2.8 und A2.9 behandelt.

A2.2.2.2 Berechnung der Zustandswahrscheinlichkeiten und der Verweilzeiten in einer bestimmten Klasse von Zuständen

Für die Anwendungen in der Zuverlässigkeitstheorie spielen vor allem die Zustandswahrscheinlichkeiten und die Verteilung der Verweilzeit in einer bestimmten Klasse von Zuständen eine wichtige Rolle. Die Zustandswahrscheinlichkeiten erlauben die Berechnung der *Punkt-Verfügbarkeit*. Mit der Verteilungsfunktion der Verweilzeit in einer bestimmten Klasse kann die *Zuverlässigkeitsfunktion* angegeben werden. Eine Kombination dieser beiden Grössen ermöglicht ferner die Berechnung der *Intervall-Zuverlässigkeit*.

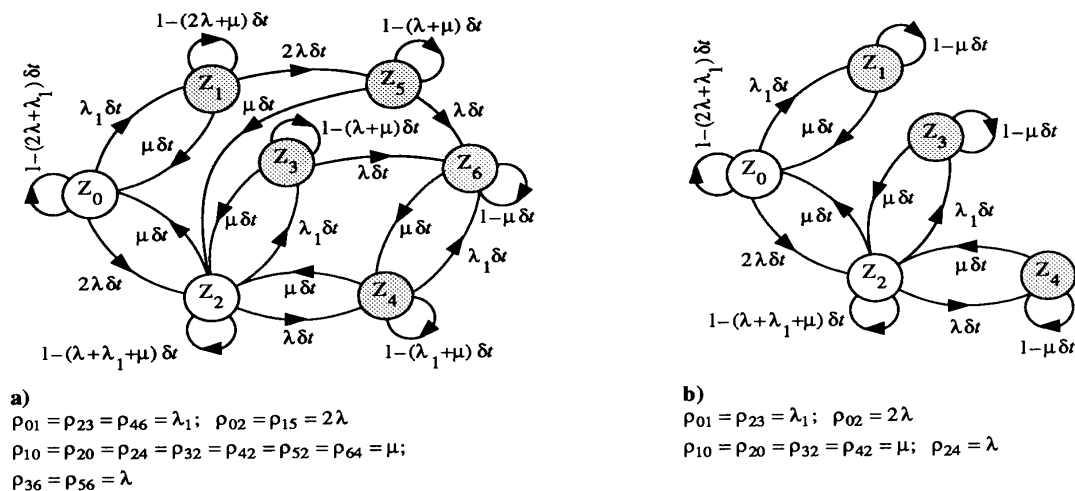
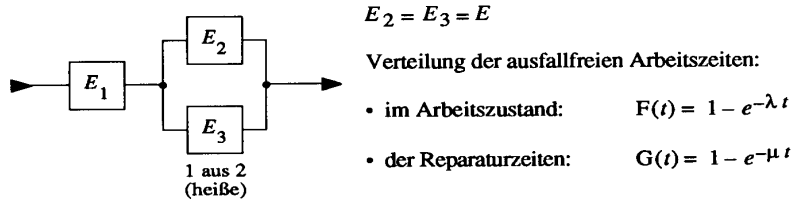


Bild A2.2 Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ einer reparierbaren Serien-/ Parallelstruktur mit $E_2 = E_3$ und Priorität der Reparatur auf E_1 (λ, λ_r = Ausfallrate, μ = Reparaturrate): **a)** die Betrachtungseinheit arbeitet bis zum Ausfall des letzten Elements; **b)** im Reparaturzustand auf Niveau Betrachtungseinheit kann kein weiterer Ausfall auftreten (t beliebig, z. B. $t = 0$, $\delta t \rightarrow 0$)

Es ist zweckmässig, für solche Betrachtungen die Zustände des Prozesses in zwei komplementären Teilmengen U und $\bar{U}$ aufzuteilen:

U = Menge der Zustände, in welchen das System (die Betrachtungseinheit) funktionstüchtig ist (*up - states*)

$\bar{U}$ = Menge der Zustände, in welchen das System (die Betrachtungseinheit) als ausgefallen gilt (*down - states*).

Die Berechnung der Zustandswahrscheinlichkeiten und der Verweilzeiten in der Menge U kann mit den Methoden der Differential- oder der Integralgleichungen erfolgen. Im folgenden wird man sich hier auf die Methode der Differentialgleichungen beschränken.

Die *Methode der Differentialgleichungen* ist die klassische Methode zur Untersuchung von Markoff-Prozessen. Sie basiert auf dem Diagramm der Übergangswahrscheinlichkeiten in $(t, t + \delta t]$. Es sei $\xi(t)$ ein zeithomogener Markoff-Prozess mit der beliebigen Anfangsverteilung $P_i(0) = \Pr\{\xi(0) = Z_i\}$ und den Übergangsraten ρ_{ij} und ρ_i . Die Zustandswahrscheinlichkeit gemäss Gl. (A2.44) erfüllt folgendes System von Differentialgleichungen [2.1991]

$$\dot{P}_j(t) = -\rho_j P_j(t) + \sum_{i=0}^m P_i(t) \rho_{ij}, \quad j = 0, \dots, m, \quad \rho_{ii} = 0. \quad (\text{A2.50})$$

Die *Punkt-Verfügbarkeit* $PA_S(t)$ ist für beliebige Anfangsbedingungen zur Zeit $t = 0$ gegeben durch

$$PA_S(t) = \Pr\{\xi(t) \in U\} = \sum_{Z_j \in U} P_j(t). \quad (\text{A2.51})$$

In der Regel ist man aber in den praktischen Anwendungen an den Resultaten für bestimmte Anfangsbedingungen zur Zeit $t = 0$ interessiert. Setzt man

$$P_i(0) = 1 \quad \text{und} \quad P_j(0) = 0 \quad \text{für} \quad j \neq i, \quad (\text{A2.52})$$

d. h. das System ist zur Zeit $t = 0$ in Z_i (in der Regel in Z_0 , wo alle Elemente funktionstüchtig (*up*) sind), so sind die Zustandswahrscheinlichkeiten $P_j(t)$ *identisch* mit den *Übergangswahrscheinlichkeiten* $P_{ij}(t)$ gemäss den Gln. (A2.40) und (A2.41)

$$P_{ij}(t) \equiv P_j(t), \quad (\text{A2.53})$$

mit $P_j(t)$ als Lösung der Gln. (A2.50) und (A2.52). Die *Punkt-Verfügbarkeit*, diesmal mit $PA_{Si}(t)$ bezeichnet, folgt dann aus

$$PA_{Si}(t) = \Pr\{\xi(t) \in U \mid \xi(0) = Z_i\} = \sum_{Z_j \in U} P_{ij}(t). \quad (\text{A2.54})$$

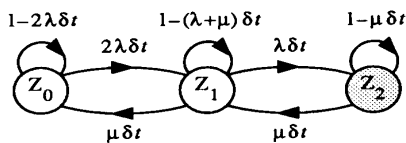
$PA_{Si}(t)$ ist die Wahrscheinlichkeit, das System zur Zeit t in Betrieb (in einem der Zustände der Menge U zu finden, wenn das System zur Zeit $t = 0$ in Z_i war. Beispiel A2.8 zeigt die Berechnung der Punkt-Verfügbarkeit im Falle einer Redundanz 1 aus 2.

Beispiel A2.8

Gegeben sei eine heisse Redundanz 1 aus 2, bestehend aus 2 identischen Elementen E_1 und E_2 mit konstanter Ausfallrate λ und Reparaturrate μ . Für das System gebe es nur eine Reparaturmannschaft. Man untersuche die Zustandswahrscheinlichkeiten des dazugehörigen stochastischen Prozesses (E_1 und E_2 sind neu zur Zeit $t = 0$).

Lösung

Das folgende Bild zeigt das entsprechende Diagramm der Übergangswahrscheinlichkeiten in $(t, t + \delta t]$ (t beliebig, z. B. $t = 0$ und $\delta t \rightarrow 0$)



Betrachtet man die zeitliche Entwicklung des Prozesses im Intervall $(t, t + \delta t]$, so erhält man für die Zustandswahrscheinlichkeiten (Gl. (A2.44))

$$P_0(t + \delta t) = P_0(t)(1 - 2\lambda\delta t) + P_1(t)\mu\delta t$$

$$P_1(t + \delta t) = P_1(t)(1 - (\lambda + \mu)\delta t) + P_0(t)2\lambda\delta t + P_2(t)\mu\delta t$$

$$P_2(t + \delta t) = P_2(t)(1 - \mu \delta t) + P_1(t) \lambda \delta t,$$

und für $\delta t \rightarrow 0$

$$\begin{aligned}\dot{P}_0(t) &= -2\lambda P_0(t) + P_1(t)\mu \\ \dot{P}_1(t) &= -(\lambda + \mu)P_1(t) + P_0(t)2\lambda + P_2(t)\mu \\ \dot{P}_2(t) &= -\mu P_2(t) + P_1(t)\lambda.\end{aligned}\tag{A2.55}$$

Die Lösung dieses Systems von Differentialgleichungen, mit den entsprechenden Anfangsbedingungen (z. B. $P_0(0) = 1, P_1(0) = P_2(0) = 0$), liefert die Zustandswahrscheinlichkeiten $P_0(t)$, $P_1(t)$ und $P_2(t)$, bzw. die Übergangswahrscheinlichkeiten $P_{00}(t)$, $P_{01}(t)$ und $P_{02}(t)$ des Prozesses und daraus die Punkt-Verfügbarkeit gemäss Gl. (A2.54).

Eine weitere wichtige Grösse für Zuverlässigkeitsuntersuchungen ist die *Zuverlässigkeitsfunktion* $R_S(t)$, d. h. die Wahrscheinlichkeit für keinen Systemausfall im Intervall $(0, t]$. Sie folgt unmittelbar aus der Verteilungsfunktion der Verweilzeit in der Menge U der *up-states* und kann mit der Methode der Differentialgleichungen berechnet werden, indem man alle Zustände der Menge $\bar{U}$ (*down-states*) *absorbierend* macht. Ist z. B. der Zustand $Z_k \in \bar{U}$ absorbierend, so wird der Prozess den Zustand Z_k nicht mehr verlassen wenn er einmal in Z_k gekommen ist. Es ist nicht schwer zu erkennen, dass in diesem Fall die Summe der Wahrscheinlichkeiten für die Zustände in U gleich der gesuchten Zuverlässigkeitsfunktion ist. Um diese Betrachtungen allgemein darzustellen, sei der modifizierte Markoff-Prozess $\xi'(t)$ mit Übergangswahrscheinlichkeiten $P'_{ij}(t)$ und Übergangsraten

$$\rho'_{ij} = \rho_{ij} \quad \text{für } Z_i \in U, \quad \rho'_{ij} = 0 \quad \text{für } Z_i \in \bar{U}, \quad \rho'_{ii} = 0, \quad \rho'_i = \sum_{j=0}^m \rho'_{ij}\tag{A2.56}$$

angenommen. Die Zustandswahrscheinlichkeiten $P'_j(t)$ von $\xi'(t)$ erfüllen folgendes System von Differentialgleichungen

$$\dot{P}'_j(t) = -\rho'_j P'_j(t) + \sum_{i=0}^m P'_i(t) \rho'_{ij}, \quad j = 0, \dots, m, \quad \rho'_{ii} = 0, \quad \rho'_{ij} = 0 \quad \text{für } Z_i = \bar{U}.\tag{A2.57}$$

Mit der Anfangsbedingung $P'_i(0) = 1$ und $P'_j(0) = 0$ für $j \neq i, Z_i \in U$, folgen aus Gl. (A2.57) die Zustandswahrscheinlichkeiten $P'_j(t)$ und dann die Übergangswahrscheinlichkeiten

$$P'_{ij}(t) \equiv P'_j(t).\tag{A2.58}$$

Für die *Zuverlässigkeitsfunktion*, hier mit $R_{Si}(t)$ bezeichnet, gilt damit

$$R_{Si}(t) = \Pr\{\xi(u) \in U \text{ für } 0 < u \leq t \mid \xi(0) = Z_i\} = \sum_{Z_j \in U} P'_{ij}(t), \quad Z_i \in U.\tag{A2.59}$$

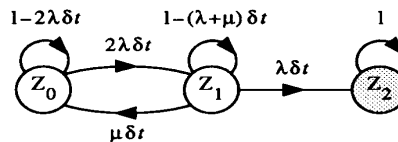
Beispiel A2.9 zeigt die Berechnung der Zuverlässigkeitsfunktion für eine reparierbare heisse Redundanz 1 aus 2.

Beispiel A2.9

Gegeben sei das gleiche System wie in Beispiel A2.8. Gesucht ist die Zuverlässigkeitsfunktion, d. h. die Wahrscheinlichkeit, dass der Prozess die Zustände Z_0 und Z_1 bis zum Zeitpunkt t nicht verlässt.

Lösung

Das Diagramm der Übergangswahrscheinlichkeiten in $(t, t + \delta t]$ modifiziert sich folgendermassen:



Damit folgt für die Zustandswahrscheinlichkeiten (vgl. Beispiel A2.8)

$$\begin{aligned}
 \dot{P}_0(t) &= -2\lambda P_0(t) + P_1(t)\mu \\
 \dot{P}_1(t) &= -(\lambda + \mu)P_1(t) + P_0(t)2\lambda \\
 \dot{P}_2(t) &= P_1(t)\lambda.
 \end{aligned} \tag{A2.60}$$

Die Lösung dieses Systems von Differentialgleichungen, mit den entsprechenden Anfangsbedingungen (z. B. $P_0(0) = 1$ und $P_1(0) = P_2(0) = 0$) liefert die Zustandswahrscheinlichkeiten $P_0(t)$, $P_1(t)$ und $P_2(t)$. Die Summe $P_0(t) + P_1(t)$ ist gleich der Wahrscheinlichkeit, dass sich zur Zeit t der Prozess im Zustand Z_0 oder Z_1 befindet, ohne zuvor den Zustand Z_2 angenommen zu haben. (Der Strich bei den Zustandswahrscheinlichkeiten soll eine Verwechslung mit den Werten aus dem Gleichungssystem (A2.55) verhindern).

Die Gln. (A2.54) und (A2.59) erlauben die Berechnung der Wahrscheinlichkeit, dass der Prozess zur Zeit t in einem Zustand der Menge U ist und in dieser Menge während des Intervalls $(t, t + \theta]$ bleibt, gegeben $\xi(0) = Z_i$. Diese Wahrscheinlichkeit wird als *Intervall-Zuverlässigkeit* definiert und mit $IR_{Si}(t, t + \theta)$ bezeichnet. Wegen der *Nachwirkungsfreiheit* des Markoff-Prozesses gilt

$$IR_{Si}(t, t + \theta) = \sum_{Z_j \in U} P_{ij}(t) R_{Sj}(\theta), \quad i = 0, \dots, m \tag{A2.61}$$

A2.2.2.3 Stationäres und asymptotisches Verhalten

Die Berechnung der zeitabhängigen Zustandswahrscheinlichkeiten und der Punkt-Verfügbarkeit eines Systems mit konstanten Ausfall- und Reparaturraten ist mit Hilfe von Differential- bzw. Integralgleichungen prinzipiell immer möglich, sie kann jedoch sehr aufwendig werden. Sind die Zustandswahrscheinlichkeiten unabhängig von der Zeit, d. h. ist der betrachtete Prozess stationär, so hat man nur ein *lineares Gleichungssystem* zu lösen. Ein zeithomogener Markoff-Prozess $\xi(t)$ mit den Zuständen $Z_0, \dots, Z_m$ ist genau dann stationär, wenn seine Zustandswahrscheinlichkeiten $P_i(t) = \Pr\{\xi(t) = Z_i\}$, $i = 0, \dots, m$, nicht von t abhängen. Aus $P_i(t+a) = P_i(t)$ folgt auch $P_i(t) = P_i(0) = p_i$ und insbesondere $\dot{P}_i(t) = 0$. Folglich ist unter Berücksichtigung der Gl. (A2.50) ein zeithomogener Markoff-Prozess $\xi(t)$ genau dann *stationär*, wenn seine Anfangsverteilung $p_i = P_i(0) = \Pr\{\xi(0) = Z_i\}$, $i = 0, \dots, m$ für alle $t > 0$ die Gleichungen

$$\rho_j p_j = \sum_{i=0}^m p_i \rho_{ij}, \quad j = 0, \dots, m \quad \text{mit} \quad p_j \geq 0, \quad \sum_{j=0}^m p_j = 1 \quad \text{und} \quad \rho_{ii} = 0. \quad (\text{A2.62})$$

erfüllt. Jede Lösung der Gl. (A2.62) mit $p_i \geq 0$, $i = 0, \dots, m$ heisst eine *stationäre Anfangsverteilung* des betrachteten Markoff-Prozesses.

Ein Markoff-Prozess heisst *irreduzibel*, wenn zu jedem Paar $i, j \in \{0, \dots, m\}$ ein t existiert, so dass $P_{ij}(t) > 0$ ist (d. h. jeder Zustand kann aus jedem anderen Zustand erreicht werden). Der Markoff-Prozess $\xi(t)$ ist *genau dann irreduzibel*, wenn seine *eingebettete Markoff-Kette* irreduzibel ist. Für einen irreduziblen Markoff-Prozess existieren Zahlen $p_j > 0$, $j = 0, \dots, m$, mit der Eigenschaft $p_0 + \dots + p_m = 1$ so, dass unabhängig von der Anfangsverteilung $P_i(0)$ gilt (Satz von Markoff)

$$\lim_{t \rightarrow \infty} P_j(t) = p_j, \quad j = 0, \dots, m. \quad (\text{A2.63})$$

Aus der Gl. (A2.63) folgt dann unmittelbar für jedes $i = 0, \dots, m$

$$\lim_{t \rightarrow \infty} P_{ij}(t) = p_j, \quad j = 0, \dots, m. \quad (\text{A2.64})$$

Man bezeichnet $p_0, \dots, p_m$ aus Gl. (A2.64) als die *Grenzverteilung* des Markoff-Prozesses. (Aus Gl. (A2.64) erkennt man, dass die Grenzverteilung die einzige stationäre Anfangsverteilung ist, d. h. die einzige Lösung der Gl. (A2.62).)

Für die Punkt-Verfügbarkeit kann damit der stationäre und asymptotische Wert PA_S

$$\lim_{t \rightarrow \infty} PA_{Si}(t) = PA_S = \sum_{Z_j \in U} p_j \quad (\text{A2.65})$$

in den meisten Fällen leicht berechnet werden. Definiert man die *durchschnittliche Verfügbarkeit* $AA_S(t)$ als

$$AA_S = \frac{1}{t} E[\text{gesamte Arbeitszeit in } (0, t)] = \frac{1}{t} \int_0^t PA_S(x) dx, \quad (\text{A2.66})$$

so folgt gemäss den Gln. (A2.54) und (A2.64)

$$AA_S = \lim_{t \rightarrow \infty} AA_S(t) = PA_S = \sum_{Z_j \in U} p_j. \quad (\text{A2.67})$$

Ausdrücke der Form $\sum_k p_k$ können verwendet werden, um die erwartete Anzahl der Elemente in Reparatur, der besetzten Reparaturmannschaften, der Elemente im Reservezustand usw. zu bestimmen.

A2.2.2.4 Wichtige Beziehungen für Markoff-Modelle

Für ein gegebenes Diagramm der Übergangswahrscheinlichkeiten in $(t, t+\delta t]$ gibt die Tabelle A2.2 die allgemeinen Formeln zur Berechnung der Zuverlässigkeitsfunktion $R_{Si}(t)$, des Mittelwertes der ausfallfreien Arbeitszeit $MTTF_{Si}$, der Punkt-Verfügbarkeit $PA_{Si}(t)$, des stationären und asymptotischen Wert PA_S und der Intervall-Zuverlässigkeit $IR_{Si}(t, t+\theta)$. Die Indizes S und i

beziehen sich dabei auf das System (S) und auf die Ausgangslage zur Zeit $t = 0$ (Zustand Z_1 , Z_0 falls das System zur Zeit $t = 0$ neu ist).

Tabelle A2.2 Wichtige Beziehungen für Markoff-Modelle [2, 1991]

	Zuverlässigkeitsfunktion	Punkt-Verfügbarkeit	Intervall-Zuverlässigkeit
Zeithomogene Markoff-Prozesse (Methode der Differentialgleichungen)	$R_{Si}(t) = \sum_{Z_j \in U} P'_{ij}(t), \quad Z_i \in U$	$PA_{Si}(t) = \sum_{Z_j \in U} P_{ij}(t), \quad i = 0, \dots, m$	$IR_{Si}(t, t + \theta) = \sum_{Z_j \in U} P_{ij}(t) R_{Sj}(\theta),$ $i = 0, \dots, m$
	$MTTF_{Si} = \frac{1}{\rho_i} + \sum_{Z_j \in U} \frac{\rho_{ij}}{\rho_i} MTTF_{Sj}, \quad Z_i \in U$ $P'_{ij}(t) \equiv P'_j(t), \quad \text{mit } P'_j(t) \text{ aus}$ $\dot{P}'_j(t) = -\rho_j P'_j(t) + \sum_{Z_i \in U} P'_i(t) \rho_{ij},$ $j = 0, \dots, m, \quad \rho_{ii} = 0$ mit $P'_i(0) = 1$ und $P'_j(0) = 0$ für $j \neq i$	$PA_S = \sum_{Z_j \in U} p_j$ $P_{ij}(t) \equiv P_j(t), \quad \text{mit } P_j(t) \text{ aus}$ $\dot{P}_j(t) = -\rho_j P_j(t) + \sum_{i=0}^m P_i(t) \rho_{ij}$ $j = 0, \dots, m, \quad \rho_{ii} = 0$ mit $P_i(0) = 1$ und $P_j(0) = 0$ für $j \neq i$	p_j aus $p_j \rho_j = \sum_{i=0}^m p_i \rho_{ij},$ mit $\rho_{ii} = 0, \quad p_j > 0$ und $p_0 + \dots + p_m = 1$
$R_{Si}(t)$	$= \Pr\{\text{im Arbeitszustand in } (0, t] Z_i \text{ ist bei } t = 0 \text{ eingetreten}\}, \quad Z_i \in U$		
$MTTF_{Si}$	$= E[\text{ausfallfreie Arbeitszeit} Z_i \text{ ist bei } t = 0 \text{ eingetreten}] = \int_0^\infty R_{Si}(t) dt = \tilde{R}_{Si}(0), \quad Z_i \in U$		
$PA_{Si}(t)$	$= \Pr\{\text{im Arbeitszustand zur Zeit } t Z_i \text{ ist bei } t = 0 \text{ eingetreten}\}, \quad i = 0, \dots, m$		
PA_S	$= \Pr\{\text{im Arbeitszustand zur Zeit } t \text{ im asymptotischen } (t \rightarrow \infty) \text{ und im stationären Zustand}\} = \lim_{s \rightarrow 0} s \tilde{PA}_{Si}(s), \quad i = 0, \dots, m$		
AA_S	$= \lim_{t \rightarrow \infty} \frac{1}{t} E[\text{totale Betriebszeit in } (0, t)] = PA_S$		
$IR_{Si}(t, t + \theta)$	$= \Pr\{\text{im Arbeitszustand in } (t, t + \theta] Z_i \text{ ist bei } t = 0 \text{ eingetreten}\}, \quad i = 0, \dots, m$		
$IR_S(\theta)$	$= \Pr\{\text{im Arbeitszustand in } (t, t + \theta] \text{ im asymptotischen } (t \rightarrow \infty) \text{ und im stationären Zustand}\}$		
U	$=$ Menge der Zustände, in welchem das System funktionstüchtig ist (up states)		
$\bar{U}$	$=$ Menge der Zustände, in welchem das System als ausgefallen gilt (down states)		
$P_{ij}(t)$	$= \Pr\{\text{in } Z_j \text{ zur Zeit } t Z_i \text{ ist bei } t = 0 \text{ eingetreten}\}$		
ρ_{ij}	$= \lim_{\delta t \rightarrow 0} \frac{1}{\delta t} \Pr\{\text{Übergang von } Z_i \text{ nach } Z_j \text{ in } (t, t + \delta t] Z_i \text{ zur Zeit } t\}, \quad t \text{ beliebig, z.B. } t = 0 \text{ (gilt nur für Markoff-Prozesse)}$		
S	bezieht sich auf System; $\tilde{u}(s) = \int_0^\infty \tilde{u}(t) e^{-st} dt =$ Laplace-Transformation von $u(t)$		

beziehen sich dabei auf das System (S) und auf die Ausgangslage zur Zeit $t = 0$ (Zustand Z_i , Z_0 falls das System zur Zeit $t = 0$ neu ist).

Tabelle A2.2 Wichtige Beziehungen für Markoff-Modelle [2, 1991]

	Zuverlässigkeitsfunktion	Punkt-Verfügbarkeit	Intervall-Zuverlässigkeit
Zeithomogene Markoff-Prozesse (Methode der Differentialgleichungen)	$R_{Si}(t) = \sum_{Z_j \in U} P_{ij}^*(t), \quad Z_i \in U$	$PA_{Si}(t) = \sum_{Z_j \in U} P_{ij}^*(t), \quad i = 0, \dots, m$	$IR_{Si}(t, t+\theta) = \sum_{Z_j \in U} P_{ij}(t) R_{Sj}(\theta), \quad i = 0, \dots, m$
	$MTTF_{Si} = \frac{1}{\rho_i} + \sum_{Z_j \in U} \frac{\rho_{ij}}{\rho_i} MTTF_{Sj}, \quad Z_i \in U$ $\dot{P}_{ij}^*(t) = \dot{P}_{ij}^*(t), \quad \text{mit } \dot{P}_{ij}^*(t) \text{ aus}$ $\dot{P}_{ij}^*(t) = -\rho_j P_{ij}^*(t) + \sum_{Z_k \in U} P_{ki}^*(t) \rho_{kj},$ $j = 0, \dots, m, \quad \rho_{ii} = 0$ mit $\dot{P}_{ij}^*(0) = 1$ und $\dot{P}_{ij}^*(0) = 0$ für $j \neq i$	$PA_S = \sum_{Z_j \in U} p_j$ $P_{ij}(t) = P_{ij}(t), \quad \text{mit } P_{ij}(t) \text{ aus}$ $\dot{P}_{ij}(t) = -\rho_j P_{ij}(t) + \sum_{i=0}^m P_{ij}(t) \rho_{ij}$ $j = 0, \dots, m, \quad \rho_{ii} = 0$ mit $P_{ij}(0) = 1$ und $P_{ij}(0) = 0$ für $j \neq i$	$IR_S(\theta) = \sum_{Z_j \in U} p_j R_{Sj}(\theta)$ p_j aus $p_j \rho_j = \sum_{i=0}^m p_i \rho_{ij},$ mit $\rho_{ii} = 0, \quad p_j > 0 \quad \text{und} \quad p_0 + \dots + p_m = 1$
$R_{Si}(t)$	$= \Pr(\text{im Arbeitszustand in } (0, t) Z_i \text{ ist bei } t = 0 \text{ eingetreten}), \quad Z_i \in U$		
$MTTF_{Si}$	$= E[\text{ausfallfreie Arbeitszeit} Z_i \text{ ist bei } t = 0 \text{ eingetreten}] = \int_0^\infty R_{Si}(t) dt = \bar{R}_{Si}(0), \quad Z_i \in U$		
$PA_{Si}(t)$	$= \Pr(\text{im Arbeitszustand zur Zeit } t Z_i \text{ ist bei } t = 0 \text{ eingetreten}), \quad i = 0, \dots, m$		
PA_S	$= \Pr(\text{im Arbeitszustand zur Zeit } t \text{ im asymptotischen } (t \rightarrow \infty) \text{ und im stationären Zustand}) = \lim_{s \rightarrow 0} s \bar{P}A_{Si}(s), \quad i = 0, \dots, m$		
AA_S	$= \lim_{t \rightarrow \infty} \frac{1}{t} E[\text{totale Betriebszeit in } (0, t)] = PA_S$		
$IR_{Si}(t, t+\theta)$	$= \Pr(\text{im Arbeitszustand in } (t, t+\theta) Z_i \text{ ist bei } t = 0 \text{ eingetreten}), \quad i = 0, \dots, m$		
$IR_S(\theta)$	$= \Pr(\text{im Arbeitszustand in } (t, t+\theta) \text{ im asymptotischen } (t \rightarrow \infty) \text{ und im stationären Zustand})$		
U	$= \text{Menge der Zustände, in welchem das System funktionstüchtig ist (up states)}$		
$\bar{U}$	$= \text{Menge der Zustände, in welchem das System als ausgefallen gilt (down states)}$		
$P_{ij}(t)$	$= \Pr(\text{in } Z_j \text{ zur Zeit } t Z_i \text{ ist bei } t = 0 \text{ eingetreten})$		
ρ_{ij}	$= \lim_{\delta t \rightarrow 0} \frac{1}{\delta t} \Pr\{\text{Übergang von } Z_i \text{ nach } Z_j \text{ in } (t, t+\delta t) \text{in } Z_i \text{ zur Zeit } t\}, \quad t \text{ beliebig, z.B. } t = 0 \text{ (gilt nur für Markoff-Prozesse)}$		
S	bezieht sich auf System; $\bar{u}(s) = \int_0^\infty \bar{u}(t) e^{-st} dt = \text{Laplace-Transformation von } u(t)$		

Literaturverzeichnis

- [I] Bemet R., *Modellierung reparierbarer Systeme durch Markoff- und Semi-regenerative Prozesse*. Dissertation ETH Nr. 9682, 1992.
- [2] Birolini A., *Qualität und Zuverlässigkeit technischer Systeme: Theorie, Praxis, Management*. Berlin: Springer Verlag, 3. Aufl. 1991 (Engl. Ed. in prep.); Zuverlässigkeitssicherung von Automatisierungssystemen und -prozessen. e&i, 107 (1990) 5, pp. 25827 1; *Entwicklungs- und Konstruktionsrichtlinien für Zuverlässigkeit, Instandhaltbarkeit und Softwarequalität*. Bericht Z4 an der Professur für Zuverlässigkeitstechnik der ETH Zürich, 1992, 14 pp.; *Modelle zur Berechnung der Rentabilität der Qualitäts- und Zuverlässigkeitssicherung von komplexen Waffensystemen*. Bericht für die GRD Bem, 1986, 46 pp.
- [3] Clemson E., "Corporate Strategies for Information Technology: A Resource-Based Approach", *IEEE Computer*, Vol. 24, Nov 1991, pp. 23-32
- [4] IEEE, Special issue on: Reliability. *IEEE Spectrum*, Oct. 198 1; - : USAF R&M 2000 Initiative. *IEEE Trans. Rel.*, 36(1987)3.
- [5] Irland E.A., Assuring Quality and Reliability of Complex Electronic Systems: Hardware and Software. *Proc. IEEE*, 76(1988)1, pp. 5-18.
- [6] Juran J.M. and Gryna F.M. (Eds.), *Quality Control Handbook*. New York: McGrawHill, 4th Ed., 1988.
- [7] Miteff L., Reduktion der Standbyleistungsaufnahme und Auswirkung auf die Zuverlässigkeit eines Videorekorders und eines Telefax-Geräts, Berichte Z1 und Z2 der Professur für Zuverlässigkeitstechnik der ETH Zürich, 1991, 13 und 1 1 pp.
- [8] MIL-HDBK-338, *Electronic Reliability Design Handbook*, Vol. I and II. 1984.; -HDBK-217: *Reliability Prediction of Electronic Equipment*. Ed. F 1991.
- [9] Quiswater J., Desmedt Y., "Chinese Lotto as an Exhaustive Code-Breaking Machine", *IEEE Computer*, Vol. 24, Nov. 1991, pp. 14-22.
- [10] RADC, *RADC Reliability Engineer's Toolkit*. 1988; *SOAR-7: A Guide for Implementing Total Quality Management*. 1990.
- [II] Shooman M.L., *Probabilistic Reliability - an Engineering Approach*. New York: McGraw-Hill, 1968.
- [12] Taguchi G., *System of Experimental Design - Engineering Methods to Optimize Quality and Minimize Costs*. Vol. 1 & 2. White Plains NY: Unipub., 1987.