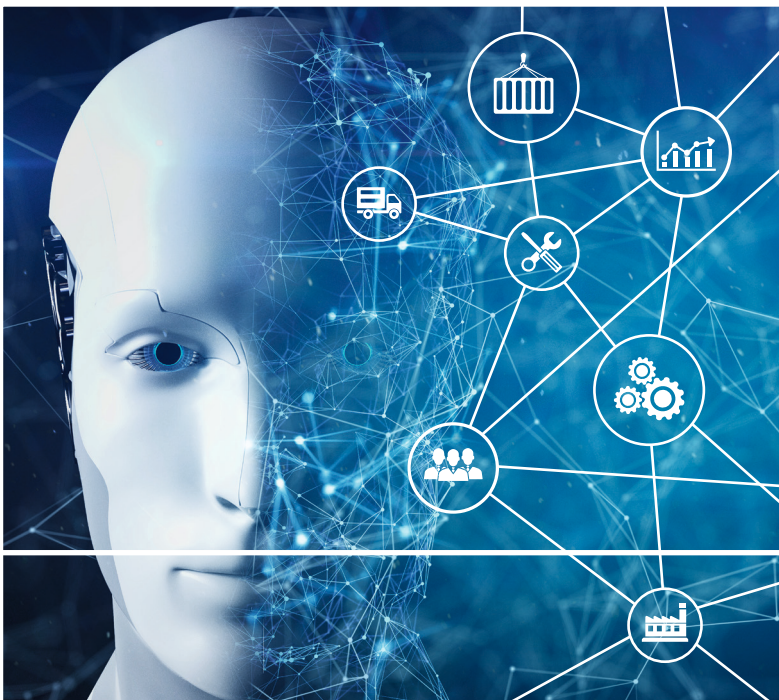


Arno Meyna · Dirk Althaus
Andreas Braasch · Fabian Plinke
Marco Schlummer

Sicherheit und Zuverlässigkeit technischer Systeme



HANSER

Meyna/Althaus/Braasch/Plinke/Schlummer
Sicherheit und Zuverlässigkeit technischer Systeme



Bleiben Sie auf dem Laufenden!

Hanser Newsletter informieren Sie regelmäßig über neue Bücher und Termine aus den verschiedenen Bereichen der Technik. Profitieren Sie auch von Gewinnspielen und exklusiven Leseproben. Gleich anmelden unter

www.hanser-fachbuch.de/newsletter

Arno Meyna
Dirk Althaus
Andreas Braasch
Fabian Plinke
Marco Schlummer

Sicherheit und Zuverlässig- keit technischer Systeme

HANSER

Die Autoren:

Univ.-Prof. Dr.-Ing. habil. Arno Meyna, Emeritus „Sicherheitstheorie und Verkehrstechnik“, Bergische Universität Wuppertal; zugleich Geschäftsführer des Instituts für Qualitäts- und Zuverlässigkeitsmanagement GmbH (IQZ), Wuppertal

Dr.-Ing. Dirk Althaus, Geschäftsführer des Instituts für Qualitäts- und Zuverlässigkeitsmanagement GmbH (IQZ), Wuppertal

Prof. Dr.-Ing. Andreas Braasch, Institut Naturwissenschaften, Hochschule Ruhr West; zugleich Geschäftsführer des Instituts für Qualitäts- und Zuverlässigkeitsmanagement GmbH (IQZ), Wuppertal

Dr.-Ing. Fabian Plinke, Handlungsbevollmächtigter des Instituts für Qualitäts- und Zuverlässigkeitsmanagement GmbH (IQZ), Niederlassung Hamburg

Dr.-Ing. Marco Schlummer, Geschäftsführer des Instituts für Qualitäts- und Zuverlässigkeitsmanagement GmbH (IQZ), Wuppertal

Alle in diesem Buch enthaltenen Informationen wurden nach bestem Wissen zusammengestellt und mit Sorgfalt getestet. Dennoch sind Fehler nicht ganz auszuschließen. Aus diesem Grund sind die im vorliegenden Buch enthaltenen Informationen mit keiner Verpflichtung oder Garantie irgendeiner Art verbunden. Autoren und Verlag übernehmen infolgedessen keine Verantwortung und werden keine daraus folgende oder sonstige Haftung übernehmen, die auf irgendeine Weise aus der Benutzung dieser Informationen – oder Teilen davon – entsteht, auch nicht für die Verletzung von Patentrechten, die daraus resultieren können.

Ebenso wenig übernehmen Autor und Verlag die Gewähr dafür, dass die beschriebenen Verfahren usw. frei von Schutzrechten Dritter sind. Die Wiedergabe von Gebrauchsnamen, Handelsnamen, Warenbezeichnungen usw. in diesem Werk berechtigt also auch ohne besondere Kennzeichnung nicht zu der Annahme, dass solche Namen im Sinne der Warenzeichen- und Markenschutz-Gesetzgebung als frei zu betrachten wären und daher von jedermann benutzt werden dürften.

Bibliografische Information der deutschen Nationalbibliothek:

Die Deutsche Nationalbibliothek verzeichnet diese Publikation in der Deutschen Nationalbibliografie; detaillierte bibliografische Daten sind im Internet unter <http://dnb.d-nb.de> abrufbar.

Dieses Werk ist urheberrechtlich geschützt.

Alle Rechte, auch die der Übersetzung, des Nachdruckes und der Vervielfältigung des Buches, oder Teilen daraus, vorbehalten. Kein Teil des Werkes darf ohne schriftliche Genehmigung des Verlages in irgendeiner Form (Fotokopie, Mikrofilm oder ein anderes Verfahren), auch nicht für Zwecke der Unterrichtsgestaltung, reproduziert oder unter Verwendung elektronischer Systeme verarbeitet, vervielfältigt oder verbreitet werden.

© 2023 Carl Hanser Verlag München

www.hanser-fachbuch.de

Lektorat: Dipl.-Ing. Volker Herzberg, Julia Stepp

Herstellung: Melanie Zinsler

Titelmotiv: © [gettyimages.de/](https://www.gettyimages.de/)Yuichiro Chino und © [shutterstock.com/](https://www.shutterstock.com/)VectorForever

Coverkonzept: Marc Müller-Bremer, www.rebranding.de, München

Coverrealisation: Max Kostopoulos

Satz: Eberl & Koesel Studio, Kempten

Druck und Bindung: CPI books GmbH, Leck

Printed in Germany

Print-ISBN: 978-3-446-46003-4

E-Book-ISBN: 978-3-446-46808-5

Inhalt

Vorwort	XVII
Das Institut für Qualitäts- und Zuverlässigkeitsmanagement (IQZ)	XIX
Teil I: Essenzielle Anforderungen an die Sicherheits- und Zuverlässigkeitstechnik	1
1 Herausforderungen für Staat, Gesellschaft und Unternehmen	3
1.1 Fragmente der Explikation der Sicherheits- und Zuverlässigkeitstechnik	3
1.2 Aktuelle Herausforderungen	11
2 Zuverlässigkeit bei Haftungs- und Gewährleistungsfragen ...	17
2.1 Haftungsgrundlage	17
2.1.1 Außervertragliche Haftung	18
2.1.2 Vertragliche Haftung	19
2.1.3 Stand der Technik	20
2.1.4 Gewährleistungsmanagement zwischen Unternehmen	21
2.2 Schadteilanalyse Feld und No-Trouble-Found-Prozess	23
3 Normative Anforderungen in der Sicherheits- und Zuverlässigkeitstechnik	26
3.1 Einleitung zu normativen Anforderungen im rechtlichen Kontext	26
3.2 Überblick über die Begrifflichkeiten: Safety, Security und Reliability ..	31

Teil II: Zuverlässigkeit im Produktentstehungsprozess	37
4 Zuverlässigkeit im Produktentwicklungsprozess	39
4.1 Zuverlässigkeitsprozess	39
4.2 Bereiche, Rollen und Verantwortlichkeiten	46
4.3 Wirtschaftlichkeitsaspekte und Zuverlässigkeitsziele	48
4.4 Reifegrad, Musterstände und Freigabeprozesse in der Automobilindustrie	51
5 Funktionale Sicherheit im Produktentwicklungsprozess	55
5.1 Der sicherheitstechnische Prozess	55
5.1.1 Allgemeine Einführung in die Funktionale Sicherheit	55
5.1.2 Die Sicherheitsgrundnorm IEC 61508	58
5.1.3 Der sicherheitstechnische Prozess in der zivilen Luftfahrtindustrie	62
5.1.4 Funktionale Sicherheit für Straßenfahrzeuge	74
5.2 Unterstützende und begleitende Prozesse als Grundvoraussetzung für die Funktionale Sicherheit	91
6 Datenquellen und -management	100
6.1 Rohdatenerfassung und -management	101
6.2 Nutzungsdaten	101
6.3 Felddaten	102
6.4 Datenschutz am Beispiel Automotive	105
7 Nutzungs- und belastungsabhängige Produktentwicklung	108
7.1 User Experience, Marktanalysen und Use Cases	110
7.2 Kundencluster und -profilerstellung	112
7.3 Design for Reliability, nutzungs- und belastungsabhängige Zuverlässigkeit	113
Teil III: Grundlagen der Zuverlässigkeits- und Sicherheitsanalyse	115
8 Mathematische Grundlagen aus der Wahrscheinlichkeitsrechnung	117
8.1 Mengenalgebra	117
8.1.1 Grundbegriffe und Definitionen	117

8.1.2	Mengenoperationen	118
8.2	Grundbegriffe der Wahrscheinlichkeitsrechnung	120
8.2.1	Wahrscheinlichkeitsbegriff	120
8.2.2	Axiomsystem von Kolmogorov	121
8.2.3	Die bedingte Wahrscheinlichkeit	124
8.2.4	Unabhängige Ereignisse	126
8.2.5	Regel von der totalen Wahrscheinlichkeit	126
8.2.6	Satz von Bayes	127
8.3	Zufallsgrößen und ihre Wahrscheinlichkeitsverteilung	129
8.3.1	Grundbegriffe	129
8.3.2	Erwartungswert und Momente einer Verteilungsfunktion	133
8.3.3	Quantil, Median und Modalwert	138
9	Zuverlässigkeits- und Sicherheitskenngrößen	141
9.1	Zuverlässigkeitskenngrößen nicht reparierbarer Systeme	141
9.2	Empirische Zuverlässigkeitskenngrößen und weitere Zuverlässigkeitsmerkmale	150
9.3	Zuverlässigkeitskenngrößen reparierbarer Systeme, Instandhaltung	153
9.4	Sicherheitskenngrößen	156
10	Wichtige Verteilungsfunktionen	160
10.1	Wichtige Lebensdauerverteilungen und ihre Zuverlässigkeits- kenngrößen	160
10.1.1	Exponentialverteilung	160
10.1.2	Weibull-Verteilung	164
10.1.3	Die spezielle Erlang-Verteilung	172
10.1.4	Die Normalverteilung	176
10.1.5	Die logarithmische Normalverteilung	179
10.1.6	Asymptotische Extremwertverteilung	184
10.2	Wichtige diskrete Verteilungsfunktionen	190
10.2.1	Binomialverteilung	190
10.2.2	Poisson-Verteilung	193
10.2.3	Hypergeometrische Verteilung	195

11	Ausfallratenmodelle	201
11.1	Zeitliches Verhalten der Ausfallrate	201
11.2	Ausfallratenangaben	202
11.3	Ausfallratendatenhandbücher	204
11.4	Ausfallratenmodelle	207
11.5	Zeitliche Schwankungen der Ausfallrate	214
Teil IV: Methoden der Zuverlässigkeits- und Sicherheitsanalyse		217
12	Einführung in die Methoden der Zuverlässigkeits- und Sicherheitstechnik	219
12.1	Allgemeine Einführung	219
12.2	Methodenvergleich	223
13	Zuverlässigkeitsanalyse einfacher Systemstrukturen	229
13.1	Grafische Darstellung von Systemkonfigurationen	230
13.1.1	Zuverlässigkeits-Blockschaltbild	230
13.1.2	Fehler- oder Funktionsbäume: Darstellung mithilfe logischer Symbole der Booleschen Algebra	230
13.1.3	Zustandsdiagramme (Zustandsübergangsgraphen)	231
13.2	Logisches Seriensystem	232
13.3	Logisches Parallelsystem	233
13.4	Parallel-Seriensystem	237
13.5	Brückenkonfiguration	239
13.6	Berücksichtigung mehrerer Ausfallarten	242
13.6.1	Logisches Seriensystem bei zwei Ausfallarten	244
13.6.2	Logisches Parallelsystem bei zwei Ausfallarten	245
13.6.3	Logisches Parallel-Seriensystem bei zwei Ausfallarten	247
13.6.4	Beliebige Konfigurationen	250
14	Zuverlässigkeitserhöhung in Planung und Praxis	252
14.1	Allgemeine Maßnahmen zur Zuverlässigkeitserhöhung	252
14.2	Begriff und Definition der Redundanz	255
14.3	Redundanzarten, Grundprinzipien	256
14.4	Aktive Redundanz	257

14.5	mvn-System	258
14.6	nvN-System	262
14.7	Standby-System – passive Redundanz	265
15	Systembetrachtung	269
15.1	Begriffliche Exemplifikationen	269
15.2	Technisches System	270
15.3	Elektrik/Elektronik/Software-Systemarchitekturen	272
15.4	Ausfallverhalten von Elektrik/Elektronik/Software-Konfigurationen ..	273
16	Boolesche Modellbildung	275
16.1	Begriffe und Regeln der Booleschen Algebra	275
16.1.1	Die Boolesche Funktion	275
16.1.2	Grundverknüpfungen	277
16.1.3	Axiome der Booleschen Algebra	280
16.1.4	Karnaugh-Veitch-Diagramm	282
16.1.5	Kanonische Darstellung von Booleschen Funktionen	284
16.1.6	Shannonsche Zerlegung	290
16.1.7	Die Boolesche Funktion mit reellen Variablen	292
16.2	Die Systemfunktion	294
16.3	Einführung von Wahrscheinlichkeiten	297
16.4	Fehlerbaumanalyse	299
16.4.1	Einführung	299
16.4.2	Darstellung monotoner Strukturen durch Minimalpfade und Minimalschnitte	302
16.4.3	Quantitative Fehlerbaumauswertung	306
16.5	Importanzkenngrößen	315
16.5.1	Strukturelle Importanz	315
16.5.2	Marginale Importanz	318
16.5.3	Fraktionale Importanz	320
16.5.4	Barlow-Proschan-Importanz	321
16.6	Bestimmung der mittleren Häufigkeit von Systemausfällen sowie der mittleren Ausfall- und Betriebsdauer	324
16.7	Induktive Zuverlässigkeits- und Sicherheitsanalyse	327

17	Zuverlässigkeitsbewertung mithilfe der Fuzzy-Logik	330
17.1	Grundlagen der Fuzzy-Logik	331
17.1.1	Verknüpfung unscharfer Mengen	334
17.1.2	Fuzzy-Relation	336
17.1.3	Erweiterungsprinzip	340
17.2	Prinzipieller Ablauf einer Fuzzy-Anwendung	341
17.2.1	Fuzzifizierung	342
17.2.2	Fuzzy-Inferenz	342
17.2.3	Defuzzifizierung	343
17.3	Anwendung der Fuzzy-Logik bei der FMEA	348
17.3.1	Eingangsgrößen	348
17.3.2	Fuzzifizierung	351
17.3.3	Verarbeitungsregeln	354
17.3.4	Berechnung der Zugehörigkeitsgrade	355
17.3.5	Defuzzifizierung	357
17.4	Fuzzy-Fehlerbaumanalyse	357
17.4.1	Das Fuzzy-Modell	358
17.4.2	Praktisches Anwendungsbeispiel	362
18	Einführung in die stochastischen Prozesse	366
18.1	Beurteilungskriterien stochastischer Prozesse	368
18.1.1	Definitionsspezifische Beurteilungskriterien	369
18.1.1.1	Markov-Bedingungen	369
18.1.1.2	Regenerationspunkte des Prozesses	369
18.1.2	Anwendungsspezifische Beurteilungskriterien	370
18.1.2.1	Akzeptanz von stochastischen Abhängigkeiten zwischen den Elementen des Prozesses	370
18.1.2.2	Anwendbare Verteilungsfunktionen der Zufallszeiten	370
18.1.3	Klassifizierung stochastischer Prozesse anhand der Beurteilungskriterien	371
18.2	Analysemöglichkeiten eines Parallelsystems mit zwei identischen Einheiten	373

19	Markovsche Modellbildung	380
19.1	Der Markovsche Prozess mit diskretem Parameterbereich und endlich vielen Zuständen (Markov-Kette)	380
19.1.1	Zustandsgleichung	380
19.1.2	Zustandsklassen	383
19.1.3	Die absorbierende homogene Markov-Kette	385
19.1.4	Ergodensatz für Markovsche Ketten	389
19.2	Der Markovsche Prozess mit kontinuierlichem Parameterraum und diskretem Zustandsraum	392
19.2.1	Zustandsgleichungen	392
19.2.2	Laplace-Transformation der Zustandsgleichung	398
19.3	Der Semi-Markov-Prozess	405
19.3.1	Einführung	405
19.3.2	Definition und Grundbegriffe	405
19.3.3	Der absorbierende Semi-Markov-Prozess	412
19.3.4	Der ergodische Semi-Markov-Prozess	416
20	Monte-Carlo-Simulation	421
20.1	Einführung	421
20.2	Grundlagen der Monte-Carlo-Simulation	423
20.3	Generierung von Zufallszahlen	425
20.4	Methoden zur Generierung beliebig verteilter Funktionen	429
20.5	Direkte Monte-Carlo-Simulation	432
20.5.1	Generierung eines Zustandsübergangs	432
20.5.2	Last-Event-Schätzer	434
20.5.3	Free-Flight-Schätzer	434
20.6	Anwendungsbeispiel	437
21	Zuverlässigkeitsbewertung mithilfe der Graphentheorie	444
21.1	Gerichteter Graph	445
21.1.1	Einige Grundbegriffe	445
21.1.2	Lineare Flussgraphen	447
21.1.3	Auswertung der linearen Flussgraphen mithilfe der Mason-Formel	450

21.2	Anwendung der linearen Flussgraphen auf diskrete Markov-Prozesse	453
21.2.1	Inhomogene Prozessdarstellung	453
21.2.2	Homogene Prozessdarstellung	454
21.2.3	Asymptotisches Verhalten	457
21.2.4	Erwartungswert und Eintrittswahrscheinlichkeit	457
21.3	Anwendung der linearen Flussgraphen auf stetige Markov-Prozesse	458
22	Neuronale Netze	466
22.1	Grundlagen	467
22.1.1	Das biologische Paradigma	467
22.1.2	Aufbau und Arbeitsweise eines künstlichen Neurons	468
22.1.3	Aufbau eines neuronalen Netzes	472
22.1.4	Arbeitsweise neuronaler Netze	473
22.2	Anwendung in der technischen Zuverlässigkeit	477
22.2.1	Neuronale Schätzung der Parameter einer Verteilungsfunktion	477
22.2.2	Neuronale Zuverlässigkeitsprognose	481
	Teil V: Zuverlässigkeitsprüfung und -bewertung	487
23	Stichprobenverteilung	489
23.1	Stichprobenverteilung des Mittelwertes	489
23.2	Stichprobenverteilung der Varianz	493
23.3	Stichprobenverteilung der Mittelwerte bei unbekannter Varianz	494
23.4	Stichprobenverteilung für die Differenz und Summe zweier arithmetischer Mittelwerte	495
23.5	Stichprobenverteilung des Quotienten zweier Varianzen	497
24	Grenzwertsätze und Gesetze der großen Zahlen	498
24.1	Grenzwertsätze und Approximationen	498
24.1.1	Approximation der Binomialverteilung durch die Poisson-Verteilung	498
24.1.2	Approximation der hypergeometrischen Verteilung durch eine Binomialverteilung	498
24.1.3	Approximation der Poisson-Verteilung durch eine Normalverteilung	499

24.1.4	Approximation der Binomialverteilung durch die Normalverteilung	499
24.1.5	Approximation der hypergeometrischen Verteilung durch die Normalverteilung	500
24.1.6	Zentraler Grenzwertsatz	501
24.2	Gesetz der großen Zahlen	502
24.2.1	Tschebyscheffsche Ungleichung	502
24.2.2	Satz von Bernoulli	504
25	Statistische Schätzung von Parametern	505
25.1	Eigenschaften von Schätzfunktionen	505
25.2	Vertrauensintervalle	507
25.3	Konfidenzintervall für den Erwartungswert und die Varianz bei normalverteilter Grundgesamtheit und Bestimmung des Stichprobenumfangs	508
25.3.1	Konfidenzintervall für den Erwartungswert	508
25.3.2	Konfidenzintervall für die Varianz	512
25.3.3	Bestimmung des Stichprobenumfangs	513
25.4	Die Maximum-Likelihood-Methode (M-L-M)	517
25.4.1	Maximum-Likelihood-Schätzer für die Parameter der Binomial- und Poisson-Verteilung	520
25.4.2	Maximum-Likelihood-Schätzer für den Parameter einer Exponentialverteilung	521
25.4.3	Maximum-Likelihood-Schätzer für die Parameter der Normal- und Lognormalverteilung	521
25.4.4	Maximum-Likelihood-Schätzer für die Parameter der Weibull-Verteilung	521
25.5	Maximum-Likelihood-Methode bei zensierter und gestutzter Stichprobe	524
25.6	Die Momentenmethode	532
25.6.1	Momentenschätzer für den Parameter einer Exponentialverteilung	534
25.6.2	Momentenschätzer für die Parameter einer Lognormalverteilung	536
25.6.3	Momentenschätzer für die Parameter einer Weibull-Verteilung	536
25.7	Lineare Regression und die Methode der kleinsten Quadrate	536

26	Bestimmung des Verteilungstyps	539
26.1	Wahrscheinlichkeitsnetz der Weibull-Verteilung	539
26.1.1	Konstruktion des Wahrscheinlichkeitsnetzes	539
26.1.2	Gebrauchsanweisung für das Wahrscheinlichkeitsnetz der Weibull-Verteilung nach Stange und Gumbel (DGQ-Lebensdauernetz)	540
26.2	Tests zur Überprüfung des Verteilungstyps – Anpassungstests	547
26.2.1	Der Chi-Quadrat-Anpassungstest	547
26.2.2	Der Kolmogorov-Smirnov-Test (K-S-T)	552
26.3	Vergleich der beiden Anpassungstests	559
27	Test- und Prüfplanung – Testverfahren	560
27.1	Statistische Verfahren	564
27.1.1	Der Binomialprüfplan als attributiver Abnahmeprüfplan	564
27.1.2	Sequenzialprüfung	566
27.1.3	Success Run	569
27.1.4	Sudden-Death	574
27.1.5	Vorwissen und Test	581
27.1.6	End-of-Life-Test	581
27.2	Beschleunigte Lebensdauertests	583
27.2.1	Das Arrhenius-Modell	583
27.2.2	Das Eyring-Modell und dessen Modifikation	584
27.2.3	Das Peck-Modell	585
27.2.4	Dauerschwingversuch nach Wöhler	586
27.2.5	Hochbeschleunigte Testmethoden	590
28	Felddatenanalyse	593
28.1	Allgemeine Einführung	593
28.2	Schichtliniendiagramme und Beanstandungsverläufe	596
28.3	Zuverlässigkeitsprognosen für mechatronische Systeme in Kraftfahrzeugen bei nicht vollständigen Daten	604
28.3.1	Zuverlässigkeitsprognosen für Systeme im Kraftfahrzeug während der Nutzungsphase	604
28.3.2	Zuverlässigkeitsprognosen für zeitnahe Garantiedaten	610

Literaturverzeichnis	616
Anhang A	631
Anhang B	649
Anhang C	655
Anhang D	657
Anhang E	660
Index	665

Der Verlag und die Autoren haben sich mit der Problematik einer gendergerechten Sprache intensiv beschäftigt. Um eine optimale Lesbarkeit und Verständlichkeit sicherzustellen, wird in diesem Werk auf Gendersternchen und sonstige Varianten verzichtet; diese Entscheidung basiert auf der Empfehlung des Rates für deutsche Rechtschreibung. Grundsätzlich respektieren der Verlag und die Autoren alle Menschen unabhängig von ihrem Geschlecht, ihrer Sexualität, ihrer Hautfarbe, ihrer Herkunft und ihrer nationalen Zugehörigkeit.

Vorwort

Sicherheit und Zuverlässigkeit zählen zu den wichtigsten Anforderungen, die heute an ein technisches System gestellt werden. Das vorliegende Werk zu dieser vielseitigen Thematik wurde durch ein Autorenteam des Instituts für Qualitäts- und Zuverlässigkeitsmanagement GmbH (IQZ) erstellt, das seit vielen Jahren auf dem Gebiet der Sicherheit und Zuverlässigkeit tätig ist und daher über ein breit gefächertes und fundiertes Spezialwissen verfügt.

Das Werk basiert einerseits auf Vorlesungs- und Seminarunterlagen sowie Abschlussarbeiten einschließlich Dissertationen, wobei Teilbereiche davon auch im Lehrbuch von Meyna/Pauli (siehe Lit.) enthalten sind, das ebenfalls im Carl Hanser Verlag publiziert wurde, und andererseits auf Erkenntnissen der vielfältigen industriellen Projekte und dem damit verbundenen Erfahrungsschatz der Autoren.

Die Implementierung, die Erstellung der Bilder, Tabellen etc. sowie die umfangreiche Korrespondenz mit den Autoren und dem Verlag erfolgte durch Frau Xenia Rein, die mit Verve und vielen Anregungen am Gelingen des Werks einen erheblichen Anteil hatte. Dafür möchten sich die Autoren ganz herzlich bei Frau Rein bedanken.

Dem Lektor des Carl Hanser Verlags, Herrn Dipl.-Ing. Volker Herzberg, möchten die Autoren für seine vielen Anregungen und, insbesondere aufgrund der Covid-19-Pandemie, für das entgegengebrachte Verständnis und Vertrauen ihren Dank aussprechen.

Möge das Werk dazu dienen, die dargelegten Erkenntnisse der hochinteressanten Fachdisziplin *Zuverlässigkeits- und Sicherheitstechnik* in Forschung und Lehre sowie insbesondere in der Praxis zu nutzen, weiterzuentwickeln und damit zur Verbesserung der Qualität technischer Produkte einen Beitrag zu leisten.

Wuppertal/Hamburg im Dezember 2022

Arno Meyna

Dirk Althaus

Andreas Braasch

Fabian Plinke

Marco Schlummer

Das Institut für Qualitäts- und Zuverlässigkeits- management (IQZ)

Alle Autoren dieses Buches sind Geschäftsführer bzw. leitende Mitarbeiter des Instituts für Qualitäts- und Zuverlässigkeitsmanagement GmbH (IQZ).



Ihr Qualitäts-Zulieferer.

Institut für Qualitäts- und Zuverlässigkeitsmanagement GmbH

www.iqz-wuppertal.de

Das IQZ wurde als innovatives Beratungsunternehmen aus der Bergischen Universität Wuppertal heraus gegründet und hat seine Geschäftstätigkeit im Juni 2012 aufgenommen. Das IQZ bietet seinen Kunden wissenschaftsnahe, lösungsorientierte und ganzheitliche Beratungsdienstleistungen, welche den Stand von Wissenschaft und Technik widerspiegeln, unter anderem in folgenden Bereichen:

Data Science

Funktionale
Sicherheit

Maschinen-Sicherheit

Warranty
Management

Zuverlässigkeits-
Management

Risikomanagement

Qualitätsmanagement

Durch die enge Kooperation mit der Bergischen Universität Wuppertal und der Hochschule Ruhr West kann das IQZ nicht nur auf innovative Methoden nach Stand von Wissenschaft und Technik zurückgreifen, sondern deckt, durch jahrelange Projekterfahrung im Kontext namhafter Unternehmen, auch mehrere Jahrzehnte Fach Erfahrung im Sicherheits- und Zuverlässigkeitsmanagement ab. Daher können selbst komplexe Beratungs- und Forschungs-

aufträge durch das IQZ erbracht werden, immer mit dem Fokus, maßgeschneiderte Konzepte für die Kunden zu entwickeln und diese Konzepte pragmatisch in die Praxis umzusetzen.

Neben den Forschungs- und Entwicklungsprojekten im Rahmen des Produktentstehungs- und Produktbeobachtungsprozesses von der Planung, Entwicklung, Produktion, Nutzung bis zur Entsorgung in vielen industriellen Bereichen ist das IQZ seit vielen Jahren auch im Haftungs- und Versicherungsbereich mit Tätigkeiten betraut. Dies beinhaltet z.B. die Haftungsabwehr und Verhandlungsunterstützung im Regressfall oder Gutachtertätigkeiten für Versicherungskonzerne. Namhafte, weltweit agierende Versicherer sowie führende Anwaltskanzleien im Bereich Produkthaftung gehören hier zu unseren langjährigen Partnern. Das Kundenspektrum des IQZ reicht vom Kleinstunternehmen, kleinen und mittleren Unternehmen (KMU) bis zum DAX-30-Unternehmen z.B. in der Automobil- und Luftfahrt-industrie.

Kontakt

Institut für Qualitäts- und Zuverlässigkeitsmanagement GmbH

Heinz-Fangman-Str. 4

42287 Wuppertal

Tel.: +49(0)202/515 616 90

Fax: +49(0)202/515 616 89

Mail: info@iqz-wuppertal.de

Teil I:

Essenzielle Anforderungen an die Sicherheits- und Zuverlässigkeitstechnik

„Der Mangel an Gewissheit ist die wichtigste Quelle unseres Wissens.“

Carlo Rovelli (1956)*

Herausforderungen für Staat, Gesellschaft und Unternehmen

■ 1.1 Fragmente der Explikation der Sicherheits- und Zuverlässigkeitstechnik

Ein kurzer kulturhistorischer Rückblick

Sicherheits- und Zuverlässigkeitstechnik als interdisziplinäres Wissenschaftsgebiet ist eng mit der Entwicklung der Naturwissenschaft und Technik, aber auch mit der Gesellschaft und dem einzelnen Menschen selbst verbunden.

Der moderne Mensch (*Homo sapiens*, lat. „der weise, kluge Mensch“) wird gegenwärtig auf ein Alter von ca. 233 000 Jahren datiert, bezogen auf Knochen und Schädelfragmente des „Kibish Oma I“ in Ägypten (neue Erkenntnisse von Celine Vidal et al., publiziert in der englischsprachigen renommierten Fachzeitschrift „nature“), und breitete sich vom afrikanischen Kontinent ausgehend nach Europa aus. Nachgewiesen sind der *Homo sapiens* und der Neandertaler in Vorderasien vor etwa 40 000 Jahren und in Mitteleuropa vor ca. 10 000 Jahren. Die Begleitfunde, wie Pfeilspitzen, Faustkeile, Steinbeile etc., zeigen die enge Verbundenheit des modernen Menschen mit der Technik. Demnach war der Mensch bei seinem Erscheinen auf der Erde sogleich Techniker. Technik ist also menschliche Uralage (sinngemäß nach Spengler).

Oswald Spengler schreibt weiter (siehe Lit.):

„Um das Wesen des Technischen zu verstehen, darf man nicht von der Maschinenteknik ausgehen, am wenigsten von dem verführerischen Gedanken, daß die Herstellung von Maschinen und Werkzeugen der Zweck der Technik sei.“

Mit der Entwicklung der menschlichen Sprache und der Hochkulturen in Ägypten und Mesopotamien entstand eine neue Qualität der Naturwissenschaft, Technik und Gesellschaft. Es entstanden technische Glanzleistungen, die heute noch bewundert werden, wie z. B. die Pyramiden von Gizeh oder die Megalithbauten in Spanien, deren Errichtung nur durch die systematische Nutzung der damaligen Naturwissenschaft und Technik, verbunden mit einem großen Einsatz von Mensch, Tier und Organisation, ermöglicht wurde.

Im Gegensatz zur „*Techné*“ als Inbegriff allen Könnens, das auf Wissen begründet ist, z. B. Malen und Musizieren als persönliche Fertigkeit, ist die Technik durch Ziele, z. B. Totenkult, und Zwecke, z. B. Bau einer Pyramide, charakterisiert. Nach Dessauer (siehe Lit.) sind technische Objekte durch Raumformen (räumliche Gebilde, ein Gerät, eine Maschine) und Zeit-

formen (eine Methode, ein Verfahren) gekennzeichnet, die aufgrund ihrer Finalität über den technischen Bereich hinausweisen können.

„Technik kann also mit Recht ein Real-Sein und Real-Werden aus Ideen genannt werden, ein dauerndes Hinüberschreiten von Zweckformen aus der Immanenz in die Erfahrungswelt, die hierdurch im Zeitverlauf bereichert wird“ (Dessauer).

Nach Bringmann ist Technik mehr als eine Maschine, ein Produktionsprozess; Technik ist selbst Wissenschaft. Donald Bringmann schreibt hierzu im Werk von Dessauer mit dem Titel „Streit um Technik“ (siehe Lit.):

„Wer in der Technik nur eine angewandte Wissenschaft sieht, lässt das Wesentliche unbetrachtet, das in den technischen Erfindungen und Konstruktionen steckt: jene irrationale seelische Triebkraft, die sich in allen noch so zweckrationalen technischen Gebilden zugleich verbirgt und kundgibt.“

Betrachtet man nun die Entwicklung der Sicherheits- und Zuverlässigkeitstechnik als wissenschaftliche Teildisziplin der Technik, so war und ist diese zunächst final mit der technologischen Entwicklung selbst verbunden. Die frühen Menschen kannten den Nutzen ihrer Waffen und Gerätschaften, aber auch die Gefahren, die mit einer unsachgemäßen Benutzung verbunden waren. Durch die Weiterentwicklung der Technik gelang es den Menschen jedoch, den vielfältigen Urgefahren der Natur und Tierwelt entgegenzutreten und ihre Lebensgrundlage zu sichern.

Nachgewiesen ist auch, dass es zur Zeit der ersten Hochkulturen bereits ein – wie wir heute sagen würden – sicherheitstechnisches Recht gab. Bekannt ist in diesem Zusammenhang der Codex, der unter der Herrschaft des Hammurapi (auch Hammurabi), 6. König der ersten Dynastie (1792 bis 1750 v. Chr.), im altbabylonischen Reich (Babylon wurde 1894/1830 v. Chr. vom semitischen Stamm der Amoriter unter Sumu-abum gegründet) entstand, als älteste bekannte Gesetzessammlung – mit 270 Rechtsgrundsätzen für das tägliche Leben – der Welt, aufgezeichnet auf einer ca. 2,25 m hohen Stele, die 1902 bei Ausgrabungen in Susa gefunden wurde und sich heute im Louvre in Paris befindet (Wikipedia).

Für den Sicherheitstechniker interessant ist in diesem Zusammenhang der Codex der Haftung, welcher in den §§ 228 bis 233 (DIN Mitteilungen, 10/78) beschrieben ist:

- *„Wenn ein Baumeister ein Haus für einen Mann baut und es für ihn vollendet, so soll dieser ihm als Lohn zwei Sekel Silber geben für je ein Sar Haus: (Anm.: 1 Sekel = 360 Weizenkörner = 9,1 Gramm, 1 Sar = 14,88 m²).“*
- *„Wenn ein Baumeister ein Haus baut für einen Mann und macht seine Konstruktion nicht stark, so dass es einstürzt und verursacht Tod des Bauherrn: dieser Baumeister soll getötet werden.“*
- *„Wenn der Einsturz den Tod eines Sohnes des Bauherrn verursacht, so sollen sie einen Sohn des Baumeisters töten.“*
- *„Kommt ein Knecht des Bauherrn dabei um, so gebe der Baumeister einen Knecht von gleichem Wert.“*
- *„Wird beim Einsturz Eigentum zerstört, so stelle der Baumeister wieder her, was immer zerstört wurde; weil er das Haus nicht fest genug baute, baut er es auf eigene Kosten wieder auf.“*
- *„Wenn ein Baumeister ein Haus baut und macht die Konstruktion nicht stark genug, so dass eine Wand einstürzt, dann soll er sie auf eigene Kosten verstärkt wieder aufbauen.“*

Weiter findet sich in der Bibel der Juden (600 bis 200 v. Chr.) (5. Buch Moses, 22/8):

„Wenn du ein neues Haus baust, sollst Du um die Dachterrasse eine Brüstung ziehen. Du sollst nicht dafür, dass jemand herunterfällt, Blutschuld auf Dein Haus laden.“

Der vorangehend aufgeführte Codex der Haftung zeigt, dass der Baumeister eines Hauses bereits vor 4000 Jahren die Sicherheit implementierte, d.h. nach heutigem Verständnis „in ein Produkt/einen Produktionsprozess hineinentwickelte“. Neben den Gefahren durch die Technik war der Mensch seit jeher einer Vielzahl natürlicher Gefahren einschließlich gesellschaftsbedingten und politischen Gefahren – wie die vielen Kriege und die damit verbundenen schrecklichen Auswirkungen zeigen – ausgesetzt. Nach Heidegger¹ ist die menschliche Existenz ein „*Sein zum Tode*“. Unbestritten ist jedoch, dass die stetig fortschreitende technische Zivilisation umfassende ökonomische, medizinische und soziale Vorteile mit sich brachte, ein „menschenwürdiges Leben“ ermöglichte und zu einer erheblichen Verlängerung des menschlichen Lebens führte. Diese Kohärenz von Mensch und Technik als Grundpfeiler von Lebensqualität, Freiheit und Entfaltung wurde mit der Entwicklung großtechnischer Systeme und deren möglichen negativen Auswirkungen für Mensch und Umwelt, die sich aufgrund neuer naturwissenschaftlicher und technischer Erkenntnisse feststellen ließen, zu Beginn des 20. Jahrhunderts zur Inkohärenz und ist heute Gegenstand kontroverser Diskussionen über das „Für und Wider der Technik“.

Ratzinger schrieb hierzu (siehe Lit.):

„Wenn es zutrifft, daß der innere Ausgangspunkt der Technik in der Erringung von Freiheit durch Gewähren von Sicherheit lag, so muß es von da aus auch die innere Forderung jeder technischen Entwicklung und ihr eigenes Leitmaß sein, sie so zu gestalten, daß daraus nicht größere Unsicherheit und in der Steigerung von Abhängigkeiten größere Unfreiheit entsteht. Sie wird dabei erkennen müssen, daß nicht (wie es anfangs aussah) technisches Tun als solches schon befreiendes und damit sittliches Tun ist, daß aber technisches Tun von sittlichen Maximen geleitet sein muß, um seinem eigenen Ursprung zu genügen, der in einer sittlichen Idee lag.“

Die systematische Analyse von Gefahrenpotenzialen gleich welcher Art, denen der Mensch in seinem Dasein ausgesetzt ist, deren Risiken er bewusst oder unbewusst eingeht, und die Entwicklung von entsprechenden Schutz- und Verhaltensmaßnahmen zu deren Bewältigung sind seit jeher Aufgabe der entsprechenden wissenschaftlichen Fachdisziplinen der Naturwissenschaft, Ingenieurwissenschaft, Mathematik, Geistes- und Gesellschaftswissenschaft, Medizin und Psychologie sowie deren Spezialisierungen, da nur in diesen das entsprechende Fachwissen vorhanden ist.

Allerdings zeigte es sich, dass die methodische und inhaltliche Fokussierung auf eine oder mehrere Fachdisziplinen der Bedeutung der Sicherheit für den einzelnen Menschen, aber auch die Gesellschaft mit ihrem ethischen, humanen, wirtschaftlichen und politischen Grundverständnis nicht gerecht wird. Hierzu schrieben Peters und Meyna im Vorwort (siehe Lit.):

„Die Sicherheitstechnik ist eine interdisziplinäre Wissenschaft, deren Schwerpunkt traditionell die Ingenieurwissenschaften bilden. Aber auch die Human- und Sozialwissenschaften, Recht, Ökonomie, Management, Personen- und Objektschutz, Rettungswesen, Umweltschutz, Datenschutz u. a., leisten heute einen nicht unerheblichen Beitrag zur Reduzierung und Bewertung des Risikos und der sicheren Nutzung eines Mensch-Maschine-Umwelt-Systems.“

¹⁾ Martin Heidegger (1889 – 1976)

Risiko in der modernen globalen Industriegesellschaft

Wie bereits dargelegt, sind der Mensch in seinem Dasein und die Gesellschaft vielfältigen natürlichen und zivilisationsbedingten Risiken ausgesetzt. Die Herkunft des Begriffs Risiko ist nicht eindeutig geklärt. Vielfach wird von einem arabischen Ursprung mit dem Ausdruck „risqu“ (von Gottes Gnaden abhängiger Lebensunterhalt) bzw. vom im Mittelalter ins Italienische übernommenen „risico“ (Wagnis, Gefahr bei einer Schiffsreise oder militärischen Unternehmen) ausgegangen.

Mit Risiko wird heute in der Regel das Produkt aus der Eintrittswahrscheinlichkeit eines bestimmten Ereignisses und dem Schadensausmaß (Ereignisschwere) bezeichnet (siehe Abschnitt 9.4).

Durch Logarithmierung und Darstellung als Geradengleichung erhält man eine Gerade gleichen Risikos (Bild 9.5). Das heißt, ein formal gleiches Risiko tritt bei geringer Eintrittswahrscheinlichkeit und großen Auswirkungen (z. B. bei einem Flugzeugabsturz), aber auch bei einer großen Eintrittswahrscheinlichkeit und geringeren Auswirkungen (z. B. im Straßenverkehr) ein.

Interessant ist in diesem Zusammenhang, dass in der Bevölkerung bezüglich des Erstgenannten eine Risikoaversion besteht. Es wurde deshalb vorgeschlagen, einen Risikoaversionsfaktor in die Risikobeurteilung einzufügen. Allerdings fehlt hierfür die wissenschaftliche Grundlage. Problematisch ist jedoch, dass vielfach von „punktförmigen“ Risiken, z. B. Toten/Jahr und Einwohnern, ausgegangen wird und entsprechende Vergleiche z. B. für tödliche Arbeitsunfälle nach Wirtschaftszweigen oder Vergleiche von statistisch abgesicherten natürlichen Risiken mit probabilistisch ermittelten technischen Risiken (z. B. Kernkraftwerke) erfolgen, um so eine gewisse Akzeptanz zu ermöglichen (siehe deutsche Risikostudie, „Kernkraftwerke“, siehe Lit.). Geht man davon aus, dass die Eintrittshäufigkeit (Wahrscheinlichkeit) durch eine entsprechende Verteilungsfunktion und die Ereignisschwere durch eine weitere entsprechende Funktion gegeben ist, so ist auch das Risiko durch eine entsprechende Zeit-Orts-Funktion gegeben.

Das zivilisationsbezogene Risiko ist immer mit einer menschlichen Handlung verbunden, deren Folgen nicht absehbar sind. Es müssen vielfach Entscheidungen getroffen werden, auch dann, wenn nicht alle Fakten und Sachverhalte überprüfbar und bekannt sind. Als weiterer Faktor ist die Kontrollierbarkeit der möglichen Folgen (z. B. von Atomkraft, Gentechnik, neuen Technologien → Technikfolgenabschätzung) von großer Bedeutung. Problematisch ist aber auch die Bewertung der zukünftigen Eintrittswahrscheinlichkeit eines Ereignisses aus den bisher ermittelten, zumal sich die Randbedingungen ständig ändern können. Ferner ist die Freiwilligkeit, mit der der Einzelne oder bestimmte gesellschaftliche Gruppierungen ein Risiko eingehen, von Bedeutung. So gesehen stellt das Risiko eine mehrdimensionale Größe dar.

Der Begriff des Risikos wird oft synonym zum Begriff der Gefahr (gevare, mittelhochdeutsch „Hinterhalt“, „Betrug“) verwendet. Die Gefahr als Möglichkeit eines zukünftigen Schadens oder Nachteils ist im Gegensatz zum vordefinierten objektiven Risiko eine qualitative Größe für die Kennzeichnung einer natürlichen oder zivilisationsbedingten Gefahrenquelle. Diese sind nicht immer latent vorhanden und können unbekannt sein.

Als komplementäre Größe wird vielfach der Begriff der Sicherheit verwendet. Im Sinne der Sicherheitstheorie ist die additive Verknüpfung von Sicherheitswahrscheinlichkeit und Gefährdungswahrscheinlichkeit durch eins gegeben (siehe Abschnitt 9.4). Das heißt, wenn

keine Gefährdung vorliegt, z.B. keine Lawinengefahr im Mittelgebirge im Sommer, so ist die Sicherheit durch 100 % (d.h. eine absolute Sicherheit) gegeben.

Mit dem Begriff der Gefahr werden häufig zeitliche Abstände, z.B. eine konkrete Gefahr, eine aktuelle Gefahr, eine Gefahr im Verzug, und Wahrscheinlichkeiten, wie z.B. eine abstrakte Gefahr, miteinander verknüpft. In diesen Fällen ist die Quantifizierung über das objektive Risiko aussagefähiger. Dabei ist zu beachten, dass ein Risiko nur dann vorhanden ist, wenn eine Gefahr und eine Exposition gegenüber derselben gegeben sind. So sind beispielsweise die Gefahren des Straßenverkehrs für den Nichtteilnehmer irrelevant, für die Gesellschaft aber sehr bedeutend. Das heißt, eine Gefahr ist in der Regel immer durch einen örtlich und zeitlich begrenzten Gefahrenwirkungsbereich charakterisiert, wobei dieser jedoch fließend und global wirkend sein kann (z.B. Atomkatastrophen von Tschernobyl, 1986, Fukushima, 2011, Giftgasunfälle in Seveso, 1976, und Bophal, 1984).

Die örtlichen und zeitlichen Auswirkungen können auch den Lebensraum für nachfolgende Generationen stark einschränken. Das bedeutet aber auch, dass zu den global wirkenden natürlichen Katastrophen (Klima, Krankheiten, Seuchen und andere) und den durch die Gesellschaft selbst verursachten Katastrophen (z.B. Krieg, Terrorismus) eine neue Qualität des Risikos, begründet durch den naturwissenschaftlichen und technischen Fortschritt, hinzugekommen ist.

Nach Bieri ist Risiko ein Bestandteil des menschlichen Daseins, eine risikofreie Lebensform ist nicht denkbar und gleichbedeutend mit dem Tod. Ernst Bieri (siehe Lit.) schreibt hierzu treffend:

„Kein Leben und damit auch keine Technik ohne Risiko: Man kann und soll das Risiko durch den Einsatz des verfügbaren Wissens und Könnens auf ein Minimum herabdrücken, und diese Marschrichtung ist von der Technik eingeschlagen worden, von der Verbesserung der Maschinen in den Fabriken über die Verkehrsmittel aller Art, bis zu den Anlagen für Energiegewinnung und den Umweltschutz. Aber jedes Risiko vollständig und für alle Zukunft ausschalten zu wollen, das wäre Hybris, wäre die Anmaßung der Allmacht durch den Menschen. Das Ende der von uns überblickbaren Risiken tritt mit dem Tod ein. Wer also jedes Risiko ausschalten will, beendet das menschliche Leben, das ohne Risiko, ohne Angst, aber auch ohne den dauernden und über weite Strecken erfolgreichen Kampf dagegen nicht möglich ist.“

Die Ermittlung von Gefahrenpotenzialen und Risiken bis zu einem gewissen Grad beherrschbar zu machen, ist eine interdisziplinäre Aufgabe der Sicherheits- und Zuverlässigkeitstechnik. Allerdings tritt in diesem Zusammenhang sofort die Frage nach der Akzeptanz eines verbleibenden Risikos auf, d.h. das Abwägen zwischen dem Nutzen (z.B. einer neuen Technologie, eines neuen Medikamentes) und den möglichen Schäden für den Menschen und seinen Lebensraum (Arbeitswelt, Umwelt, Gesundheit und anderes). Als Beispiel hierzu sei die kontrovers geführte Diskussion über das Für und Wider der Kernenergie aufgeführt. Das Beispiel zeigt auch die zeitliche Veränderbarkeit der Bewertung von Risiken durch die Gesellschaft und insbesondere durch bestimmte Gruppierungen der Gesellschaft, unabhängig von neuen wissenschaftlichen Erkenntnissen. Karl Steinbuch schreibt hierzu (siehe Lit.):

„Die Akzeptanz technischer Risiken hat drei Dimensionen:

Erstens die technische Dimension: Die Verminderung technischer Risiken durch Menschen, Methoden oder Material.

Zweitens die pädagogische Dimension: Die Aufklärung über die Zwangsläufigkeit mancher Risiken.

Drittens die ethische Dimension: Die Erzeugung von Verantwortung und Vertrauen.

Auf diesem Gebiet wurde in den letzten Jahren vieles versäumt – wir alle werden hierunter leiden: Das Vertrauenspotential entscheidet über den erreichbaren Wohlstand – die Zerstörung des Vertrauenspotentials (z. B. durch Konfliktideologie oder Skandalpublizistik) führt zwangsläufig zur Verminderung unseres Wohlstandes.“

Das Erkennen einer Gefahr und die Beurteilung des Risikos mit seiner Eintrittswahrscheinlichkeit und den möglichen Schadensauswirkungen ist a priori immer mit mehr oder weniger großen Unsicherheiten verbunden, da besonders bei seltenen Ereignissen eine statistische Verifizierung nicht möglich ist. Folglich gibt es kein sogenanntes „Restrisiko“ und „Sicherheitsrisiko“. Dies gilt auch für die Schadensauswirkungen.

Thomas A. Jäger schreibt hierzu (siehe Lit.):

„Die lange Inkubationszeit der physischen Schädigung des Menschen durch schleichend akkumulierende Umwelteinflüsse, wie Wasser- und Luftverschmutzung oder durch in Nahrungsmitteln enthaltene Schadstoffe, macht das Erkennen der Zusammenhänge schwierig. Das gleiche gilt für die physische und psychische Schädigung durch Lärm. Der Nachweis gesundheitlicher Schädigungen ist umso schwieriger, als die Symptome schleichend und nicht eindeutig und ihre Quellen überaus vielfältig sind. Wenn jemand durch schädigende Umwelteinflüsse in seiner Vitalität herabgesetzt oder gar krank ist, so fällt dies lange Zeit nicht weiter auf und die Ursachen bleiben im Nebel.“

Ungeachtet dessen zeigen die in diesem Zusammenhang vielfältig eingesetzten wissenschaftlichen Methoden, Verfahren, Simulationen und anderes zur Sicherheit, Zuverlässigkeit, Gesundheit, Umweltverträglichkeit, Schutz, dass Risiken objektiviert und minimiert werden können. Das Grundgesetz der Bundesrepublik Deutschland sagt hierzu in Artikel 1 unmissverständlich:

„(1) Die Würde des Menschen ist unantastbar. Sie zu achten und zu schützen ist Verpflichtung aller staatlichen Gewalt.“

Zum Paradigma der Sicherheits- und Zuverlässigkeitstechnik

Aufgrund der zuvor erläuterten engen Verknüpfung der Sicherheits- und Zuverlässigkeitstechnik mit den Natur- und Technikwissenschaften und deren deterministischem Weltbild (jeder Vorgang hat eine Ursache, die Wirkung kann sich höchstens mit der endlichen Lichtgeschwindigkeit ausbreiten) ist dieser Kausalnexus des klassischen Determinismus auch Paradigma der Sicherheits- und Zuverlässigkeitstechnik.

Die Determiniertheit der Welt und des Menschen im Kontext von „Willensfreiheit und bedingtem freien Willen“ ist bis in unsere heutige Zeit Gegenstand philosophischer, psychologischer, neurophysiologischer und anderer Betrachtungen.

Ungeachtet dessen ist der naturwissenschaftliche und technische Fortschritt seit Newton² und Laplace³ durch den Determinismus geprägt. Nach Laplace ist die Zukunft durch die Gegenwart (Anfangsbedingung z. B. einer Differenzialgleichung) eindeutig bestimmt (siehe

²⁾ Isaac Newton (1643 – 1727)

³⁾ Pierre-Simon Laplace (1749–1827)

auch das Gedankenkonstrukt des „Laplaceschen Dämons“ als einer „Intelligenz“, die befähigt ist, alle Determinationen und Faktoren in Vergangenheit und Zukunft genauestens zu kennen). Dem gegenüber steht das indeterministische Weltbild der Quantenmechanik (Heisenberg⁴, Born⁵, Schrödinger⁶ und andere), d.h. die Nichtbestimmtheit der Ursachen bei physikalischen Vorgängen – Zukunft ist nur durch stochastische Modellbildung (Zufalls-Prozess) bestimmbar –, welches in Widerspruch zum Laplaceschen Weltbild steht (siehe auch 2. Hauptsatz der Thermodynamik). Man beachte in diesem Zusammenhang aber auch die streng deterministische Position von Albert Einstein⁷ „Gott würfelt nicht“ (sinngemäß aus dem Briefwechsel mit Max Born (siehe Lit.)), die heute jedoch für den Mikrokosmos als widerlegt angesehen wird.

Ähnlich wie in der Physik fand unter anderem durch die Arbeiten von Wöhler, Weibull (siehe Lit.) und weiteren Autoren im Bereich der Werkstoffermüdung und insbesondere in den 40er- und 50er-Jahren durch Pieruschka und Lusser die Stochastik als „Systemdenken“ Einzug in die Ingenieurwissenschaften. So schreibt Pieruschka, der sein Buch Robert Lusser widmete, im Vorwort (siehe Lit.):

„In the year 1943, the author became, for the first time, aware of the reliability problem during flight testing of the German V-1 missile. Intensive studies of this subject were started in 1954 upon coming to the United States. In the subsequent eight years, the author has completed a great amount of study in the field of reliability theory as research scientist with the Army Rocket and Guided Missile Agency, Redstone Arsenal, Alabama, and with the Lockheed Missile and Space Company, Sunnyvale, California.“

Weiter heißt es auf Seite 47:

„The Swiss mathematician Jakob Bernoulli developed the rule which published 1713. Robert Lusser has been stressing the serious consequence of the product rule in the achievement of reliability in complex equipments since the year 1950, and by this insistence has brought into being reliability as a serious profession. Therefore, Formula (siehe Formel 13.1) is called Robert Lusser's reliability formula.“

Die zweite Auflage des 1962 erschienenen Buches von Igor Bazovsky (siehe Lit.) stellt ebenfalls die Pionierleistungen von Lusser (Seite 275) und Pieruschka (Seiten 60 und 71) heraus.

„Another historical development in the same direction took place in Germany during the Second World War. Robert Lusser, one of the reliability pioneers, narrates how he and his colleagues, while working with Wernher von Braun on the V1 missile, met with the reliability problem. The first approach they took towards V1 reliability was that a chain cannot be stronger than its weakest link. Thus, the missile will be as reliable as the weakest link can make, or as strong.“

Nicht unerwähnt sei, dass bereits 1954 (danach jährlich) in den USA das 1. National Symposium on Reliability and Quality Control der IEEE und in Deutschland 1961 (danach alle zwei Jahre) die erste Tagung über Zuverlässigkeit stattfand. Auf Anregung von Ludwig Bölckow, einem der großen Förderer der Deutschen Zuverlässigkeit (siehe Lit.), wurde bereits

⁴) Werner Heisenberg (1901–1976)

⁵) Max Born (1882–1970)

⁶) Erwin Schrödinger (1887–1961)

⁷) Albert Einstein (1879–1955)

1964 der Fachausschuss „Zuverlässigkeit und Qualitätskontrolle“ beim VDI gegründet. Es entwickelte sich die Zuverlässigkeitstheorie als stochastische Theorie zur Bewertung komplexer Mensch-Maschine-Umweltsysteme. Das Paradigma der Zuverlässigkeitstheorie ist das stochastische Ausfall- bzw. Lebensdauerverhalten von Komponenten und Systemen. Das heißt, die Lebensdauer einer Komponente – wie die des Menschen – ist eine Zufallsvariable. Die Zuverlässigkeitstheorie verfügt wie in diesem Buch dargestellt über ein theoretisches Konzept, einschließlich Methoden, Verfahren und Kenngrößen zur qualitativen und quantitativen Bewertung technischer Systeme gleich welcher Art. Da die Sicherheit als Teilmenge der Ausfall- und Betriebszustände eines Systems angesehen werden kann, ist diese isomorph; das bedeutet, das theoretische mathematische Gebäude der Zuverlässigkeitstheorie bildet auch das theoretische Gebäude der stochastischen Sicherheitstechnik. Die stochastische Modellbildung kann allerdings eine deterministisch orientierte Sicherheits- und Zuverlässigkeitstechnik nicht ersetzen, sondern ermöglicht a priori eine Bewertung im frühen Entwicklungsstadium und im Rahmen des Produktentstehungsprozesses (siehe Kapitel 4) und der Test- und Prüfplanung (siehe Kapitel 27). Die vorangegangenen Betrachtungen zeigen, dass die Paradigmen der Sicherheits- und Zuverlässigkeitstechnik durch den Nexus von Determinismus und Indeterminismus geprägt sind. Allerdings fehlt in diesen Paradigmen eine Unbekannte, und dies ist der Mensch als Individuum und Mittelpunkt des Seins, eingebunden in Gesellschaft, Umwelt, Technik mit all ihren kulturellen und traditionellen Unterschieden. Albert Kuhlmann schlägt deshalb einen „kybernetischen“ Ansatz der Sicherheitswissenschaft vor (siehe Lit.):

„Die Verknüpfung zwischen Kybernetik und Sicherheitswissenschaft ergibt sich daraus, dass technische Einrichtungen von Menschen bedient und kontrolliert werden, die sich selbst im Wirkungsbereich dieser technischen Einrichtung befinden. Sie nutzen die Maschine, sind aber auch ihren Gefahren für Leib, Leben und Sachen ausgesetzt. Das menschliche Verhalten sowie das Verhalten der Maschine sind wiederum von den Bedingungen der Umwelt abhängig. Die Umwelt technischer Einrichtungen wird ihrerseits oft von ihnen vielfältig beeinflusst, z.B. durch die Abgabe von Abfällen, Abwässern, Lärm und luftfremden Stoffen. Der Mensch wiederum kann Einfluss nehmen auf derartige Umweltfaktoren. Jede technische Einrichtung ist deshalb eingebettet in ein durch Wechselwirkungen gekennzeichnetes Mensch-Maschine-Umwelt-System.“

Laut Kuhlmann befasst sich die Sicherheitswissenschaft mit der Sicherheit vor möglichen Gefahren bei der Nutzung der Technik, nicht aber mit militärischer und sozialer Sicherheit oder mit der Sicherheit vor Krankheit, die nicht von der Technik verursacht wird.

„Das Ziel der Sicherheitswissenschaft besteht darin, Schadwirkungen bei der Nutzung der Technik so klein wie möglich oder wenigstens in vorgegebenen Grenzen zu halten. Unter Schadwirkungen sollen dabei sowohl unfallartige als auch solche Schäden verstanden werden, die z.B. infolge von Umweltbelastungen durch technische Anlagen entstehen können.“

Der Begriff der Sicherheitswissenschaft ist bei Kuhlmann technikbezogen. Das heißt, ein Schutz vor Gefahren, die nicht durch die Technik verursacht sind, ist nicht Gegenstand der Sicherheitswissenschaft. Hierzu ist anzumerken, dass eine klare Trennung zwischen den einzelnen zivilisationsbedingten und den natürlichen Risiken nicht immer möglich ist (zum Beispiel beim Klimawandel). Die Verwendung des Terminus Sicherheitswissenschaft statt Sicherheitstechnik ist strittig. Wie zuvor erläutert, ist die Technik allgemein und damit auch die Sicherheitstechnik eine Wissenschaftsdisziplin, die allerdings als eine interdisziplinäre multifakultative Disziplin anzusehen ist. Unstrittig ist jedoch, wie bereits erwähnt,

dass der Mensch eine besondere Herausforderung für die sicherheits- und zuverlässigkeitstechnische Gestaltung und Nutzung eines Mensch-Maschine-Umwelt-Systems darstellt. Neben den Paradigmen Determinismus und Indeterminismus sei hier von einem **Humankompatibilismus** als drittes Paradigma ausgegangen. Im Gegensatz zum sogenannten „human factor“ – gekennzeichnet durch die Anpassung „Mensch/Maschine“, geprägt durch psychologische, medizinische und andere Erkenntnisse – soll das hier neu eingeführte Paradigma „Humankompatibilismus“ auch die Anpassung „Maschine/Mensch“ (human engineering) durch das Wissensgebiet der Ergonomie und Anthropotechnik mitberücksichtigen. Des Weiteren sollen die Wissenschaftsdisziplinen, die sich mit den Gefahrenpotenzialen und deren Reduzierung bei der Herstellung und Nutzung technischer Systeme gleich welcher Art auseinandersetzen, Bestandteil des Humankompatibilismus sein. Hierzu zählen insbesondere die Verkehrssicherheit, Arbeitssicherheit, Arbeitsmedizin, Arbeitsphysiologie, Pädagogik und andere Disziplinen. Dieses hier neu eingeführte Paradigma „Humankompatibilismus“ berücksichtigt die physischen, psychischen, geistigen und anderen Eigenschaften des Menschen im Sinne einer teleologischen Interaktion mit der Technik und der damit verbundenen Umwelt. Alle drei Paradigmen zusammen charakterisieren aber auch die Bereiche Umweltschutz, Personen- und Objektschutz, Rettungswesen (Sicherungswesen, Brand- und Explosionsschutz, Unfallrettungswesen, persönliche Schutzausrüstung, Katastrophen- und Zivilschutz, Evakuierung und anderes), Qualitätssicherung, Risk-Management, Datenschutz und das sicherheitstechnische Recht als sich stetig verändernde Festschreibung sicherheitstechnischer Erkenntnisse.

■ 1.2 Aktuelle Herausforderungen

Die aktuellen Herausforderungen im Kontext Sicherheit und Zuverlässigkeit sind durch die medienwirksamen Schlagwörter „Digitalisierung“, „Industrialisierung 4.0 und Robotik“ und in diesen eingebettet „Künstliche Intelligenz“ (KI) geprägt. Die industriellen Entwicklungen und Innovationen in diesen Bereichen bilden ein großes Potenzial zur Verbesserung der Lebensqualität in allen Bereichen des menschlichen Daseins. Damit eng verbunden sind jedoch auch vielfältige Fragestellungen im Sinne von Safety und Security.

Aufgabe von Staat, Gesellschaft, Industrie, Wissenschaft ist es, diese zu konstatieren und die mit jeder neuen technologischen Entwicklung verbundenen Gefahren, Risiken für den Einzelnen und die Gesellschaft zu bewerten, zu minimieren und in diesem Zusammenhang entsprechende gesetzliche Rahmenbedingungen zu schaffen.

Mit diesen Herausforderungen sind insbesondere die in Forschung und Entwicklung tätigen Wissenschaftler und Ingenieure konfrontiert. Hier gilt es neue Ansätze und Verfahren zur Bewertung, Validierung und Verifizierung der Sicherheit, Zuverlässigkeit, Verfügbarkeit und anderer Kategorien bereitzustellen.

Dabei bildet die sogenannte „Künstliche Intelligenz“ (KI) einen zentralen Baustein, der durch ständig wachsende Entwicklungen und Anwendungen im Kontext des maschinellen Lernens, bei Perzeption, Entscheiden, Handeln, Kommunikation etc. als „schwache KI“ geprägt ist.

Hierzu gehören aktuell unter anderem die Themenkomplexe

- Robotik, Steuerung, Regelung, Expertensysteme,
- autonomes Fahren,
- Sprach- und Textverarbeitung,
- Bild- und Spracherkennung,
- Mustererkennung
- etc.

Für den Bereich der Technischen Zuverlässigkeit gehören hierzu die Themen

- Zuverlässigkeitsplanung und -prüfung (allgemein),
- Predictive Maintenance,
- Fehlerdetektion,
- Ausfallratenbestimmung, Schwachstellenanalyse,
- Struktur- und Systemanalyse
- etc.

Heute wird KI bereits in vielen Branchen wie

- Industrie (Produktion, Fertigung, Qualitätssicherung, Steuerungs- und Regelungstechnik, Robotik und Vernetzung),
- Verkehrsträger (Schiene-, Luft-, Straßenverkehr und autonomes Fahren sowie deren Vernetzung),
- Energiewirtschaft (Steuerung und Optimierung, Netzbetrieb, Kraftwerksteuerung, prädiktive Wartung und Instandhaltung, Vertrieb etc.),
- Wirtschaft und Finanzwelt,
- Medien und Bildungsbereich (Textanalyse, digitale Assistenten, Crawling und Indexierung, Blockchain-Technologie, Video Content Marketing etc.),
- Medizin (Bildgebung, Diagnose, CT, Telemedizin, Neuroprothetik, Exoskelett etc.),
- Militär (autonome Waffensysteme, Strategieentwicklung, Entscheidungsfindung etc.),
- Jurisprudenz (Vertragsgestaltung und -prüfung, Recherche, Prozessvorbereitung, Handlungs- und Entscheidungsempfehlungen)
- etc.

erfolgreich eingesetzt.

Für das multi-fakultative Fachgebiet *Künstliche Intelligenz (KI)* gibt es, aufgrund des nicht fassbaren Begriffs der natürlichen Intelligenz im etymologischen Sinne, zahlreiche Definitionen. Prinzipiell verfolgt KI das Ziel die Attribute der kognitiven analogen Eigenschaften des Menschen durch eine digitale lernfähige Maschine nachzubilden.

Als Begründer der KI wird allgemein der Informatiker John McCarthy⁸ angesehen, der 1956 am Dartmouth College in Hanover (New Hampshire, USA) einen Workshop mit dem Titel „Dartmouth Summer Research Project on Artificial Intelligence“ organisierte. McCarthy verweist allerdings selbst auf Alan Turing⁹, einen der Gründungsväter der modernen Computertechnologie, der Informatik und damit der KI (Turing-Maschine, Turing-Test). Im deutschsprachigen Raum wird allgemein Karl Steinbuch¹⁰ als einer der Pioniere der Kybernetik (Begründer Norbert Wiener¹¹) und Künstlichen Intelligenz (maschinelles Lernen) und Namensgeber der Informatik als neue akademische Fachdisziplin, hervorgegangen aus der Kybernetik, angesehen (zur Geschichte der Entwicklung siehe Erhard Konrad, siehe Lit.).

Auch wenn die besonders in den letzten Jahren aufgrund verbesserter Rechentechnologien erzielten KI-Erfolge beeindruckend sind, ist der Weg zur sogenannten „starken KI“, d.h. der Adaption der gesamten kognitiven Fähigkeiten des Menschen im Sinne des Analogon Mensch-Maschine, noch durch viele Meilensteine gekennzeichnet. So schreibt Bernd Vowinkel in seinem Buch *Maschinen mit Bewusstsein* (siehe Lit.), Kapitel 1, Natürliche Intelligenz:

„Bei Fragen der Leistungsfähigkeit und der Eigenschaft von künstlicher Intelligenz wird in der Regel der Mensch als Maß aller Dinge zum Vergleich herangezogen. Das trifft insbesondere auf die Eigenschaften zu, die es bei der künstlichen Intelligenz noch nicht gibt, nämlich Bewusstsein, Vernunft und freier Wille.“

Ob es jemals eine „starke KI“ geben wird, ist strittig, obwohl ernst zu nehmende Wissenschaftler sich bereits mit dem nächsten Schritt einer „super KI“ und dessen Auswirkungen auf die Menschheit gedanklich auseinandersetzen. Das Gleiche gilt für neurobiologische Forschungen und die Prämissen zur Entwicklung einer „humanen Überintelligenz“ aufgrund von Quantenvorgängen im Gehirn. Unstrittig ist jedoch die Weiterentwicklung von KI und deren Verknüpfung mit der menschlichen Willenskraft (Nervensystem und Gehirnaktivität) zu einer gewissen „hybriden Intelligenz“, die es bereits heute ermöglicht, Schwerstbehinderten eine gewisse Beweglichkeit zurückzugeben, z.B. als Exoskelett-Hand (Neuroprothetik, neuronal gesteuerte Robotik). In diesem Kontext zeigt Bild 1.1 eine mögliche Illustration.

Auch sind sich Wissenschaftler darin einig, dass die Weiterentwicklung und Nutzung von Quantencomputern die KI stark beeinflussen wird (siehe Förderprogramm des BMBF zur Quantentechnologie).

⁸) John McCarthy (1927–2011)

⁹) Alan Turing (1912–1954)

¹⁰) Karl Steinbuch (1917–2005)

¹¹) Norbert Wiener (1894–1964)

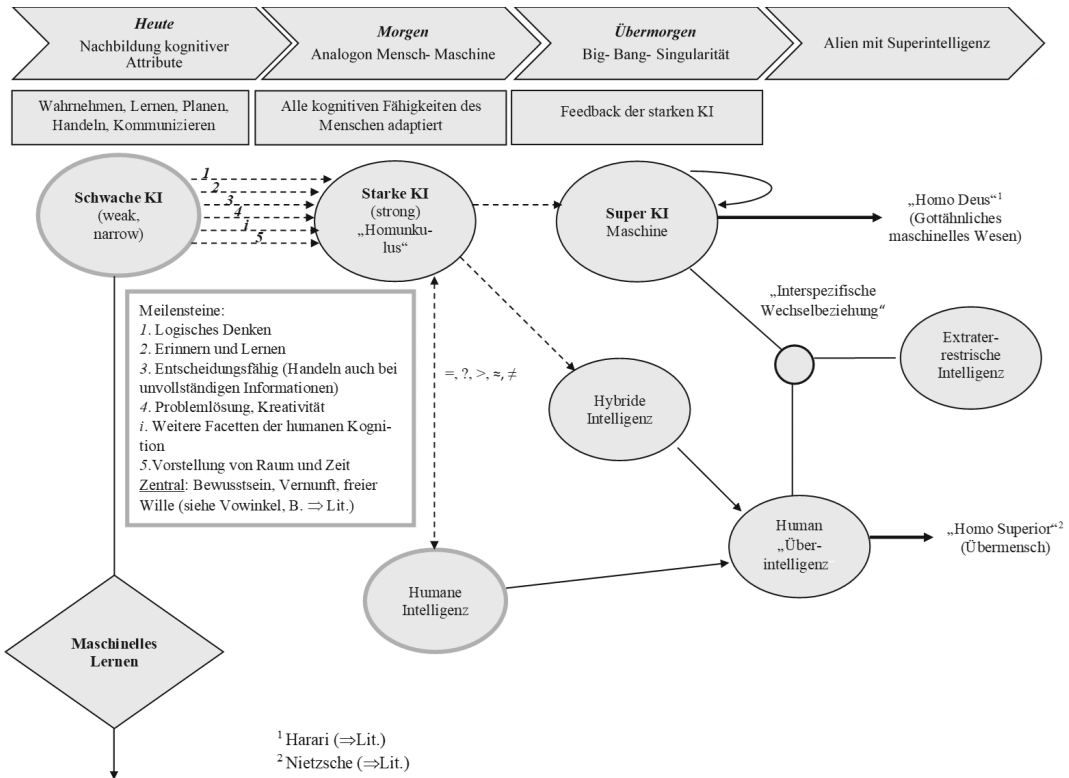


Bild 1.1 Evolution der Künstlichen Intelligenz

Dass digitale lernfähige Maschinen in vielen speziellen, determinierten Bereichen den intellektuellen Fähigkeiten eines Menschen überlegen sind, ist unstrittig. So gewann der IBM Computer Deep Blue 1997 gegen den damaligen Schachweltmeister Garry Kasparov mit 3,5 zu 2,5 Punkten (d.h., von sechs Partien hat Kasparov zwei verloren, eine gewonnen, drei endeten mit Remis).

Doch zurück zur „schwachen KI“: Der Kernbereich der „schwachen KI“ ist durch das *maschinelle Lernen* geprägt. Hier wurden in den letzten Jahren aufgrund der immensen Weiterentwicklung der Computertechnologie, verbunden mit einer enormen Steigerung der Rechenleistung durch entsprechende Prozessoren und Vernetzung, mit der Möglichkeit sehr große Datenmengen auch in Echtzeit sehr schnell zu verarbeiten, enorme Fortschritte erzielt. *Machine Learning* ist durch lernfähige Algorithmen wie *neuronale Netze* und *Deep Learning* gekennzeichnet. Im Gegensatz hierzu nutzt *Data Mining* statistische Verfahren mit dem Ziel aus den Daten (auch bei einer geringen Anzahl) allgemeine Zusammenhänge und Wissen zu generieren. Bild 1.2 zeigt hierzu einen groben Überblick hinsichtlich *Machine Learning* und *Deep Learning* als inklusive Ausprägung sowie *Data Mining* und die hierzu genutzten algorithmischen Verfahren.

„Niemand formuliert es so, aber für mich ist KI fast eine Geisteswissenschaft. Es ist wirklich der Versuch, menschliche Intelligenz und Kognition zu verstehen.“

Sebastian Thrun (CEO bei Kitty Hawk Coporation and Innovation)

Für weitere Ausführungen hierzu siehe Kapitel 22.

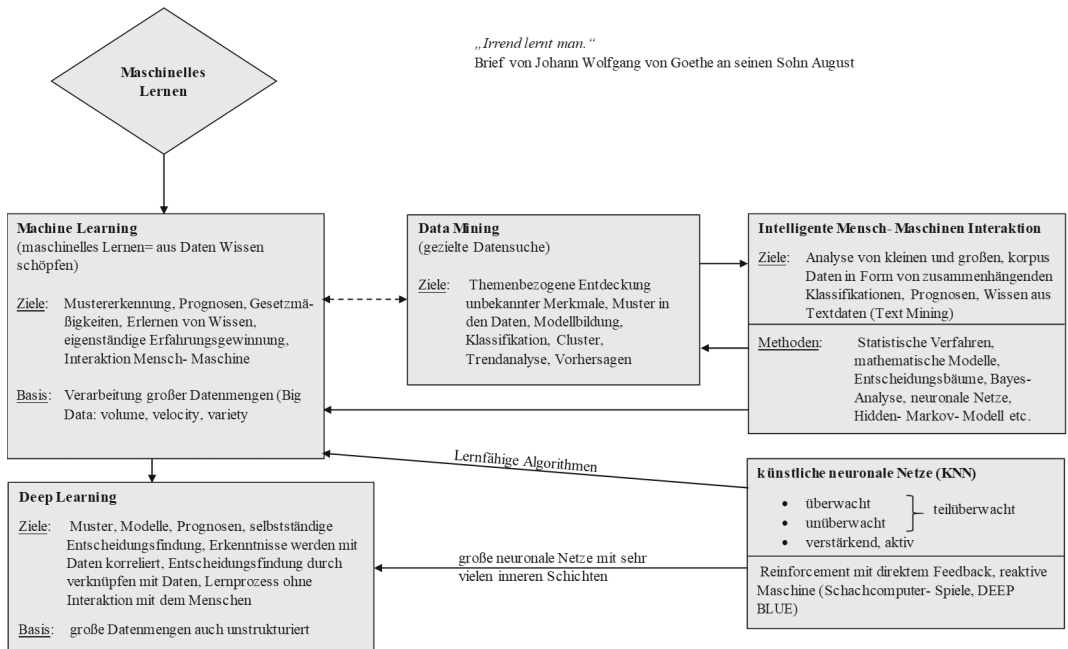


Bild 1.2 Gegenwärtige Praxis der Künstlichen Intelligenz

Bei aller Euphorie bezüglich der Weiterentwicklung und stetig wachsenden Einsatzbereiche von KI bleibt ein zentrales Thema die Absicherung, d.h. die qualitative und quantitative Bewertung hinsichtlich Sicherheit, Zuverlässigkeit, Verfügbarkeit der durch das maschinelle Lernen geprägten sicherheitskritischen Anwendungsbereiche, wie z.B. Mensch-Roboter-Interaktion, autonomes Fahren, aufgrund der nicht mehr nachvollziehbaren Output-Ergebnisse der zugrunde gelegten und implementierten neuronalen Netze allgemein und insbesondere von *Deep Neural Networks*. Ein Ziel ist es deshalb, neuronale Netze erklärbar zu gestalten. Unter dem Terminus „*Explainable Artificial Intelligence*“, kurz *Explainable AI*, werden hierzu gegenwärtig weltweit Förderprojekte und Forschungen durchgeführt (siehe hierzu unter anderem das Fraunhofer-Institut für Kognitive Systeme, IKS).

Des Weiteren spielt die Qualität der zugrunde gelegten Input-Daten eine entscheidende Rolle. Eine sorgfältige Dateninterpretation und -verarbeitung gelernter Trainingsdaten, aber auch ungelernter Daten als Modellinput ist deshalb unerlässlich für die Robustheit der Modellierung. Diesbezügliche Fehler, insbesondere bei sicherheitskritischen Anwendungen, können naturgemäß mit weitreichenden Auswirkungen und Folgen verbunden sein. Für weitere Ausführungen hierzu siehe Kapitel 6.

Es liegt auf der Hand, dass der Einsatz von KI insbesondere für automatisierte Systeme wie z.B. in automatisierten Prozessen und bei der Nutzung von Robotern sowie bei deren Vernetzung in der Produktion (Industrie 4.0) oder beim autonomen Fahren, aber auch im medizinischen Bereich (Medizintechnik), mit einem stringenten Sicherheitsnachweis verbunden ist. So sind beispielsweise die prinzipiellen Dilemmata beim autonomen Fahren zwar weitgehend bekannt (Meyna/Heinrich, siehe Lit.), allerdings ab Level 3 (Hochautoma-

tisiert) aufgrund der erforderlichen Absicherungen allgemein – von Forschungsfahrzeugen bis Level 5 abgesehen – noch nicht verfügbar. Es ist deshalb äußerst anspruchsvoll, geeignete Strategien wie Fail-Operational Systems für die Behandlung von komplexen Fehlerzuständen wie z.B. bei der Perzeption, der Fahrzeugregelung, der implementierten Hard- und Software, ungewollter oder gewollter Manipulation im Kontext des dynamischen, adaptiven Fehlermanagements inklusive einer Rückfallebene zur Degradation zu entwickeln und diese qualitativ und quantitativ zu bewerten. Nicht zuletzt gilt es, eine systembezogene Resilienz und Verfügbarkeit in Echtzeitanforderungen umzusetzen und zu gewährleisten (Forschungen hierzu siehe unter anderem Fraunhofer IKS und IQZ).

Nach wie vor ungeklärt sind auch die rechtlichen Fragestellungen und Probleme bei den Lern- und Trainingsprozessen der Software im Kontext der damit verbundenen „autonomen“ Handlungs- und Entscheidungsprozesse (Ambivalenz Hersteller vs. Anwender), siehe hierzu unter anderem die juristische Dissertation von Justin Grapentin (siehe Lit.).

Mit dem zunehmenden Einsatz von KI und deren Weiterentwicklung in vielen Bereichen der humanen Existenz spielen auch die damit verbundenen ethischen Fragestellungen eine immer größere Rolle. In diesem Zusammenhang sei auf die in jeder Hinsicht lesenswerte kleine Monographie von Christoph Bartneck et al. des Carl Hanser Verlags (siehe Lit.) verwiesen. Empfehlenswert ist außerdem das Buch „Maschinelles Lernen – Grundlagen und Algorithmen in Python“ von Jörg Frochte (2. Auflage 2019, Carl Hanser Verlag, München).

„KI ist wahrscheinlich das Beste oder das Schlimmste, was der Menschheit passieren kann.“

Stephen Hawking (1942 – 2018)

2

Zuverlässigkeit bei Haftungs- und Gewährleistungsfragen

Kürzere Entwicklungszeiten, hohe Systemkomplexität, verlängerte Gewährleistungszeiträume – das sind die Schlagworte, aus denen sich die Herausforderungen für Hersteller und Zulieferer komplexer Systeme ergeben und denen es sich in heutiger Zeit zu stellen gilt. Zudem zeigt sich fast täglich im Rahmen öffentlichkeitswirksamer Rückrufe, dass nicht alle Systeme vor dem Inverkehrbringen die notwendige Reife hinsichtlich Sicherheit und Zuverlässigkeit erreichen. Verschärft wird diese Thematik durch den Einsatz von Gleichteil- und Plattformstrategien, die auf den ersten Blick durch erhebliche Skaleneffekte attraktiv erscheinen, jedoch hinsichtlich des Risikopotenzials von großen Serienschäden einer ganzheitlichen Risikobetrachtung bedürfen.

Eine weitere Herausforderung stellen neue Geschäftsmodelle dar, bei denen die klassischen Original Equipment Manufacturer (OEM) sich immer mehr zum Plattformanbieter entwickeln und statt einem Produkt eine Dienstleistung verkaufen. Damit übernehmen sie in Teilen das Betreiberisiko für den Kunden und versuchen dieses Risiko entlang der Lieferkette zu diversifizieren. Die Lieferanten spüren diese Entwicklung in Form von steigenden Gewährleistungszeiten sowie erhöhten Anforderungen hinsichtlich der Lebensdauer. Dabei muss berücksichtigt werden, dass viele Zusagen der Lieferanten weit über den regulatorischen Rahmen der gesetzlichen Sachmängelhaftung hinausgehen und folglich auch der Versicherungsrahmen neu bewertet werden kann. In diesem Kapitel sollen die rechtlichen Grundlagen einfach und verständlich dargestellt werden.

■ 2.1 Haftungsgrundlage

Grundsätzlich sind für einen Hersteller von Produkten bzw. jemanden, der solche Produkte in den Verkehr bringt, Haftungsansprüche aus dem Zivilrecht, dem öffentlichem Recht, aber auch aus dem Strafrecht denkbar. Haftungsansprüche aus dem öffentlichen Recht sowie dem Strafrecht sind im alltäglichen Bereich jedoch nicht die Regel.

Zuverlässigkeitstechnische Einflüsse auf die Haftung finden sich vornehmlich im Zivilrecht wieder, welches noch einmal die vertragliche Haftung (Gewährleistung) sowie die außervertragliche Haftung unterscheidet. Diese Bereiche decken nicht sämtliche erdenklichen Risiken im Zusammenhang mit fehlerhaften Produkten ab, umfassen jedoch einen großen Teil der für Hersteller, „Quasi-Hersteller“ und Importeure praxisrelevanten Rechtsbereiche.

2.1.1 Außervertragliche Haftung

Bezüglich der außervertraglichen Haftung sind besonders der § 823 Bürgerliches Gesetzbuch (BGB) (siehe Lit.) sowie der § 1 Produkthaftungsgesetz (ProdHaftG) (siehe Lit.) zu nennen (Bild 2.1).

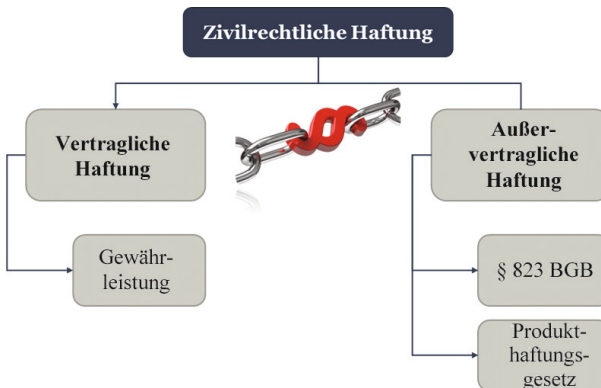


Bild 2.1 Zivilrechtliche Haftung

Eine Haftung nach § 1 ProdHaftG beinhaltet/bedeutet eine verschuldensunabhängige Haftung. Es kommt hier also nicht auf eine schuldhaftes Handeln des Herstellers des Produkts an. Im Rahmen dieser Gefährdungshaftung kann ein Unternehmen auch für Ausreißer in der Produktion haften, wenn dadurch jemand verletzt oder eine Sache beschädigt wird. Verursacht das fehlerhafte Produkt einen Schaden an einer anderen Sache, liegt ein Sachschaden vor. Für solche Sachschäden haftet der Hersteller aber nur, wenn die beschädigte Sache üblicherweise dem privaten Gebrauch oder privaten Verbrauch dient und auch hauptsächlich so verwendet wird. Durch § 1 ProdHaftG ist somit der Verbraucherschutz für Endkunden gewährleistet. Eine Haftung ist nur dann ausgeschlossen, wenn der Fehler des Produktes zum Zeitpunkt des Inverkehrbringens nach dem Stand von Wissenschaft und Technik nicht zu erkennen war. Ein solcher Nachweis ist für komplexe, technische Produkte extrem schwer. Um hingegen einen Haftungsanspruch gemäß § 823 BGB geltend zu machen, muss dem Verursacher ein Verschulden nachgewiesen werden. Er muss also vorsätzlich bzw. fahrlässig gehandelt haben. Kann ihm dies nicht nachgewiesen werden, kann ein Haftungsanspruch abgewehrt werden. Bedeutsam für diesen Punkt aus zuverlässigkeitstechnischer Sicht sind z.B. das Einhalten von Normen und Standards, das Arbeiten nach dem Stand der Technik sowie eine gute Dokumentation über den gesamten Produktlebenszyklus. Auch der Themenkomplex „Erprobung“ spielt hier eine bedeutende Rolle. Bei der Haftung nach § 823 BGB spricht man auch von der Produzentenhaftung. Diese erstreckt sich auf den gesamten Lebenszyklus des Produkts und umfasst im Wesentlichen die Phasen Design/Konstruktion, Fabrikation, Instruktion sowie Produktbeobachtung. In all diesen Phasen sind zuverlässigkeitstechnische Tätigkeiten ein Garant dafür, den Beweis führen zu können, dass den Hersteller kein Verschulden trifft.

Aus § 823 BGB lässt sich somit auch die Produktbeobachtungspflicht ableiten, welche fälschlicherweise oftmals aus dem Produkthaftungsgesetz hergeleitet wird. Der Hersteller

muss systematisch und kontinuierlich prüfen, wie sich seine Produkte in der Praxis bewähren, und er muss etwaigen Auffälligkeiten nachgehen. Dabei erfolgt eine Unterteilung in eine aktive sowie passive Marktbeobachtung. Aufgaben im Rahmen der Marktbeobachtung sind z. B. die Analyse von Schadteilen, das Screening und Auswerten von Feldbeanstandungen (Massendaten), der aktive Rückkauf von Systemen und Bauteilen aus dem Feld zur anschließenden Befundung, aber auch das kontinuierliche Überwachen von einschlägigen Internetforen oder Social Media. Besonders der letzte Punkt hat sich in den letzten Jahren als probates Mittel erwiesen, um ein schnelles Feedback aus dem Feld zu bekommen. Zwar ersetzen diese Informationen keine professionelle Analyse von Schadteilen oder Feldbeanstandungen, jedoch lassen sich damit weltweit Diskussionen über Fehler oder Enttäuschungen der Kunden bezüglich eines Produkts überwachen. Die Produktbeobachtungspflicht umfasst weiterhin die Verpflichtung des Herstellers, alles zu tun, was ihm den Umständen nach zugemutet werden kann, um konkrete Gefahren für Leib und Leben abzuwenden. Sollte von einem Produkt eine solche konkrete Gefahr ausgehen, so kann der Hersteller dazu verpflichtet sein, das Produkt zurückzurufen. Rückruf bedeutet damit in letzter Konsequenz nicht, dass das Produkt auch physisch zurückgeschickt oder repariert werden muss. Man versteht darunter auch die Warnung des Endkunden oder gegebenenfalls die Aufforderung, das Produkt nicht weiter zu nutzen. Vgl. hierzu das Pflegebettenurteil des BGH aus dem Jahr 2008 (BGH, Urt. v. 16. 12. 2008 – VI ZR 170/07) und das Urteil des BGHZ aus dem Jahr 2009 (BGHZ 179, 157 = BB 2009, 627); dazu Reusch, Pflichtenkreis von Unternehmen im Umgang mit unsicheren Produkten – Thesen zum Produktrückruf, BB 2017, 2248 ff. (siehe Lit.).

2.1.2 Vertragliche Haftung

Im Rahmen der vertraglichen Haftung muss zuerst einmal in die Begriffe „Garantie“ sowie „Gewährleistung“ unterschieden werden. Bei einer „Garantie“ handelt es sich um eine freiwillige Leistung eines Herstellers oder Händlers für sein Produkt. Dabei wird i. d. R. ein besonderes Haltbarkeits- oder Funktionsversprechen gegeben, welches über die Anforderungen der gesetzlich geregelten Gewährleistung hinausgeht. Manchmal betrifft die Garantie nur einen bestimmten Teil des Gesamtproduktes oder es werden spezifische Teile ausgenommen. Garantien sind frei gestaltbar, sofern sie geltendes Recht nicht verletzen. Bei einer Garantie spielt der Zustand der Ware zum Zeitpunkt der Übergabe an den Kunden keine Rolle, da die Funktionsfähigkeit für den gesamten Zeitraum garantiert wird. Die Garantie ist folglich ein optionales Versprechen und keine Verpflichtung.

Die Gewährleistung, auch gesetzliche Sachmängelhaftung, fällt unter den Bereich der vertraglichen Haftung und muss durch den Verkäufer verpflichtend gewährt werden. Die Rechte und Pflichten sind in § 433 BGB (Kaufvertrag), § 434 BGB (Sachmangel) und § 437 ff. (Rechte des Käufers bei Mängeln) geregelt. Damit eine Haftung eintreten kann, bedarf es eines Kaufvertrags (z. B. Lieferabrufe), eines Sach- oder Rechtsmangels beim Zeitpunkt des Gefahrübergangs sowie des Nachweises des Käufers, dass ein Mangel vorliegt. Zudem muss der Mangel innerhalb der Gewährleistungsfrist (gesetzlich Business-to-Consumer (B2C) – Bereich zwei Jahre, vertraglich im Business-to-Business (B2B) – Bereich oft deutlich höher) eintreten und vor Ablauf der Verjährungsfrist gemeldet werden.

Hinsichtlich technischer Spezifikationen ist vor allem die Definition des Sachmangels von großer Bedeutung. Es gilt nach § 434 BGB:

„Eine Sache ist frei von Sachmängeln, wenn sie bei Gefahrübergang die vereinbarte Beschaffenheit hat.“

Falls keine Beschaffenheit vereinbart wurde, gilt:

„sich für die im Vertrag vorausgesetzte Verwendung eignet.“

Ansonsten gilt:

„sich zur gewöhnlichen Verwendung eignet und eine Beschaffenheit hat, die für Sachen gleicher Art und Güte üblich ist und die der Käufer hätte erwarten dürfen.“

Um sich langwierige und nervenaufreibende Diskussionen und Verhandlungen im Rahmen einer Gewährleistungsauseinandersetzung zu sparen, sollte ein besonderes Augenmerk auf die Ausgestaltung von Lastenheften und technischen Spezifikationen gelegt werden. Somit lässt sich der Sachmangel im Haftungsfall einfacher feststellen oder ablehnen.

Tritt ein Sachmangel im Rahmen der Gewährleistungsfrist auf, so hat der Anspruchsteller ein Recht auf Nacherfüllung. Dies kann durch eine Reparatur oder eine Neubelieferung erfüllt werden. Kommt der Verkäufer seinen Pflichten nicht nach, so hat der Käufer die Möglichkeit, vom Kaufvertrag zurückzutreten, den Kaufpreis zu mindern oder Aufwendungsersatz bzw. Schadensersatz zu fordern. Hat der Käufer durch den Mangel weitere Einbußen (z. B. Betriebsausfall), kann er zusätzlich Schadensersatz geltend machen. Im B2B-Bereich sei darauf hingewiesen, dass nach Änderung des BGB zum 01.01.2018 auch die Ein- und Ausbaurkosten für mangelhafte Produkte geltend gemacht werden können. Vor dieser Regelung war dieser Punkt Streitgegenstand zahlreicher Verfahren (z. B. Parketturteil des BGH aus dem Jahr 2008).

2.1.3 Stand der Technik

Bei vielen haftungsrelevanten Fragestellungen wird immer wieder der Begriff „Stand der Technik“ herangezogen, welcher oftmals falsch interpretiert wird. Für sicherheitsrelevante Technologien bietet es sich an, nach der Definition zu arbeiten, welche höchstrichterlich anerkannt wurde. Beispielhaft kann dazu das Urteil des Bundesverfassungsgerichts (BVerfG) aus dem Jahr 1978 zum Atomgesetz genannt werden (Kalkar, Urteil vom 08.08.1978 – 2 BvL8/77). Dort wurde in die drei Stufen

- allgemein anerkannte Regeln der Technik,
- Stand der Technik und
- Stand von Wissenschaft und Technik

unterschieden.

Unter den „Anerkannten Regeln der Technik“ versteht man den Entwicklungsstand, den die meisten Fachleute für richtig erachten. Dieser ist i. d. R. in Normen, Richtlinien und Fachbüchern dargelegt. Er stellt das Mindestmaß an Sicherheit für ein technisches Produkt dar.

Den „Stand der Technik“ sieht das BVerfG als die Front der technischen Entwicklung, welche jedoch praktisch erprobt und bewährt ist. Die Europäische Union beschreibt den Stand

der Technik auch als „beste verfügbare Technik“. Der Stand der Technik geht folglich i. d. R. über die Anforderungen von Normen und Richtlinien hinaus, da diese aufgrund ihrer langen Bearbeitungszeiträume häufig schon bei der Veröffentlichung nicht mehr die „Front der technischen Entwicklung“ darstellen. Zudem sind Normen und Richtlinien ein Kompromiss der daran beteiligten Akteure, sodass der darin festgelegte Stand immer mit dem eigenen Produkt abgeglichen werden muss. Der Stand der Technik ist in zahlreichen Normen und Richtlinien als Entwicklungsziel aufgeführt. Versicherungstechnisch ist der Stand der Technik von großer Bedeutung, da er in deckungsschädlichen Klauseln, wie z. B. der Erprobungsklausel (Steinkühler, siehe Lit.), aufgeführt wird.

Diese lautet wie folgt:

„Ausgeschlossen vom Versicherungsschutz sind Ansprüche aus Sach- und Vermögensschäden durch Erzeugnisse, deren Verwendung oder Wirkung in Hinblick auf den konkreten Verwendungszweck nicht nach dem Stand der Technik oder in sonstiger Weise ausreichend erprobt waren.“

Eine Erprobung, die sich nur an Normen und Richtlinien orientiert, kann gegebenenfalls nicht den Stand der Technik darstellen und dazu führen, dass ein Verlust der Deckung droht.

Beim „Stand von Wissenschaft und Technik“ wird vorausgesetzt, dass man mögliche Gefahren schon erkennen kann, obwohl sie von der Technik noch nicht vollständig beherrscht werden. Es geht um das, was führende Fachleute in Wissenschaft und Technik aufgrund neuester wissenschaftlicher Technik für erforderlich halten, um höchsten Gefahren in der Praxis vorzubeugen. Dies umfasst wissenschaftliche Veröffentlichungen, Kongressbeiträge, Schutzrechtsanmeldungen etc.

Für weitere Ausführungen hierzu siehe Abschnitt 3.1.

2.1.4 Gewährleistungsmanagement zwischen Unternehmen

Die Gewährleistungsabwicklung zwischen Geschäftskunden (Business-to-Business, B2B) weicht z. T. deutlich von den Regelungen der gesetzlichen Sachmängelhaftung ab. Besonders im Automobilbereich wurden in den letzten Jahren Gewährleistungsvereinbarungen und -prozesse umgesetzt, die zahlreiche Rechte und Pflichten aus der gesetzlichen Regelung aushebeln oder die Verantwortlichkeiten tauschen. Initial sollten diese Abweichungen zur Vereinfachung der Gewährleistungsprozesse sowie zur Kosteneinsparung dienen. Deshalb werden z. B. Schadteile nur aus sehr begrenzten Regionen an den OEM zurückgesendet (Sendemärkte) und auch von dieser Stichprobe wird nur ein gewisser Teil zur Befundung an den Lieferanten geschickt. Dies macht grundsätzlich Sinn, da ansonsten erhebliche Logistik- und Prüfkosten anfallen würden. Der Lieferant befundet dann die Stichprobe und gibt an, welcher Anteil in seiner Verantwortung liegt. Die Regressierung erfolgt danach entweder über das sogenannte Referenzmarktprinzip, das Anerkennungsverfahren oder einen hybriden Ansatz. Um den Prozess weiter zu vereinfachen, fassen die meisten Hersteller die Produkte des Lieferanten in Warenkörben oder Produktgruppen zusammen.

Beim Referenzmarktprinzip schickt der Hersteller eine repräsentative Stichprobe an Schadteilen aus einem oder mehreren Sendemärkten an den Lieferanten zu Befundung. Stellt dieser einen Fehler am Bauteil fest und ist dieser Fehler durch ihn verursacht, so belastet

der Hersteller das Lieferantenkonto mit einem Regressbetrag, der sich aus den Teilekosten, den Kosten für den Ein- und Ausbau und den Logistikkosten multipliziert mit dem Verhältnis aus zurückgesendeten Teilen zu weltweit ausgefallenen Teilen berechnet. Mit dieser Berechnungslogik kann der Regressbetrag für ein Bauteil, welches für wenige Euro verkauft wird, bei mehreren hundert oder sogar tausend Euro liegen.

Das Beispiel in Bild 2.2 zeigt auf, dass für ein Bauteil, welches der Zulieferer für 20 € verkauft, ein Regressbetrag von 750 € fällig wird.

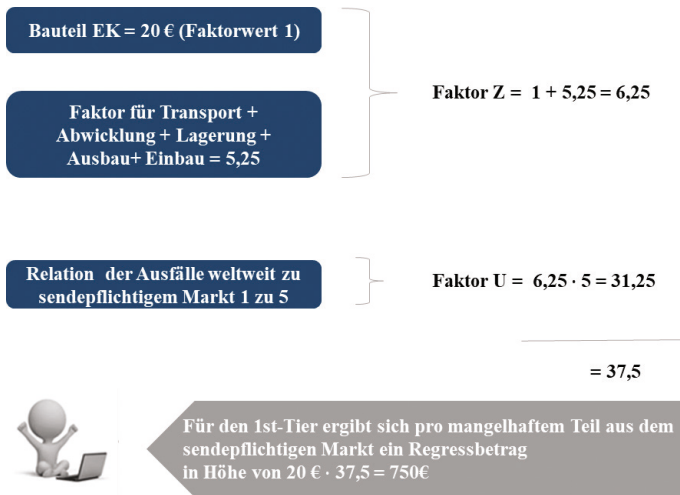
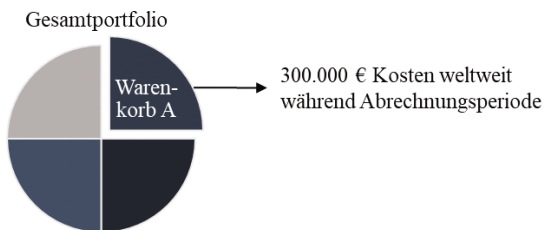


Bild 2.2

Kostenentwicklung bei Referenzmarktp Prinzip

Beim Anerkennungsverfahren wird eine repräsentative Stichprobe an Schadteilen gesammelt, welche der Lieferant befundet. Aus diesen Teilen wird ein Technischer Faktor bzw. eine Verantwortungsquote gebildet, welche dann i.d.R. für ein Jahr gilt. Der Hersteller überprüft dann unterjährig die weltweiten Feldbeanstandungen, ermittelt dafür die Kosten, multipliziert diese dann mit dem Technischen Faktor und belastet danach den Lieferanten (Bild 2.3).



- Überprüfung der „repräsentativen“ Menge (20 Teile) an Teilen bringt 12 (tatsächlich) mangelhafte Teile → Quote 12/20
- Daher: Anerkennungsquote = 60%
- Gewährleistungskosten des Lieferanten für Produktgruppe A im abgerechneten Zeitraum = $0,6 \cdot 300.000 \text{ €} = 180.000 \text{ €}$

Bild 2.3

Kostenentwicklung bei Anerkennungsverfahren

In beiden Verfahren werden häufig Kostenpauschalen eingesetzt, sodass grundsätzlich für den Einzelfall weder der konkrete Schadensgrund noch die tatsächliche Schadenshöhe ermittelt werden kann. Der Zulieferer verzichtet damit auf elementare Rechte der gesetzlich geregelten Gewährleistung, sodass hier besonders auf einen abgestimmten Versicherungsschutz geachtet werden muss, da sich Standardpolicen i. d. R. an den gesetzlichen Regelungen orientieren.

Aufgabe eines ganzheitlichen Zuverlässigkeitsmanagements ist es, die im Rahmen des Gewährleistungsprozesses anfallenden Daten zu analysieren, Muster zu erkennen und Hinweise auf unplausible Abrechnungen zu geben.

Neben der vorangehend genannten Abrechnung im Regelregress haben – zumindest im Automobilbereich – alle OEM entsprechende Klauseln, die bei vermehrt auftretenden Fehlern einen sogenannten Sonderregress ermöglichen. Wird dieser ausgerufen, so erfolgt die Regressierung i. d. R. im Rahmen eines kleinen Projektes. Oft sind Sonderregresse mit Feldaktionen verbunden, in denen der OEM nicht sicherheitsrelevante Fehler durch eine groß angelegte Tauschaktion verhindern möchte. Die Kosten für eine solche Aktion belaufen sich schnell im Bereich von mehreren Millionen Euro und können aufgrund der Plattform- und Gleichteilstrategie auch in dreistellige Millionen oder sogar Milliarden reichen. Meist wird die Schadenssumme auf einen bestimmten Zeitraum prognostiziert und dann final verhandelt. Auch in diesem Bereich leisten Zuverlässigkeitsmethoden, wie z. B. Schichtlinien oder Felddatenprognosen, einen erheblichen Beitrag, um die eigene Verhandlungsposition zu stärken. Im Kapitel 28 finden sich einige Beispiele aus dem Haftungs- und Gewährleistungsbereich, welche die Bedeutung von Zuverlässigkeitsmethoden aufzeigen.

■ 2.2 Schadteilanalyse Feld und No-Trouble-Found-Prozess

Ein weiterer, wichtiger Baustein im Haftungs- und Gewährleistungsmanagement ist die Befundung von physisch vorliegenden Schadteilen. Schon Dorian Shainin¹⁾ sagte: „Talk to the parts; they are smarter than the engineers.“

Es geht dabei um die systematische Analyse von fehlerhaften Bauteilen aus dem Feld, um zum einen den Fehlerabstellprozess anzustoßen, und zum anderen die Verantwortlichkeit für den Fehler nach dem Verursacherprinzip zu ermitteln. Einer der neusten Standards im Bereich der Schadteilanalyse Feld ist der VDA Standard „Schadteilanalyse Feld & Auditstandard“ aus dem Jahr 2018. Anhand dieses Standards soll im Folgenden die Bedeutung im Rahmen des Haftungs- und Gewährleistungsmanagements, aber auch des Zuverlässigkeitsmanagements dargestellt werden.

Grundsätzlich ergibt sich in vielen Branchen das Problem, dass aufgrund von finanziellen, logistischen, aber auch zollrechtlichen Problemen nicht alle ausgefallenen Teile aus dem Feld als Schadteile an den Hersteller zurückgesendet werden können. Folglich muss sichergestellt werden, dass aus dieser möglichst repräsentativen Stichprobe der größte Mehrwert

¹⁾ Dorian Shainin (1914 – 2000)

gezogen werden kann. Es sind somit standardisierte, mit dem Kunden abgestimmte Prüfprozesse zu entwickeln, welche eine vertrauensvolle Zusammenarbeit zwischen Kunde und Lieferant ermöglichen. Das gemeinsam abgestimmte Konzept soll sicherstellen, dass mit wirtschaftlich vertretbarem Aufwand möglichst viele Fehler gefunden werden, da hinter jeder Reklamation ein potenziell enttäuschter Endkunde steht. Gleichzeitig soll das Konzept so aufgebaut sein, dass mit Standardprüfungen die relevanten Funktionen des betroffenen Systems oder Bauteils überprüft werden können. Tritt in dieser Standardprüfung ein Fehler auf und korreliert dieser mit der Fehlerbeschreibung, wird das Bauteil als „n.i.O. gemäß Befundung“ eingestuft und in den Problem-Lösungs-Prozess (auch Fehlerabstellprozess) übergeben. Sollte in der Standardprüfung kein Fehler zu finden sein, müssen die Schachteile in die Belastungsprüfung überführt werden. Wie der Name es schon sagt, sollte in der Belastungsprüfung versucht werden, Umgebungsbedingungen gemäß Spezifikationen zu simulieren, um die im Feld auftretenden Lasten (intern und extern) auf das System zu berücksichtigen. Wird in der Belastungsprüfung ein Fehler gefunden, erfolgt die weitere Bearbeitung im Problem-Lösungs-Prozess.

Bild 2.4 zeigt den Ablauf einer Schachteilanalyse mit den unterschiedlichen Prüfstufen.

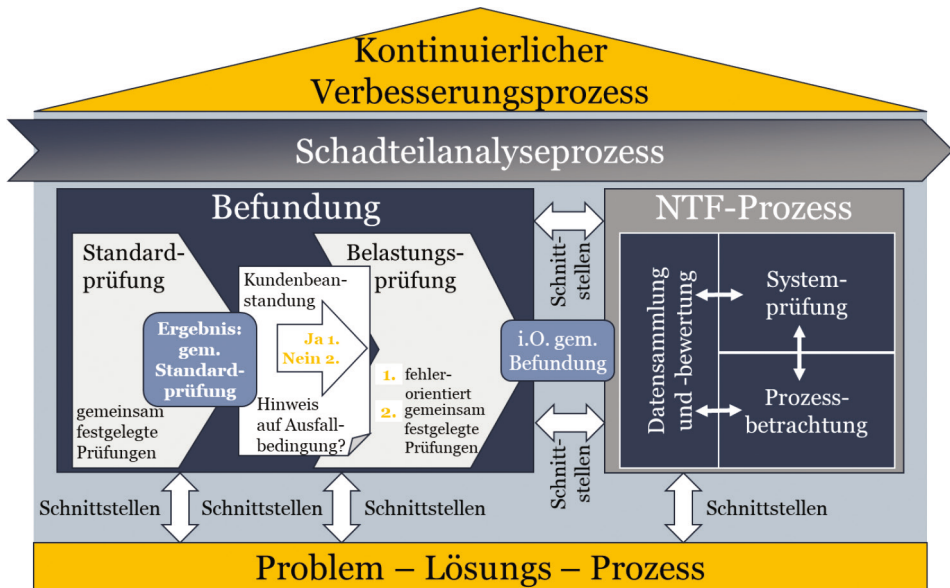


Bild 2.4 Ablauf einer Schachteilanalyse gemäß VDA

Wie aus Bild 2.4 hervorgeht, gibt es für Schachteile, welche auch nach der Belastungsprüfung ohne Befund sind, einen speziellen „No-Trouble-Found-Prozess“ (NTF). Dieser Prozess berücksichtigt die immer komplexeren und vernetzten eingesetzten Systeme im Fahrzeug. Oftmals kann nämlich ein Lieferant ohne zusätzliches Wissen von anderen Lieferanten oder seinem Kunden den Fehler nicht nachstellen, da weitere bzw. übergeordnete Systeme das Fehlerbild auslösen oder begünstigen. Um diese Fehler trotzdem finden und abstellen zu können, wurde der NTF-Prozess ins Leben gerufen. Er folgt nicht einem klar vorgegebenen

nen Prüfablauf, sondern wird individuell auf die Fragestellung hin ausgerichtet. Ein wichtiger Baustein des NTF-Prozesses ist der Bereich Datensammlung und Datenauswertung. Zahlreiche Verfahren aus diesem Buch zur Analyse von Daten eignen sich, um im NTF-Prozess wertvolle Hinweise zu Mustern oder Auffälligkeiten zu geben. Damit lassen sich kritische Herstellungsmonate eingrenzen, auffällige Märkte detektieren, organisatorische Einflüsse bei Beanstandungen visualisieren oder Prognosen zum zukünftigen Ausfallverhalten erstellen. Die mathematischen Methoden der Zuverlässigkeitstechnik liefern somit einen wertvollen Beitrag für ein erfolgreiches NTF-Projekt.

Bild 2.5 zeigt auf, wie vielfältig Auslöser für NTF-Probleme sein können. Es ist daher bedeutsam, zur Verfügung stehende Daten intensiv zu analysieren, um zeit- und ressourcenschonend das Problem zu finden.

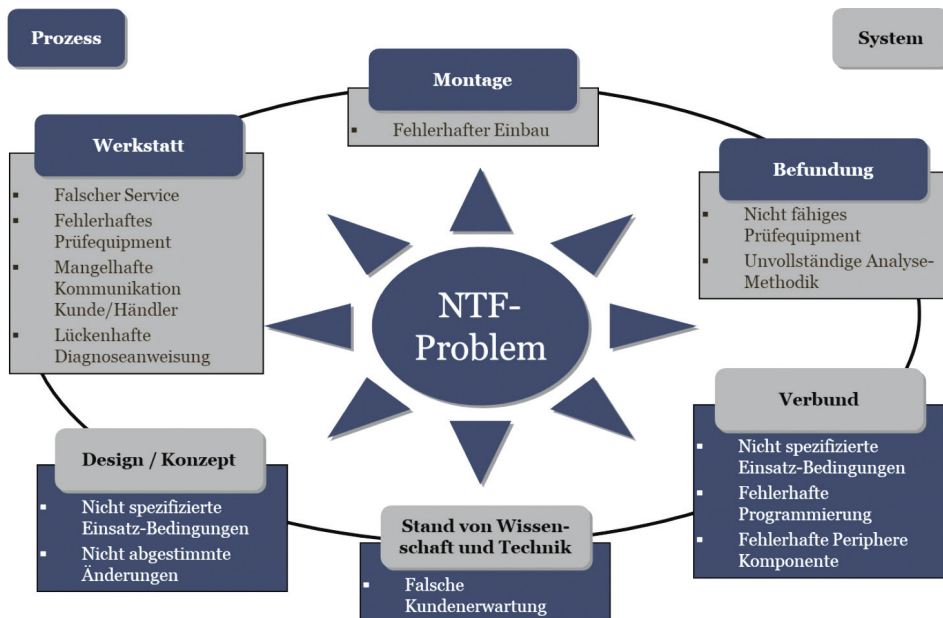


Bild 2.5 NTF-Problemlöser gemäß VDA

Für Unternehmen ist es notwendig, entsprechende Prozesse und Ressourcen in der Organisation vorzuhalten, welche die Umsetzung eines umfassenden Schadteilanalyseprozesses inklusive NTF-Prozess ermöglichen. Unternehmen aus dem Automobilbereich sind über die IATF 16949:2016 (International Automotive Task Force) sogar verpflichtet, einen entsprechenden Prozess vorzuhalten.

Normative Anforderungen in der Sicherheits- und Zuverlässigkeitstechnik

■ 3.1 Einleitung zu normativen Anforderungen im rechtlichen Kontext

Unter dem Begriff „Norm“ wird im allgemeinen Sprachgebrauch die Formulierung der „vereinbarten Art und Weise etwas zu tun“ verstanden. Seine Wortherkunft hat der Begriff im lateinischen Wort „norma“, was „Regel“, „Maßstab“ oder „Richtschnur“, aber auch „Winkelmaß“ bedeutet.

Der Begriff „Norm“ wird als Basis für eine Rechtsordnung weiterhin im Sinne von Gebot oder Verbot verstanden. Die Verletzung einer solchen Rechtsordnung stellt eine Rechtswidrigkeit dar und führt somit zu rechtlichen Sanktionen. In diesem Zusammenhang wird auch von „Rechtsnormen“ gesprochen, was für die vorliegende Betrachtung allerdings nicht weiter im Fokus steht.

Normen stellen Regeln und Leitlinien auf, die sich auf die Herstellung von Produkten, die Bereitstellung von Dienstleistungen oder auch die Verwaltung von Prozessen beziehen können. Sie sind somit eine relevante Quelle für Anforderungen an ein spezifisches Themengebiet. Die BSI Group (ein internationales Unternehmen für Normen und Zertifizierung) schreibt auf ihrer Website, dass *„Normen die destillierte Weisheit von Menschen mit Fachkenntnissen zu ihrem Thema sind, die die Bedürfnisse der Unternehmen, die sie vertreten, kennen“* (siehe BSI Group). Normen bündeln demnach den zum Erstellungszeitpunkt vorhandenen Wissensumfang zu einem konkreten Thema.

Folglich decken Normen vielfältige Themengebiete ab, wie z. B. das Bauwesen, den Arbeitsschutz und Gesundheitsschutz oder das Qualitätsmanagement. Für die Bereiche Sicherheits- und Zuverlässigkeitstechnik haben „technische Normen“ eine entsprechende Relevanz, worunter allgemein Normen im Sinne der festgelegten Vereinheitlichung von Abmessungen, Qualitäten, Herstellungsverfahren, Sicherheitsanforderungen und Bezeichnungen industrieller und gewerblicher Produkte verstanden werden (vgl. Schweizerischer Verband der Strassen- und Verkehrsfachleute (VSS, siehe Lit.). Es existiert hingegen keine einheitliche Definition des Begriffs „technische Norm“. Allerdings weisen technische Normen eine hohe Bedeutung in Bezug auf „Rationalisierung, Qualitätssicherung, Verständigung der am Wirtschaftsleben beteiligten Kreise, aber auch für die Sicherheit der Produkte der industriellen Massenfabrication“ auf (vgl. Urteil des Bundesgerichtshofs (BGH) vom 10.03.1987, siehe Lit.). Die Regelsatzung einer technischen Norm erfolgt durch privatrechtliche Organisationen, welche

alle an einem spezifischen Thema interessierten Personen und Gruppierungen (z.B. Hersteller, Forschungsinstitute, Behörden, Universitäten) zur gemeinsamen nationalen, europäischen oder internationalen Arbeit und Abstimmung an einen Tisch holen. Eine technische Norm ist laut Mohr (siehe Lit.) somit ein mit Konsens erstelltes Dokument, sodass die Gemeinschaftsarbeit der interessierten Kreise betont werden muss. Die Verabschiedung einer technischen Norm (im Gegensatz zu einem Standard) erfolgt in der Regel nach festgelegten Grundsätzen mithilfe eines öffentlichen Einspruchsverfahrens. Zu diesen Organisationen zählen u. a.

- DIN Deutsches Institut für Normung
- CEN Europäisches Komitee für Normung
(franz. Comité Européen de Normalisation)
- CENELEC Europäisches Komitee für elektrotechnische Normung
(franz. Comité Européen de Normalisation Électrotechnique)
- ETSI Europäisches Institut für Telekommunikationsnormen
(engl. European Telecommunications Standards Institute)
- IEC Internationale elektrotechnische Kommission
(engl. International Electrotechnical Commission)
- ISO Internationale Organisation für Normung
(engl. International Organisation for Standardization).

Der deutsche Alltag wird laut Angabe des DIN aktuell übrigens durch rund 34 000 Normen maßgeblich mit beeinflusst (siehe DIN 2023).

Obwohl der Normung, wie bereits angesprochen, eine hohe Bedeutung zugewiesen wird, gelten Normen als nicht zwingend, sondern sind – bis auf wenige Ausnahmen – auf freiwillige Anwendung ausgerichtete Empfehlungen (vgl. Urteil des BGHs vom 03.02.2004, siehe Lit.). Sie werden vielmehr zur Konkretisierung von unbestimmten Gesetzen herangezogen. So begründete z. B. das Bundesverwaltungsgericht in einem Beschluss aus dem Jahre 1996 (siehe Lit.):

„Das Deutsche Institut für Normung hat indes keine Rechtsetzungsbefugnisse. Es ist ein eingetragener Verein, der es sich zur satzungsgemäßen Aufgabe gemacht hat, auf ausschließlich gemeinnütziger Basis durch Gemeinschaftsarbeit der interessierten Kreise zum Nutzen der Allgemeinheit Normen zur Rationalisierung, Qualitätssicherung, Sicherheit und Verständigung aufzustellen und zu veröffentlichen. Wie weit er diesem Anspruch im Einzelfall gerecht wird, ist keine Rechtsfrage, sondern eine Frage der praktischen Tauglichkeit der Arbeitsergebnisse für den ihnen zugedachten Zweck.“

Weiter heißt es:

„Rechtliche Relevanz erlangen Normen im Bereich des technischen Sicherheitsrechts nicht, weil sie eigenständige Geltungskraft besitzen, sondern nur, soweit sie die Tatbestandsmerkmale von Regeln der Technik erfüllen, die der Gesetzgeber als solche in seinen Regelungswillen aufnimmt. Werden sie, wie dies beim Bau und beim Betrieb von Abwasseranlagen geschehen ist, vom Gesetzgeber rezipiert, so nehmen sie an der normativen Wirkung in der Weise teil, dass die materielle Rechtsvorschrift durch sie näher konkretisiert wird.“

Normen spiegeln gemäß dem Urteil des Bundesgerichtshofs von 2004 (siehe Lit.)

„den Stand der für die betroffenen Kreise geltenden anerkannten Regeln der Technik wider und sind somit zur Bestimmung des nach der Verkehrsauffassung zur Sicherheit Gebotenen in besonderer Weise geeignet“.

Weiter hat der Bundesgerichtshof im sogenannten Airbag-Urteil von 2009 festgehalten, dass der Hersteller zur Gewährleistung der erforderlichen Produktsicherheit diejenigen Maßnahmen zu treffen hat, *„die nach den Gegebenheiten des konkreten Falles zur Vermeidung einer Gefahr objektiv erforderlich und nach objektiven Maßstäben zumutbar sind“* (vgl. Urteil des BGHs vom 16.06.2009, siehe Lit.). Darin heißt es weiter, dass Sicherheitsmaßnahmen erforderlich sind, *„die nach dem im Zeitpunkt des Inverkehrbringens des Produkts vorhandenen neusten Stand der Wissenschaft und Technik konstruktiv möglich sind“*.

In zuvor erwähnten sowie weiteren gerichtlichen Beschlüssen bzw. Urteilen tauchen sogenannte „technische Generalklauseln“ (oder auch Technik Klauseln) auf, deren Verhältnis zur Normung durchaus unterschiedlich ist. Gemeint sind nach der Kommission Arbeitsschutz und Normung (KAN) (siehe Lit.) mit den technischen Generalklauseln

- die anerkannten Regeln der Technik,
- der Stand der Technik und
- der Stand von Wissenschaft und Technik,

welche nicht allesamt gesetzlich definiert sind.

Unter den **anerkannten Regeln der Technik** werden Prinzipien und Lösungen verstanden, die nach dem Beschluss des Bundesverwaltungsgerichtes aus dem Jahr 1996 (BVerwG) (siehe Lit.) *„in der Praxis erprobt und bewährt sind“*. Zahlreiche Gerichtsurteile geben an, dass Normen die Vermutung in sich tragen, dass sie den Stand der allgemein anerkannten Regeln der Technik wiedergeben (vgl. Urteil des BGHs vom 24.05.2013, siehe Lit.). Normen haben allerdings keinen Ausschließlichkeitsanspruch und können hinter den anerkannten Regeln der Technik zurückbleiben. Die „Regeln der Technik“ sind im Übrigen nicht mit „technischen Regeln“ zu verwechseln, wobei es sich nach Völkel (siehe Lit.) um einen anderen Ausdruck für das *„Verhalten auf einem technischen Gebiet handelt“*. Eine technische Regel ist beispielsweise die gängige Praxis, dass Schrauben rechtsherum eingedreht und linksherum ausgedreht werden. Es muss allerdings festgehalten werden, dass die beiden Begriffe eng miteinander verbunden sind, da beide Begriffe Verhaltens- bzw. Vorgehensweisen auf technischem Gebiet beschreiben, wobei die technischen Regeln dem Tatsachenbereich zuzuordnen sind, wohingegen es sich bei den Regeln der Technik um ein rechtliches Phänomen handelt. Alle Regeln der Technik sind auch technische Regeln, nicht alle technischen Regeln sind aber auch Regeln der Technik.

Die Betriebssicherheitsverordnung definiert in § 2 Abs. 10 den **Stand der Technik (SdT)** als

„Entwicklungsstand fortschrittlicher Verfahren, Einrichtungen oder Betriebsweisen, der die praktische Eignung einer Maßnahme oder Vorgehensweise zum Schutz der Gesundheit und zur Sicherheit der Beschäftigten oder anderer Personen gesichert erscheinen lässt“ (siehe Lit.).

Ferner sind *„bei der Bestimmung des Stands der Technik [...] insbesondere vergleichbare Verfahren, Einrichtungen oder Betriebsweisen heranzuziehen, die mit Erfolg in der Praxis erprobt worden sind“*.

Normen sind gemäß der KAN nur dann verbindlich, wenn sie den SdT in einem Codex festlegen. Harmonisierte Normen bieten einen guten Anhaltspunkt für den Stand der Technik im Zeitpunkt der Bekanntmachung. Harmonisierte europäische Normen werden im Amtsblatt der Europäischen Union veröffentlicht. Eine Norm ist allerdings nicht mit dem SdT gleichzusetzen, sie kann immer nur ein Ausdruck des Stands der Technik sein. Normen veralten ständig – dabei wird auf die Tatsache, dass einzelne Normeninhalte seit ihrem Entwicklungszeitpunkt über die Abstimmungs- und Genehmigungsdauer bis hin zur Publikation bereits wieder überholt sein können, an dieser Stelle gar nicht weiter eingegangen. Es sollte folglich regelmäßig geprüft werden, ob eine Norm noch dem SdT entspricht und diesen korrekt wiedergibt (siehe auch VSS).

Der **Stand von Wissenschaft und Technik (SWT)** ist gemäß dem sogenannten Kalkar-Beschluss des BVerwG (siehe Lit.) unter der Berücksichtigung von neuesten wissenschaftlichen Erkenntnissen das „*technisch gegenwärtig Machbare*“. Hier stellt sich das Verhältnis zur Normung als gänzlich anders dar. Nach KAN spiegeln nicht Normen den SWT wider, sondern der Stand von Wissenschaft und Technik ist bei der Erstellung von Normen zu berücksichtigen. Dies ist auch in DIN 820-1 manifestiert, in der festgehalten ist, dass Normen den „*jeweiligen Stand der Wissenschaft und Technik sowie die wirtschaftlichen Gegebenheiten*“ zu berücksichtigen haben (siehe Lit.). Die beiden Begriffe „Stand der Wissenschaft und Technik“ sowie „Stand der Wissenschaft“ werden im Übrigen synonym zu „Stand von Wissenschaft und Technik“ verwendet. Nach der Handwerkskammer für München und Oberbayern (siehe Lit.) sollen die allgemein anerkannten Regeln der Technik eine Mehrheitsmeinung der Praxis darstellen, während der Stand der Technik nicht von der herrschenden Auffassung unter Praktikern abhängig ist. Er wird vielmehr danach bestimmt, was an der Front des technischen Fortschritts für geeignet, notwendig oder angemessen gehalten wird. Somit bedeutet Stand der Technik den technisch und wirtschaftlich realisierbaren Fortschritt.

Die unbestimmten Rechtsbegriffe technischer Sachverhalte sind in Bild 3.1 zusammenfassend dargestellt.

Begriff	Anerkannte Regeln der Technik	Stand der Technik	Stand von Wissenschaft und Technik
Definition	Die von der Mehrheit der Fachleute anerkannten - wissenschaftlich begründeten - ausreichend erprobten Regeln zum Lösen technischer Aufgaben	Das Fachleuten verfügbare Fachwissen, das - wissenschaftlich begründet, - praktisch erprobt und ausreichend bewährt ist	Der neueste Stand wissenschaftlicher und technischer Erkenntnisse, der - nachprüfbar begründet, - technisch durchführbar, - allgemein zugänglich, aber noch ohne praktische Bewährung ist
Praxisbeispiele	DIN-, DIN EN-, DIN EN ISO-, DIN IEC-Normen, VDI-Richtlinien, VDE-Vorschriften, UVV, Regeln technisch-wissenschaftlicher Vereine	Einzelnachweise übereinstimmender Bewertung in - Zeitschriftenbeiträgen, - Fachliteratur, - Sachverständigen-gutachten	Zeitpunktbezogene Darstellungen in - wissenschaftlichen Veröffentlichungen, - Kongressberichten, - Schutzrechtsschriften

Bild 3.1 Unbestimmte Rechtsbegriffe technischer Sachverhalte (angelehnt an Reuschlaw Legal Consultants, Berlin)

Normen gelten nach Helmig (siehe Lit.) als komplementärer Auslegungsmaßstab für die vertrags- und haftungsrechtliche Beurteilung technischer Produkte, da der Gesetzgeber zwar die Forderung nach sicheren Produkten stellt, diese aber nicht im Detail vorgibt und auch nicht vorgeben kann, wie die Sicherheit technisch zu gewährleisten ist. Normen geben die herrschende Auffassung der technischen Praktiken wieder. Aufgrund der praktischen Bedeutung und der Verbreitung von Normen herrscht gemäß Ensthaler, Müller und Symantschke allerdings „*faktisch ein Befolgungszwang*“ (siehe Lit.).

Die beiden Begriffe „Norm“ und „Standard“ werden im allgemeinen Sprachgebrauch sehr häufig in einem Atemzug genannt, was aber bei genauerer Betrachtung nicht korrekt ist. Im Gegensatz zu Normen umfassen Standards zwar ebenfalls Festlegungen von Regeln, Leitlinien und Merkmalen für Tätigkeiten und deren Ergebnisse bei allgemeinen und wiederkehrenden Anwendungen, allerdings sind an der Schaffung eines Standards nicht alle interessierten Kreise involviert (wie bei einer Norm), sondern ein geschlossener Kreis unter Ausschluss der Öffentlichkeit. Die Verabschiedung eines Standards erfolgt ebenfalls nicht in einem öffentlichen Verfahren, sondern in diesem kleinen Kreis. Standards werden z. B. von Verbänden, Organisationen oder auch einzelnen Unternehmen publiziert – oft wird hier von Industrie-Standards gesprochen. Alle Normen sind zwar Standards, aber nicht alle Standards sind auch gleichzeitig Normen.

Der Nexus von Standards lässt sich anhand von Bild 3.2 für den Bereich Funktionale Sicherheit gut veranschaulichen.

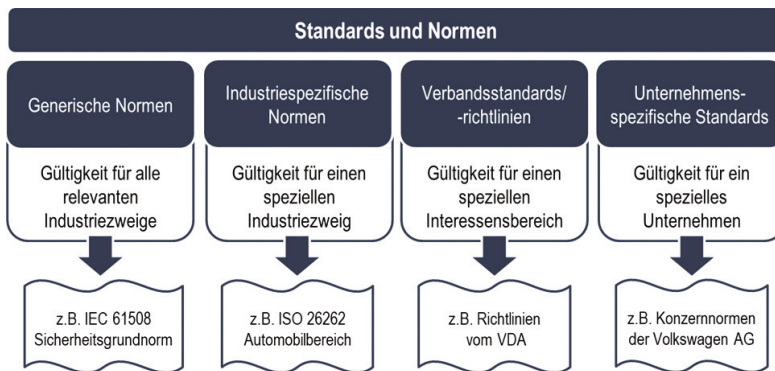


Bild 3.2 Nexus von Standards und Normen am Beispiel der Funktionalen Sicherheit

Verschiedene Unternehmen haben hauseigene *unternehmensspezifische Standards* zu speziellen Themengebieten erarbeitet. So gibt es z. B. von der Volkswagen AG die „Entwicklungsrichtlinie Funktionale Sicherheit“, in welcher die Erwartungen des Volkswagen Konzerns an den Umgang, die Auslegung und das Verständnis der ISO 26262 formuliert werden. Hierdurch soll im Rahmen einer verteilten oder unternehmensübergreifenden Entwicklung die Zusammenarbeit zwischen dem OEM und seinen Zulieferern verbessert werden, da u. a. Kommunikationsschwierigkeiten und Missverständnisse bei einer solchen Entwicklung reduziert werden. Arbeitet ein Unternehmen als Auftragnehmer mit der Volkswagen AG im Rahmen einer sicherheitsrelevanten Produktentwicklung zusammen, so

stellt diese Entwicklungsrichtlinie eine verpflichtende Anforderung dar. Dies wird in den entsprechenden Bauteillastenheften verankert.

Verschiedene Interessensverbände veröffentlichen *verbandsinterne Standards, Richtlinien und Empfehlungen*, welche die Interessen des entsprechenden Verbands darlegen und öffentlich verfügbar machen. Die darin enthaltenen Anwendungsregeln und Handlungsempfehlungen stehen oftmals deutlich schneller zur Verfügung, als dies im Rahmen einer herkömmlichen Norm möglich ist. So hat der Verband der Automobilindustrie (VDA) z. B. in einer Empfehlung einen Situationskatalog veröffentlicht, in dem Basissituationen und deren E-Parameter enthalten sind, welche für die Anwendung in einer Gefährdungsanalyse und Risikobewertung gemäß ISO 26262 bestimmt sind (siehe VDA-Empfehlung 702, siehe Lit.). Auch der Verband Deutscher Maschinen- und Anlagenbau (VDMA) greift in der Veröffentlichung seiner Einheitsblätter und Leitfäden immer wieder Themen der Funktionalen Sicherheit auf, in denen eine standardisierte Sichtweise dargelegt werden soll (siehe z. B. VDMA aus den Jahren 2009, 2012, 2013, siehe Lit.).

Industriespezifische Normen, wie z. B. im Kontext der Funktionalen Sicherheit die ISO 26262 für den Bereich der Straßenfahrzeuge oder die DIN EN ISO 13849 für den Bereich Maschinen, besitzen Gültigkeit für einen speziellen relevanten Industriezweig oder -sektor. Aus diesem Grund wird hierbei auch von sektorspezifischen Normen gesprochen. Solche industriespezifischen Normen konkretisieren allgemeinere Vorgaben auf den konkreten Bereich und berücksichtigen die dort herrschenden Gegebenheiten.

Eine Ebene höher existieren *generische Normen*, wie z. B. die DIN EN 61508, die als sogenannte Sicherheitsgrundnorm über einer Vielzahl anderer anwendungsspezifischer Normen steht und Gültigkeit für alle relevanten Industriezweige hat.

■ 3.2 Überblick über die Begrifflichkeiten: Safety, Security und Reliability

Bezüglich der Meta-Begrifflichkeiten „Safety“, „Security“ und „Reliability“ lohnt eine genauere Betrachtung und insbesondere eine Abgrenzung, da die Begriffe durchaus manchmal gleichgesetzt oder synonym verwendet werden.

Die Frage „Was ist **Sicherheit**?“ lässt sich so ohne Weiteres gar nicht so einfach beantworten. Zwar ist es unstrittig, dass Sicherheit ein zentraler Wertebegriff ist und dass sie eine der elementaren Voraussetzungen für alle Bereiche im öffentlichen Leben darstellt, aber der Begriff ist nicht genau definiert (vgl. Endreß/Petersen, siehe Lit.). Im allgemeinen Sprachgebrauch wird Sicherheit am ehesten mit der „*Abwesenheit von unvermeidbaren Risiken und drohenden Gefahren*“ erläutert. Der Sicherheitsbegriff an sich birgt jedoch eine ganze Reihe an verschiedenen und mehrschichtigen Facetten und Dimensionen (Bild 3.3), die und deren Grenzen es zu klären gilt, wenn sich mit dem Begriff Sicherheit auseinandergesetzt wird.

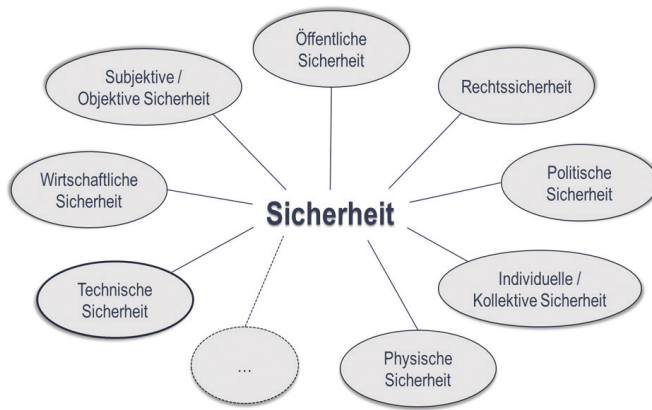


Bild 3.3 Aspekte und Facetten des Sicherheitsbegriffs

Wie in Bild 3.3 zu erkennen, kann und muss Sicherheit in verschiedenen Kontexten betrachtet werden. Die Abbildung soll dies lediglich verdeutlichen und sie ist auch in keinem Fall als vollständig anzusehen. Das vorliegende Werk hat nicht den Anspruch, eine politik- oder sozialwissenschaftliche Auseinandersetzung mit dem Begriff in der erforderlichen Detailtiefe durchzuführen. Vielmehr soll der Aspekt der „technischen Sicherheit“ näher betrachtet werden, da er für den Komplex der Zuverlässigkeits- und Sicherheitstechnik eine prägende Rolle spielt, insbesondere wenn die Begriffe „Safety“ und „Security“ ins richtige Licht gerückt werden sollen.

Laut dem Verein Deutscher Ingenieure (VDI) wird unter **Technischer Sicherheit** (siehe Lit.) begrifflich verstanden:

„dass ein technisches System, eine Anlage, ein Produkt über einen geplanten Zeitraum (gegebenfalls die geplante Lebensdauer) hinweg die vorgesehenen Funktionen erfüllt und bei bestimmungsgemäßer Nutzung keine geschützten Rechtsgüter verletzt, d. h. weder Personen noch Sachen geschädigt werden, soweit dafür das System, die Anlage oder das Produkt ursächlich sein können“. Weiter heißt es in VDI 2016, dass die Zuverlässigkeit der Funktion über die vorgesehene Lebensdauer kein Bestandteil der Sicherheit ist, sofern der Verlust der Funktion zu keinem unsicheren Zustand führt. Daran wird bereits deutlich, dass es einen Unterschied zwischen Sicherheit und Zuverlässigkeit geben muss, worauf aber an späterer Stelle noch genauer eingegangen wird.

Bei einer Betrachtung des angloamerikanischen Sprachraums fällt auf, dass es mit „Safety“ und „Security“ zwei Begriffe zu geben scheint, die heutzutage auch im deutschen Sprachraum sehr häufig im Kontext „Sicherheit“ verwendet werden. Allerdings ist die Bedeutung der beiden Begriffe durchaus unterschiedlich, was durch Bild 3.4 verdeutlicht wird.

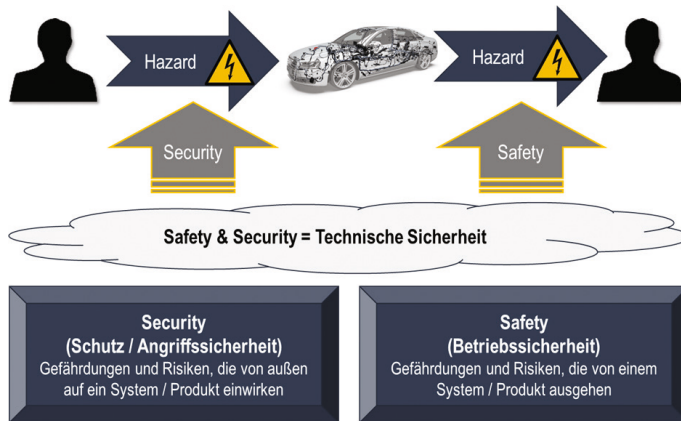


Bild 3.4 Bedeutung von Safety und Security im Kontext der Technischen Sicherheit

Von **Safety** wird gesprochen, wenn es um die Gefährdungen und Risiken geht, die von einem Produkt (in Bild 3.4 von einem Automobil) ausgehen. Dementsprechend geht es dabei um die Betriebssicherheit des Produkts und den Schutz von Personen, Sachen und der Umwelt. Die Sicht der Safety ist folglich „von innen nach außen“ gerichtet (Klauda, siehe Lit.).

Dahingegen betrachtet **Security** den umgekehrten Fall, nämlich die Sicht auf das Produkt von außen nach innen. Sie behandelt Gefährdungen und damit zusammenhängende Risiken, die von außen auf ein Produkt einwirken, weswegen hier auch von Angriffssicherheit oder vom Schutz des Produkts gesprochen wird. Manchmal wird unter Security lediglich der Schutz von Daten oder die Informationssicherheit verstanden, was allerdings zu kurz gegriffen scheint. In diesem Kontext, der ohne Zweifel eine hohe Wichtigkeit innehat, beschreibt die Informationssicherheit die Eigenschaften von informationsverarbeitenden Systemen, die für digitale Daten die Schutzziele Vertraulichkeit, Verfügbarkeit und Integrität sicherstellen (vgl. Pauli 2017, siehe Lit.). Zusammengefasst kann festgehalten werden, dass sich der zuvor bereits angesprochene Begriff der Technischen Sicherheit aus den beiden Aspekten Safety und Security zusammensetzt.

Bis vor einigen Jahren war die Grenze zwischen den beiden Domänen Safety und Security noch relativ gut zu trennen, da es wenige Überlappungsbereiche gab. Heutzutage, vor allem durch die immer stärker fortschreitende Vernetzung von Produkten im Zuge der Digitalisierung, verschwimmen die Grenzen immer stärker, sodass eine kombinierte Betrachtung dieser beiden Sicherheitsaspekte sinnvoll erscheint. Insbesondere der Bereich Automotive wird durch die steigende drahtlose Vernetzung der Fahrzeuge sowohl fahrzeugintern als auch mit der Umgebung (V2X/C2X – Vehicle/Car-to-X-Kommunikation) mit immer mehr potenziellen Angriffspunkten für externe Attacken konfrontiert, sodass das Thema Automotive Security wachsende Bedeutung genießt (Klauda, siehe Lit.). Der Bundesverband Digitale Wirtschaft (BVDW, siehe Lit.) veranschlagte den Gesamtmarkt „Connected Cars“ im Jahr 2015 auf ca. 32 Mrd. Euro. Laut BVDW sollte das Wertschöpfungspotenzial bis ins Jahr 2020 auf ca. 115 Mrd. Euro anwachsen. Andere Studien gingen für den gleichen Zielhorizont sogar von über 130 Mrd. Euro aus (siehe magility 2018). Der Informationssicherheit

wird folglich eine immer stärkere Bedeutung zukommen, welche es neben dem „bewährten Feld“ der Safety zu berücksichtigen gilt.

Ein einfaches Beispiel für den teils konträren Fokus von Safety und Security im Automobilbereich kann durch ein Türschließsystem verdeutlicht werden. Im normalen Alltagsgebrauch soll ein Türschloss das Fahrzeug durch ungewollte und unbefugte Zugriffe von außen schützen, sodass niemand das Fahrzeug oder darin befindliche Gegenstände stehlen, seine Insassen verletzen oder entführen kann (Security). Im hoffentlich sehr unwahrscheinlichen Fall eines Unfalls soll das Türschloss jedoch eine leichte Öffnung der Tür ermöglichen, um entsprechende Evakuierungs- und Rettungsmaßnahmen zu ermöglichen (Safety).

Noch gut in Erinnerung dürfte der Vorfall aus dem Sommer 2015 sein, bei dem die beiden Security-Experten Charlie Miller und Chris Valasek live demonstrierten, wie sie per Remote-Hack über das Internet die Kontrolle über ein modernes SUV-Fahrzeug (Jeep Cherokee von 2014) übernahmen – und dies während der vollen Fahrt auf einem amerikanischen Highway in St. Louis (Missouri, USA). Dabei verschafften die beiden sich nicht nur Zugriff auf die Klimaanlage, das Radio oder die Scheibenwischer, sondern konnten auch gezielt in die Fahrzeugsteuerung eingreifen und so beispielsweise die Motorkontrolle deaktivieren. Der Fahrer des Jeeps, ein US-Journalist der US-Internetseite „Wired“, der über den Hack im Vorfeld Bescheid wusste und entsprechend vorbereitet war, wurde bei seinen Reaktionen gefilmt (siehe Greenberg 2015). Er konnte gegen keinen der Angriffe etwas unternehmen; so war es ihm z. B. nicht möglich, die plötzlich geänderte Lautstärke des Radios wieder zurückzudrehen oder das Fahrzeug manuell zu beschleunigen. Eindrucksvoll und medienwirksam wurde damals gezeigt, wie ein eigentlich technologisch fortschrittliches System wie ein modernes High-End-Fahrzeug Ziel einer solchen Attacke werden kann und dabei unmittelbar Safety-Aspekte betroffen sein können. Der Angriff erfolgte damals über eine Mobilfunkverbindung auf das Fahrzeug-Entertainment-System. Der Jeep-Hersteller Fiat Chrysler musste daraufhin übrigens eine großangelegte Rückrufaktion durchführen (siehe Greenberg 2015), in der bei rund 1,4 Millionen Fahrzeugen die Software erneuert wurde.

Da beide Domänen zwar durchaus ähnliche Entwicklungsprozesse verwenden und es Synergien bei der technischen Implementierung gibt, erscheint eine gemeinschaftliche Betrachtung sinnvoll. Allerdings dürfen die ebenfalls vorhandenen konkurrierenden Aspekte dabei nicht vernachlässigt werden. Eine Fehlfunktion in einer Security-Funktion darf beispielsweise nicht zu einer Deaktivierung einer Safety-Komponente führen. Nach Klauda (siehe Lit.) ist darauf zu achten, dass als Handlungsleitlinie die Security-Maßnahmen wichtige Safety-Funktionen nicht beeinträchtigen dürfen und die Verfügbarkeit der Safety-Funktion gegenüber der Security-Funktion Vorrang haben muss.

In Bild 3.5 ist ein Vorschlag für die gemeinschaftliche Adressierung von Safety- und Security-Aspekten in einem Entwicklungsprozess zu finden. Auch Bild 3.6 zeigt die Analogien der Safety-Vorgehensweise gemäß ISO 26262 und der Security-Prozessschritte anhand eines vereinfachten V-Modells.

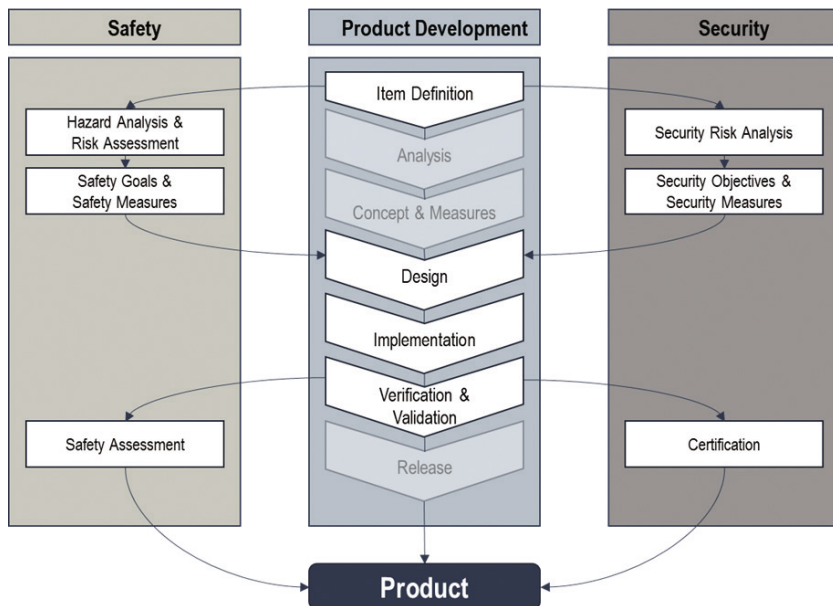


Bild 3.5 Vorschlag zur Berücksichtigung von Safety und Security in einem gemeinsamen Entwicklungsprozess nach Strasser et al. (siehe Lit.)

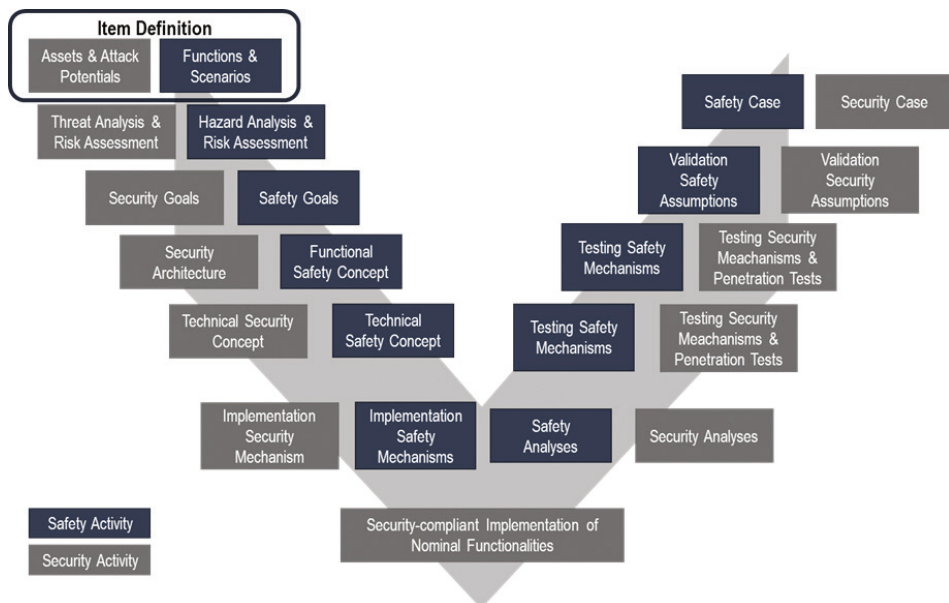


Bild 3.6 Gemeinsame Betrachtung von Safety- und Security-Anforderungen nach Ebert und Metzker (siehe Lit.)

Es wird deutlich, dass beide Themen nicht unabhängig voneinander betrachtet werden sollten. Vielmehr sollten Vorgehensweisen für eine gemeinsame Bearbeitung definiert werden, um vorhandene Synergien bestmöglich nutzen zu können.

Neben Safety und Security spielt ein weiterer Aspekt eine sehr wichtige Rolle bei der Entwicklung von Produkten: die Zuverlässigkeit. Unter dem Begriff der Zuverlässigkeit wird allgemein die Fähigkeit einer Betrachtungseinheit verstanden, eine geforderte Funktion unter den festgelegten Anwendungsbedingungen während der definierten Missionsdauer korrekt auszuführen. DIN EN 60300 definiert die Zuverlässigkeit einer Einheit kurz und knapp als „*Fähigkeit zu funktionieren wie und wann gefordert*“ (siehe Lit.). Zuverlässigkeit kann somit auch als „Qualität auf Zeit“ verstanden werden, wobei anstelle einer definierten Zeitspanne auch eine festgelegte Anzahl an Betriebszyklen oder Schaltspielen benutzt werden kann. **Reliability** wird oftmals mit dem Begriff Zuverlässigkeit gleichgesetzt (siehe z.B. VDA 2016), was allerdings ein wenig missverständlich ist. Er sollte vielmehr teils in der Bedeutung „Funktionsfähigkeit“ und teils in der Bedeutung „Überlebenswahrscheinlichkeit“ definiert werden (vgl. auch DIN 1190).

Alle drei Aspekte, also Safety, Security und Reliability stehen in Wechselwirkungen zueinander, wie Bild 3.7 visualisiert.

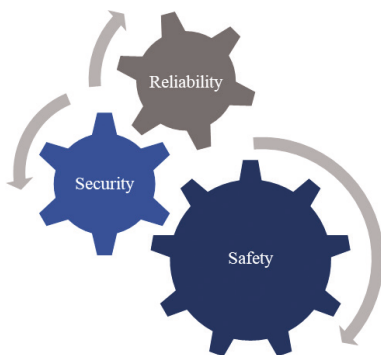


Bild 3.7

Wechselwirkung zwischen Safety, Security und Reliability nach Pauli 2017 (siehe Lit.)

Werden Security-Maßnahmen in ein technisches Produkt eingeführt, darf es in keinem Falle dazu kommen, dass die Betriebssicherheit darunter leidet und es zu einem erhöhten Risiko für die Anwender/Verbraucher kommt. Auf der anderen Seite muss bei neuen Safety-Features geschaut werden, ob sich diese in Einklang mit den bereits vorhandenen Security-Maßnahmen befinden oder diese außer Kraft setzen. Die Einführung neuer technischer Sicherheitsmaßnahmen führt unweigerlich zu einer Zunahme der Komplexität des Produkts. Dies kann grundsätzlich die Wahrscheinlichkeit für einen Ausfall erhöhen, was wiederum zu einer negativen Beeinflussung der Zuverlässigkeit führen wird. In Anhang A sind einige der wichtigsten und relevantesten Standards für die Bereiche Safety, Security und Reliability zu finden, welche im Rahmen der Entwicklung technischer Produkte zur Anwendung kommen. Die Standards werden jeweils kurz vorgestellt, wobei sich auf die für den deutschsprachigen Raum gültige Version bezogen wird. Die enthaltene Auflistung erhebt dabei keinerlei Anspruch auf Vollständigkeit; sie soll vielmehr als Anhaltspunkt für tiefergehende Recherchen dienen. Zu der Sicherheitsgrundnorm IEC 61508 sind in Abschnitt 5.1.2 vertiefende Angaben zu finden. Das entsprechende Derivat ISO 26262 für den Straßenverkehr wird in Abschnitt 5.1.4 ausführlich behandelt.

Teil II:

Zuverlässigkeit im Produktentstehungsprozess

4

Zuverlässigkeit im Produktentwicklungsprozess

■ 4.1 Zuverlässigkeitsprozess

Der Zuverlässigkeitsprozess ist ein stetiger Begleitprozess für den üblichen Produktlebenszyklus und geht zeitlich gesehen weit über die eigentliche Entwicklung eines Produkts hinaus. Es ist allerdings sehr wichtig, alle Zuverlässigkeitsthemen schon während der Entwicklung zu betrachten. Die systematische Rückführung von Feldinformationen in den Produktentstehungsprozess (PEP) steht aber ganz klar im Fokus. Der Prozess nutzt dabei insbesondere Informationen von Vorgängerbaureihen oder -produkten beziehungsweise ähnlichen Produkten, um Kenntnisse zum Nutzungsverhalten wie zum Beispiel Nutzungsintensitäten durch den Kunden in Abhängigkeit meist sehr unterschiedlicher Märkte, zu kritischen Schädigungsparametern und zum Ausfallverhalten selbst zu gewinnen. Der Zuverlässigkeitsprozess kann daher als ein sich ständig wiederholender in sich geschlossener Kreisprozess angesehen werden. In Bild 4.1 ist ein möglicher ganzheitlicher Zuverlässigkeitsprozess abgebildet.

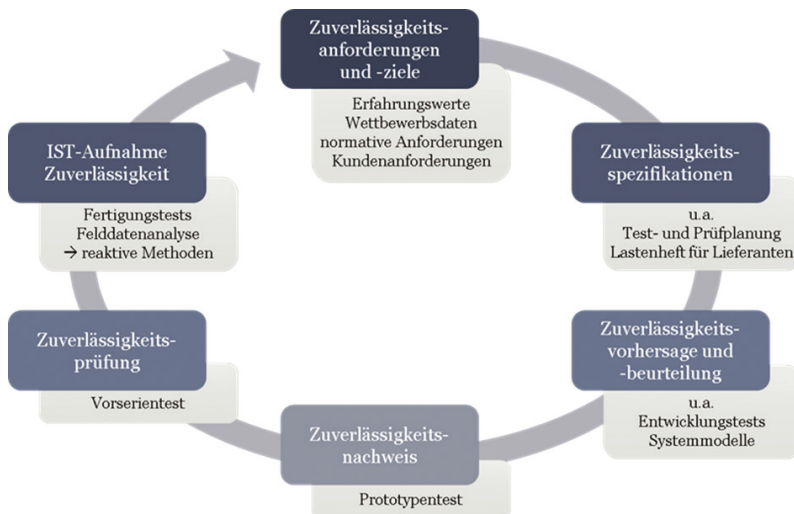


Bild 4.1 Zuverlässigkeitskreisprozess

Häufig werden allerdings nur Teilaspekte der Zuverlässigkeit im Entwicklungsprozess gesehen, sodass die Produktzuverlässigkeit im PEP nur unzureichend betrachtet wird. Etliche Themen werden dann in der Produktion und weiteren Nachfolgeprozessen (Service-/Reklamationsmanagement, Schadteilanalyseprozess, Lessons Learned) verortet und werden damit sehr spät betrachtet. Dies hat zur Folge, dass einige unternehmenskritische Themen zu spät fokussiert werden.

Ein sehr gutes Beispiel für solch ein Thema sind Gewährleistungs- und Garantiekosten in Bezug zu Entwicklungs- und Produktkosten. Jedes System hat bekanntlich eine inhärente Ausfallwahrscheinlichkeit, die unter anderem abhängig von der Qualität der Produkte, aber auch von der Auslegung der Systeme ist. Es ist leicht ersichtlich, dass kostengünstige, aber qualitativ unzureichende Produkte zu mehr Reklamationen während der Gewährleistungs-/Garantiezeit führen können. Diese Reklamationen führen unweigerlich zu Kosten, die als Abschätzung mit in die Kostenrechnung zum Entwicklungsprojekt einfließen sollten. Sind die Kosten zu hoch, kann ein ganzes Projekt unwirtschaftlich oder gar existenzbedrohend sein. Werden die Gewährleistungs-/Garantiekosten nicht dem Projekt zugerechnet, sondern auf eine andere Kostenstelle (z.B. auf ein Produktionswerk) gebucht, so ist unter Umständen nicht ersichtlich, dass solch ein Projekt so nicht hätte durchgeführt werden dürfen. Da die Entwicklung in nahezu allen Industriebereichen sehr kostengetrieben ist, gilt es hier, einen guten Kompromiss zwischen Entwicklungs-/Produktkosten und Gewährleistungs-/Garantiekosten zu finden.

Ein Zuverlässigkeitsprozess hilft dabei, derartige „Stolpersteine“ im Produktleben zu identifizieren und zu bewerten. Er sollte daher durchgängig in eine Organisation implementiert sein und unter anderem dazu genutzt werden, die in Bild 4.2 dargestellten Themen während eines Produktlebenszyklus im Blick zu behalten.



Bild 4.2 Kleeblatt des Produktionszyklus

In diesem Kontext ist es wichtig zu verstehen, dass die Themen sehr komplex miteinander verknüpft sind. Aufgrund der häufig verteilten Entwicklung und der Aufbau- und Ablauforganisation einer Firma können diese Themen von den Entwicklern und vom Projektmanager aber nicht mehr überblickt werden. So ist z.B. vielen Entwicklern in der Automobilindustrie nicht bewusst, dass aufgrund der Abrechnungsmodelle für einen einzelnen Ausfall ein Vielfaches der eigentlichen Kosten anfallen kann. Üblich sind Faktoren bis hin zu einer Größenordnung von ca. 150.

Zuverlässigkeitsmanagement

Zuverlässigkeitsmanagement (Reliability Management) bedeutet, zu einem frühen Zeitpunkt im Produktentstehungsprozess Zuverlässigkeitsvorgaben („Zuverlässigkeitsziele“) zu definieren und deren Entwicklung im Sinne eines Zuverlässigkeitswachstums (Reliability Growth) über den gesamten Produktentstehungsprozess zu überwachen. Das Zuverlässigkeitsmanagement hört allerdings nicht mit dem Projektende und dem Serienproduktionsstart (Start of Production, SOP) auf. Auch während des Feldeinsatzes werden moderne Methoden der Felddatenauswertung und Mustererkennungsverfahren genutzt, um zum einen eine passive Marktbeobachtung durchzuführen und zum anderen aber auch wichtigen Input für neue Entwicklungsprojekte zu generieren.

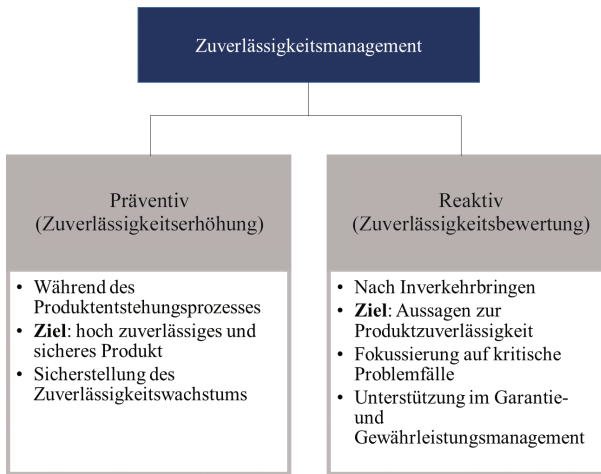
Wichtig dabei ist, dass Zuverlässigkeit in die Unternehmenskultur einfließt und kein einmaliges Ereignis ist. Eine organisatorische und prozessuale Verankerung im Unternehmen ist folglich unverzichtbar. Eine große Herausforderung bei der Verankerung spielt das Verständnis von Zuverlässigkeit.

Häufig interpretieren die am Zuverlässigkeitsprozess beteiligten Personen die Zuverlässigkeit anders. Es gilt dementsprechend diverse Fragestellungen vorab zu klären. Hierzu gehören unter anderem folgende:

- Wie zuverlässig muss mein Produkt sein?
- Wie wird das Produkt überhaupt genutzt und was wird von ihm erwartet?
- Ist ein Produkt zuverlässig, wenn es alle definierten Tests bestanden hat?
- Was sagen bestandene Tests über die Zuverlässigkeit des Produktes aus?
- Wie viel muss überhaupt getestet werden oder können Tests in Gänze entfallen?
- Ist ein Produkt automatisch unzuverlässig, wenn es Reklamationen gibt?
- etc.

Zahlreiche Standards und Richtlinien, wie z.B. VDI 4004, DIN 60300 oder VDA 3.1, 3.2 und 3.3, beschäftigen sich mit dem Themenkomplex Zuverlässigkeit und sollten im Unternehmen verstanden und beherrscht werden.

Wie bereits erwähnt, kann Zuverlässigkeitsmanagement in einen präventiven sowie in einen reaktiven Bereich unterteilt werden. Beide Bereiche sind eng miteinander verzahnt, tauschen im Optimalfall Informationen aus und stellen eine ganzheitliche Betrachtungsweise des Zuverlässigkeitsmanagements sicher (Bild 4.3). Präventives und reaktives Zuverlässigkeitsmanagement werden im Folgenden näher beschrieben.

**Bild 4.3**

Präventives und reaktives
Zuverlässigkeitsmanagement

Präventives Zuverlässigkeitsmanagement

Präventives Zuverlässigkeitsmanagement beschäftigt sich mit den Fragen der Zuverlässigkeit während des Produktentstehungsprozesses (PEP), um schon zu einem frühen Zeitpunkt potenzielle Fehler zu entdecken und kostengünstig abzustellen. Hierbei gilt es, die benötigte Zuverlässigkeit, also das Zuverlässigkeitsziel, zu bestimmen und ein Produkt so zu entwerfen, dass nicht mehr Feldausfälle auftreten als festgelegt und damit einhergehend die Garantie- und Gewährleistungskosten im kalkulierten Rahmen bleiben. Dies ist besonders im Hinblick auf steigende Gewährleistungszeiten extrem wichtig.

Folgende Themen sollten in einem präventiven Zuverlässigkeitsmanagement behandelt werden:

- Zuverlässigkeitsanforderungen und -ziele (Erfahrungswerte, Wettbewerbsdaten, normative Anforderungen, Kundenanforderungen etc.)
- Zuverlässigkeitsspezifikationen (unter anderem Test- und Prüfplanung, Lastenheft für Lieferanten, Bewertung Systemaufbau)
- Zuverlässigkeitsvorhersage und Beurteilung (unter anderem Entwicklungstests, Systemmodelle, Simulationen)
- Zuverlässigkeitsnachweis (Prototypentest)
- Zuverlässigkeitsprüfung (Vorserientest)
- IST-Aufnahme Zuverlässigkeit (Fertigungstests, Felddatenanalyse siehe reaktive Methoden)

Die Zuverlässigkeitsvorgaben werden auf der einen Seite für die Auslegung der Produkte herangezogen und auf der anderen Seite dazu genutzt, um mit einer statistisch fundierten Test- und Prüfplanung einen im Sinne des Herstellers Kosten-Nutzen-optimierten Zuverlässigkeitsnachweis zu erbringen.

Die Nutzung der vorangehend genannten Feldinformationen stellt sicher, dass relevante Testszenarien richtig erkannt werden und mit den zu erwartenden Belastungen im Einsatz auch tatsächlich korrelieren.

Sinnvoll ist es, die Zuverlässigkeitsthemen direkt in den PEP an den relevanten Meilensteinen zu integrieren (Bild 4.4).

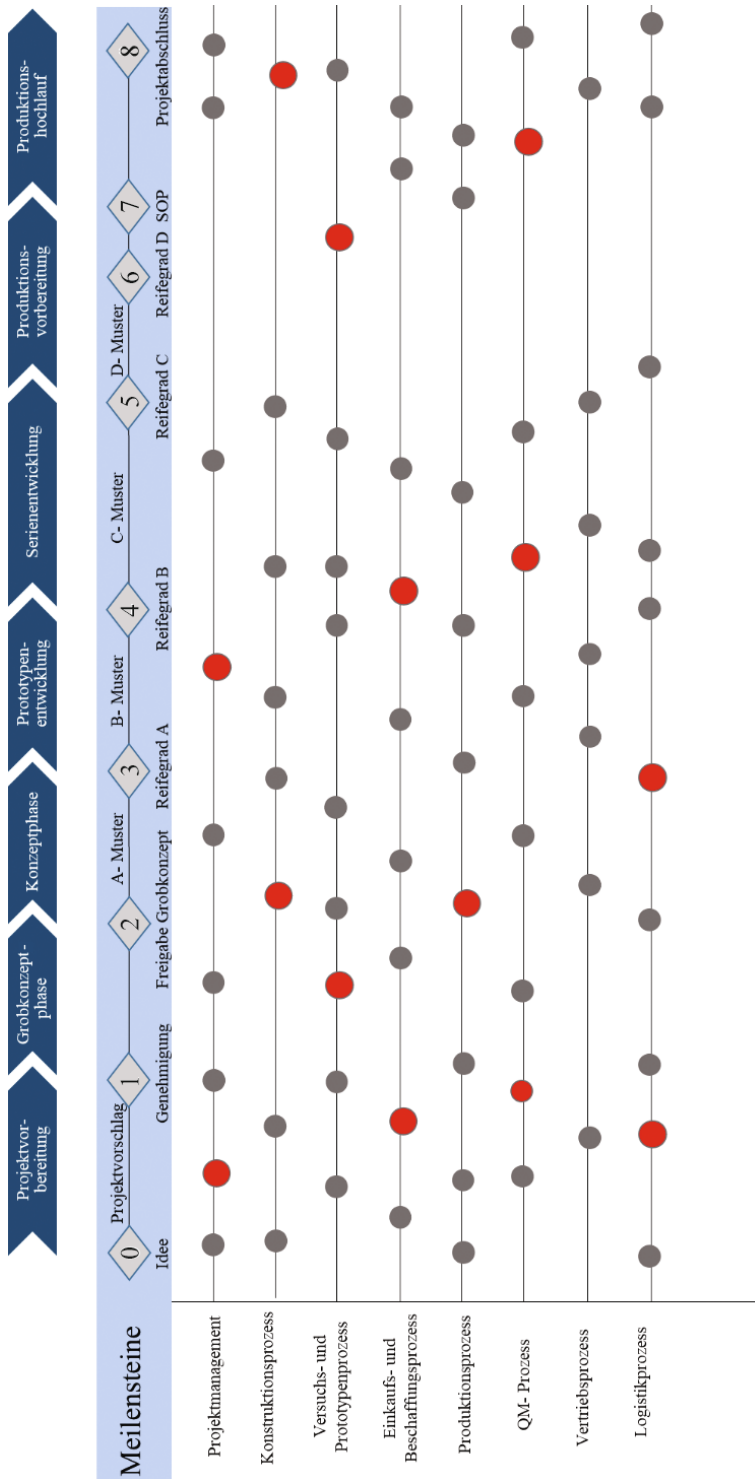


Bild 4.4 Meilensteine des Zuverlässigkeitswachstums

Zudem soll das Zuverlässigkeitswachstum (Reliability Growth Management, RGM) sichergestellt werden. Der Begriff ist etwas irreführend, da durch die Tätigkeiten im RGM die Zuverlässigkeit im eigentlichen Sinne nicht gesteigert wird, sondern nur die Nachweisführung und damit das Vertrauen in die Zuverlässigkeit. Das RGM zeigt somit, ob sich die Entwicklung auf dem richtigen Weg befindet und die vorgegebene Zuverlässigkeit erreicht wird. In Bild 4.5 ist eine mögliche Umsetzung des RGM dargestellt.

Typische Fragen, die sich bei der Entwicklung ergeben, lassen sich unter Heranziehung von Bild 4.5 beantworten. Hierzu zählen folgende:

- Wie viel Prozent Zuverlässigkeit haben wir zum aktuellen Zeitpunkt schon praktisch nachgewiesen?
- Durch welche Prüfungen wurde diese Zuverlässigkeit nachgewiesen?
- Können wir unser Zuverlässigkeitsziel unter Beibehaltung des geplanten Erprobungsumfangs noch erreichen?
- Wie viele Prüfungen müssen ggf. nachgeschoben werden, um das Ziel noch zu erreichen?

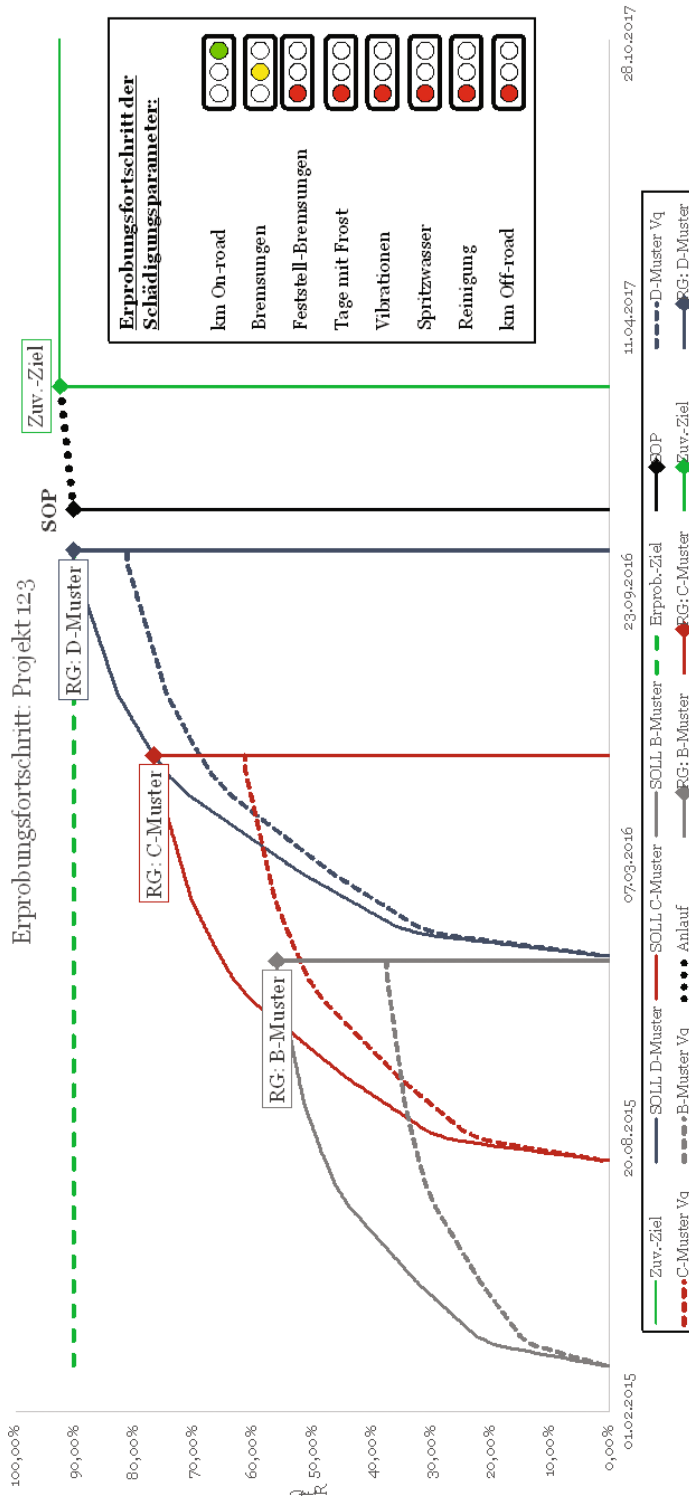


Bild 4.5 Umsetzung des Reliability Growth Management

Reaktives Zuverlässigkeitsmanagement

Das reaktive Zuverlässigkeitsmanagement betrifft sämtliche Lebensabschnitte des Produktlebenszyklus, welche **nach** Auslieferung des Produkts an den Kunden erfolgen. Dabei wird versucht, anhand realer Daten aus dem Feld eine Aussage zur Zuverlässigkeit des entsprechenden Produkts zu treffen. Zudem gibt es Methoden, mit deren Hilfe sich aus Daten Fehlerschwerpunkte ermitteln lassen (Mustererkennung, Kapitel 28). Diese Methoden sollten zu Beginn einer Zuverlässigkeitsanalyse von Felddaten immer vorgeschaltet werden, um den späteren Analyseaufwand auf kritische Problemfälle zu fokussieren. Somit werden finanzielle oder auch personelle Kapazitäten nicht nach dem Gießkannenprinzip auf alle Produkte gleichmäßig, sondern zielgerichtet verteilt.

Das reaktive Zuverlässigkeitsmanagement ist ein Unterstützungsprozess für das Risikomanagement. Mithilfe der Marktbeobachtung und der Felddatenerfassung, -auswertung und -bewertung kann sichergestellt werden, dass sich kritisch entwickelnde Feldprobleme sehr zeitnah erkannt und ihr Risiko richtig bewertet wird. Die Risikoeinstufungen helfen unter anderem dabei, Feldaktionen, wie z.B. Rückrufe, aktiv zu managen und entsprechend der Kritikalität zu steuern. Beispielhaft sei hier eine Thematik benannt, die zuerst bei Vielfahrern (größer 50 000 km pro Jahr) aufgetreten war. Hier ist schnell zu erkennen, dass eine Feldaktion mit den anderen Vielfahrern starten sollte. Im weiteren Verlauf werden dann erst die Normalfahrer in der Aktion berücksichtigt und zum Schluss wird die Aktion mit den Wenigfahrern beendet, um das Flottengesamtrisiko möglichst klein zu halten (siehe hierzu Abschnitt 28.3).

Die Ergebnisse aus Felddatenanalysen – z.B. aus einer Zuverlässigkeitsprognose – eignen sich hervorragend zur Unterstützung des präventiven Zuverlässigkeitsmanagements. Unter anderem lassen sich damit Zuverlässigkeitsanforderungen und folglich Ziele formulieren, die eine wichtige Grundlage im Entwicklungsprozess darstellen. Weiterhin sind Felddatenanalysen ein geeignetes Mittel zur Unterstützung des Garantie- und Gewährleistungsmanagements. Mit ihrer Hilfe lassen sich Garantiekosten prognostizieren, Risiken bei Garantiezeiterweiterungen abschätzen oder Rückrufaktionen verifizieren. Auch bei einer bevorstehenden Serienschadensverhandlung liefern Felddatenanalysen einen wichtigen Beitrag zur Verhandlungsvorbereitung. Zudem können sie als Entscheidungshilfen bei der Kalkulation von Serienersatzbedarf oder Endbevorratung herangezogen werden. Somit sind sie ein wichtiger Bestandteil des Obsoleszenz- und Ersatzteilmanagements (siehe Abschnitt 28.3).

■ 4.2 Bereiche, Rollen und Verantwortlichkeiten

Das Zuverlässigkeitsmanagement greift bekanntlich in nahezu alle Bereiche einer Firma ein. Viele Verantwortlichkeitsbereiche sind demnach im Zuverlässigkeitsprozess involviert. Die Verantwortlichkeiten in den einzelnen Bereichen müssen hier gut festgelegt sein, damit der Produktionsprozess geordnet ablaufen kann. Generell gilt: Zuverlässigkeit ist, wie auch die Qualität, eine Gemeinschaftsaufgabe.

Die Stakeholder im Unternehmen sind zahlreich und nicht immer identifizierbar. In Bild 4.6 sind die beteiligten Hauptbereiche mit einigen beispielhaften Verantwortlichkeiten dargestellt.

Qualitätsmanagement • Strategische Leitung • Entwicklung Methoden • Klammerfunktion	Kundendienst • Bereitstellung und Analyse Felddaten • Unterstützung Versuch/Erprobung • Voice of Customer	Versuch • Zielbestimmung • Test- und Prüfplanung • Zuverlässigkeitswachstum	Konstruktion • Zielbestimmung • Test- und Prüfplanung • Zuverlässigkeitsmonitoring
SCM • Lieferantenauswahl /-befähigung • Zielzahlen in ext. LH	Vertrieb • Analyse Märkte • Voice of Customer • Unterstützung Zielbestimmung	Produktmanagement • Analyse Märkte • Voice of Customer • Unterstützung Zielbestimmung • Benchmark	Geschäftsführung • Leitbild • Zielbestimmung • Freigaben • Lob und Tadel
Fertigung / Montage / Logistik • Unterstützung Ursachenanalyse Feldreklamationen • Fähige Prozesse	Unternehmenskommunikation • Strategie Außenkommunikation bei Rückrufen • Öffentliche Wahrnehmung	Risikomanagement / Versicherungen / Recht • Nutzung Zuverlässigkeitswerte für Risikoeinstufung • Rücklagenberechnung • Erprobungsklausel	... • ...

Bild 4.6 Zuverlässigkeitsmanagement und Verantwortlichkeitsbereiche

Die Verantwortlichkeitsbereiche können, je nach Firmenphilosophie, hierbei sehr variabel benannt und den unterschiedlichen Bereichen zugeordnet sein. Die Verknüpfung der Bereiche Rollen und Verantwortlichkeiten hängt auch stark von den gewählten Auslegungen in der Aufbau- und Ablauforganisation ab, die in jedem Fall gezielt aufeinander abgestimmt sein müssen (Bild 4.7). Bild 4.8 zeigt beispielhaft ein generelles Organisationsmodell für den Bereich Warranty.

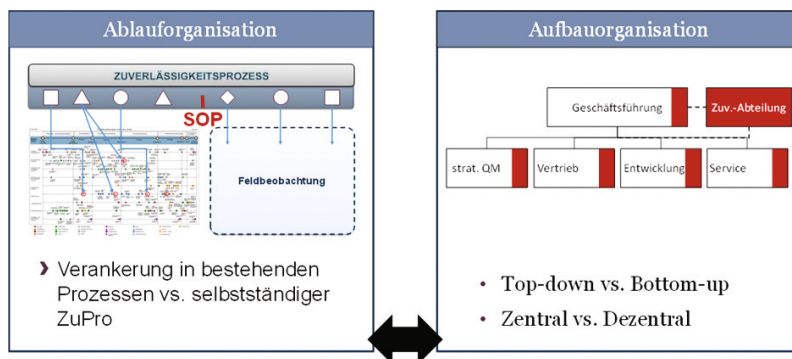


Bild 4.7 Verknüpfung der Bereiche Ablauf- und Aufbauorganisation

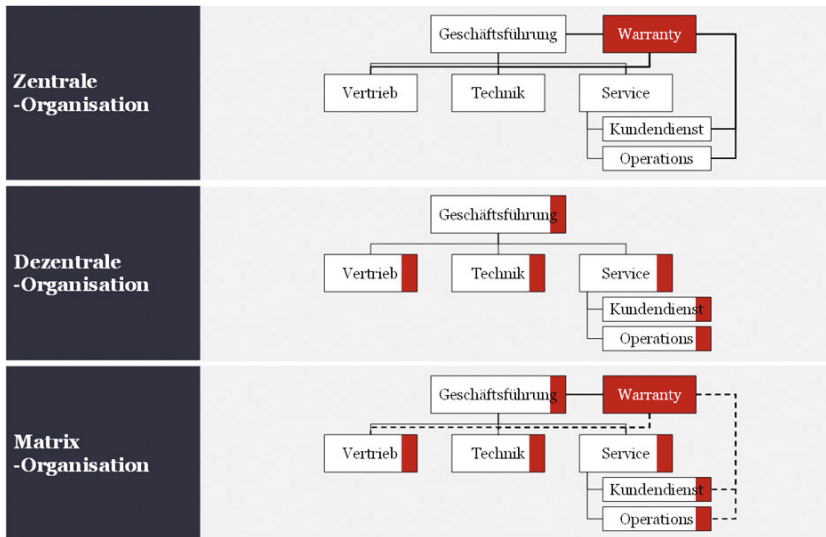


Bild 4.8 Beispiel Organisationsmodelle Warranty

Empfehlenswert kann auch eine Matrix-Organisation sein. Diese stellt sicher, dass alle Beteiligten das Thema Zuverlässigkeit in ihrem Bereich wahrnehmen, die Risikoidentifikation durch alle erfolgt und der Zentralabteilung, als unabhängige Bewertungsinstanz, bei der Risikobewertung zur Verfügung steht. Darüber hinaus bietet sie den Vorteil, dass die Methoden projektübergreifend standardisiert sind und es einen zentralen Ansprechpartner gibt, der gegebenenfalls unterstützend tätig werden kann.

Nachteilig sind der erhöhte personelle Ressourcenbedarf und der damit verbundene Schulungsaufwand.

■ 4.3 Wirtschaftlichkeitsaspekte und Zuverlässigkeitsziele

Die Wirtschaftlichkeit eines Produktes steht in den meisten Industriezweigen deutlich im Fokus. Speziell in der Automobilindustrie sind die Kosten für das Produkt eines der beherrschenden Themen. Diese stehen häufig im Widerspruch zum Verständnis der Produktentwickler zur Zuverlässigkeit, die es gilt zu optimieren.

Oft genug ist gerade dieser Zielkonflikt, Kosten auf der einen Seite, Zuverlässigkeit auf der anderen Seite, ein Grund für Feldaktionen: Die Kosten im Projekt müssen reduziert, d.h. alle Möglichkeiten zur Kosteneinsparung in Betracht gezogen werden. Ein verständliches Vorgehen ist dann oft die Beschaffung vermeintlich günstigerer Bauteile, Komponenten, Produkte, die aber unter Umständen qualitativ nicht so hochwertig sind.