

Technische Universität Berlin
Institut für Psychologie und Arbeitswissenschaft

Genehmigte Dissertation

Übersteigertes Vertrauen in Automation: Der Einfluss von Fehlererfahrungen auf Complacency und Automation Bias

zur Erlangung des akademischen Grades
Doktorin der Philosophie (Dr. phil.)
der Fakultät für Verkehrs- und Maschinensysteme (Fakultät V)
der Technischen Universität Berlin

vorgelegt von
Dipl.-Psych. Jennifer Elin Bahner

Promotionsausschuss:

Vorsitzender: Prof. Dr. Manfred Thüring
Berichter: Prof. Dr. Dietrich Manzey
Prof. Dr. Hartmut Wandke

Tag der wissenschaftlichen Aussprache: 11. September 2008

Berlin 2008

D 83

DANKSAGUNG

Die vorliegende Arbeit wäre nicht in dieser Form möglich gewesen ohne die Unterstützung durch mir wichtige Personen. Zu allererst möchte ich meinem Doktorvater Dietrich Manzey danken, der immer da war um mir wichtig erscheinende Themen und Fragen ausführlich und oft auch kontrovers zu diskutieren. Das theoretische Begriffsverständnis, dass ich heute zum Thema der „Automationspsychologie“ habe, aber auch ein Großteil der experimentellen Umsetzungs-ideen ist gerade durch diese Gespräche entstanden. Ohne diese fachliche Anbindung hätte dieser Arbeit sicherlich der „Nährboden“ gefehlt.

Besonderer Dank gilt Marcus Bleil, ohne dessen umfangreiche und dennoch geduldige Anpassungen der Experimentalumgebung diese Arbeit praktisch nicht realisierbar gewesen wäre, und Sabine Jatzev, die mir durch die Erstellung von Scripts für die Auswertung – teils in Wochenend- und Nachtarbeit – eine sehr große Hilfe war. Danken möchte ich darüber hinaus meinen Diplomandinnen Anke-Dorothea Hüper und Monika F. Elepfandt für anregende Diskussionen und die aufwändige Erhebung der Daten. Ein großes Dankeschön auch an Monika Gibler und Ulrike Wiedensohler, die meine Arbeit so gewissenhaft und trotzdem in nur kurzer zur Verfügung stehender Zeit redigiert haben.

Meinen Freunden Beatrice Viktoria Feuerberg, Dr. Sonja Geiger, Stefan Röttger, Friederike Schmidt und Doris Wiedemann danke ich sehr dafür, dass sie immer ein offenes Ohr für die im Rahmen der Arbeit entstandenen Fragen hatten und mir stets durch Anregungen und Ratschläge beiseite standen, und auch dann noch, wenn es ganz schnell gehen musste, Teile der Arbeit Korrektur lasen.

Ganz herzlich danke ich meiner Mutter Susanne Bahner, die mir half die Prioritäten stets richtig zu setzten und – „last but not least“ – meinem Mann Stefan Heyne, dem ich einen riesigen Motivationsschub in der Endphase meiner Doktorarbeit zu verdanken habe!

KURZZUSAMMENFASSUNG

Durch die zunehmende Zuverlässigkeit automatisierter Systeme konnte in den vergangenen Jahren eine Vielzahl potenzieller Fehlerquellen der Mensch-Maschine Interaktion reduziert werden. Zugleich sind dadurch jedoch auch neue Risiken entstanden: Gerade bei hoch reliablen automatisierten Systemen besteht die Gefahr eines übersteigerten Vertrauens in das System, wobei *complacency* und *automation bias* mögliche Folgen auf Verhaltensebene darstellen. *Complacency* stammt aus dem Kontext klassischer *monitoring*-Aufgaben. Zu verstehen ist darunter eine unzureichende Überwachung oder Überprüfung der Automation, die zu einem Übersehen kritischer Systemzustände führen kann. Demgegenüber stammt das Konzept *automation bias* aus dem Kontext der Nutzung von Entscheidungsassistenzsystemen. Gefasst werden darunter zwei verschiedene Fehlertypen: Während *commission* Fehler darin bestehen, dass ein Operateur einer fehlerhaften Empfehlung eines Assistenzsystems folgt, äußern sich *omission* Fehler darin, dass kritische Systemzustände übersehen werden, sofern diese vom Assistenzsystem nicht angezeigt werden. Trotz der engen konzeptuellen Nähe von *complacency* und *automation bias*, sind die Phänomene bislang nur getrennt voneinander empirisch untersucht worden. Zentrales Anliegen der vorliegenden Arbeit war es, methodische Schwächen bisheriger Studien zu überwinden und einen Beitrag zu einer empirisch fundierten Klärung der Konzepte *complacency* und *automation bias* sowie des Bezugs der Konzepte untereinander zu leisten. Darüber hinaus wurden mögliche Gegenmaßnahmen im Sinne von Automationsfehlern im Training untersucht. Zwei experimentelle Studien (jeweils $N = 24$) wurden durchgeführt, wobei eine Mikrowelt als Versuchsumgebung diente, in der die Probanden bei der Detektion, Diagnose und Behebung von Fehlfunktionen einer Prozesssteuerung durch ein Assistenzsystem unterstützt wurden. In der ersten Studie wurde der Einfluss von Fehldiagnosen im Training untersucht, während im zweiten Experiment der Einfluss von Systemausfällen im Fokus stand. Die Ergebnisse dieser beiden Studien lassen folgende Schlussfolgerungen zu: 1) *Complacency* stellt im Sinne einer unzureichenden Überwachung eine Ursache für *omission* Fehler und im Sinne einer unzureichenden Überprüfung einen beitragenden Faktor für die Entstehung von *commission* Fehlern dar. 2) Die Erfahrung von Automationsfehlern während des Trainings reduziert das Auftreten von *complacency*, verhindern es jedoch nicht gänzlich. 3) Automationsfehler wirken spezifisch: Während Fehldiagnosen im Training zu einer verbesserten Überprüfung der Diagnosefunktion des Assistenzsystems führen, bleibt die Überwachung der Alarmfunktion dadurch unbeeinflusst. Demgegenüber führen Ausfälle des Assistenzsystems im Training später zu einer vermehrten Überwachung der Alarmfunktion, nicht aber zu einer verbesserten Diagnoseüberprüfung. Dies gilt es bei der Gestaltung von Trainingsmaßnahmen zu berücksichtigen.

ABSTRACT

Over the past few years the increasing reliability of automated systems has led to a decrease of potential sources of error in human-machine interaction. Yet, at the same time, new risks have evolved: Particularly for highly reliable automation, there is a risk of over-trust in automation with complacency and automation bias as potential behavioural consequences. The concept of complacency originates from classical monitoring settings and refers to an insufficient monitoring or cross-checking of the automation which might lead to the miss of critical system states in case of automation failure. In contrast, the concept of automation bias originates from settings where decision aids are used and refers to two different error types: The first one involves so called commission errors, i. e. when operators follow a recommendation of an automated aid even though this recommendation is wrong. The second one has been described as omission error which occurs when operators do not monitor the system and fail to notice problems if the automated aid fails to alert them. Despite the close conceptual relation, complacency and automation bias have only been investigated independently of each other so far. The main objective of the present work was to overcome methodological weaknesses of the available research and to contribute to an empirical clarification of the two concepts complacency and automation bias, and their interrelation. Furthermore, preventive countermeasures in terms of automation failures during training were investigated. Two experiments (each $N = 24$) were conducted with a micro world serving as a task environment. In this task environment the detection, diagnosis and management of occurring faults in a process control task were assisted by an automated decision aid. The first experiment served to investigate the impact of experiences of false diagnoses of this aid whereas the second experiment focussed on the impact of experiences of automation misses during training. The following conclusions can be drawn from the results of these studies:

- 1) Complacency in terms of an insufficient monitoring represents a possible cause for omission errors while complacency in terms of an insufficient verification contributes to the occurrence of commission errors.
- 2) The experience of automation failures during training reduces complacency. However, this intervention does not prevent complacent behaviour completely.
- 3) Automation failures exhibit a specific effect: While false diagnoses during training enhance the operators' verification behaviour with regard to the diagnostic function of the aid, they do not affect the operators monitoring behaviour with regard to the aid's alarm function. Yet, automation misses during training enhance the monitoring of the alarm function but do not affect the verification of the diagnostic function of the decision aid. This specific effect of automation failures has to be considered in the design of training interventions.

INHALTSVERZEICHNIS

ABBILDUNGSVERZEICHNIS	8
TABELLENVERZEICHNIS.....	10
1 THEORETISCHER HINTERGRUND.....	11
1.1 VORBEMERKUNG.....	11
1.2 EINLEITUNG.....	11
1.3 VERTRAUEN IN AUTOMATION.....	15
1.3.1 Begriffliche Einordnung und Grundlagen.....	15
1.3.2 Mangelndes Vertrauen in Automation.....	22
1.3.3 Übersteigertes Vertrauen in Automation.....	25
Exkurs: die Havarie der „Royal Majesty“.....	25
1.4 COMPLACENCY	27
1.4.1 Complacency – Begriffsklärung	27
1.4.2 Entstehungsbedingungen von complacency.....	30
1.4.2.1 Einflüsse von Merkmalen der Automation	30
1.4.2.2 Einflüsse durch Merkmale der Person	34
1.4.2.3 Einflüsse durch Merkmale des situativen Kontextes.....	34
1.4.3 Integratives Rahmenmodell des Konzepts complacency	36
1.4.4 Methodische Probleme bisheriger Studien zum Phänomen complacency.....	38
1.5 AUTOMATION BIAS	40
1.5.1 Definition des Begriffs automation bias.....	40
1.5.2 Empirische Studien des Phänomens automation bias.....	42
1.5.2.1 Automation bias und Verantwortungsübernahme.....	44
1.5.2.2 Automation bias und mögliche Trainingsinterventionen	45
1.5.2.3 Automation bias: Teams versus Einzelpersonen	46
1.5.3 Methodische Probleme der bisherigen Studien zu automation bias.....	48
1.6 INTEGRATION DER KONZEPTE AUTOMATION BIAS UND COMPLACENCY.....	49
1.7 ZIEL- UND FRAGESTELLUNG DER ARBEIT	54
2 EXPERIMENTALUMGEBUNG AUTOCAMS	60
2.1 ALLGEMEINE BESCHREIBUNG DES SYSTEMS	60
2.2 BENUTZEROBERFLÄCHE UND STEUERUNG VON AUTOCAMS.....	61
2.2.1 Schematische Darstellung des Sauerstoff- und Stickstoffsystems	61
2.2.2 CO ₂ -Eingabefeld und Verbindungssymbol.....	62
2.2.3 Manuelle Steuerungsmenüs.....	63
2.2.4 Masteralarm.....	63
2.2.5 Assistenzsystem AFIRA	64
2.2.6 Verlaufsanzeigen	66

2.3	FEHLERTYPEN, -DIAGNOSE UND -MANAGEMENT	66
3	EXPERIMENT I: ZUM EINFLUSS VON FEHLDIAGNOSEN	68
3.1	VORBEMERKUNG.....	68
3.2	HYPOTHESEN EXPERIMENT I.....	68
3.3	METHODE EXPERIMENT I	70
3.3.1	<i>Stichprobe</i>	70
3.3.2	<i>Versuchsplan</i>	70
3.3.3	<i>Material</i>	71
3.3.3.1	Fragebogen zum allgemeinen Vertrauen in Automation.....	71
3.3.3.2	Handout zur Fehlerdiagnose und -behebung.....	72
3.3.4	<i>Durchführung</i>	72
3.3.4.1	Erster Untersuchungstag	73
3.3.4.2	Zweiter Untersuchungstag	74
3.3.5	<i>Abhängige Variablen</i>	77
3.3.5.1	Fehlerdiagnosezeit	77
3.3.5.2	Informationssuchverhalten allgemein in Fehlerdiagnosephasen	77
3.3.5.3	Complacency	78
3.3.5.4	Informationssuchverhalten in fehlerfreien Phasen.....	80
3.3.5.5	Commission Fehler.....	80
3.3.5.6	„Return-to-manual“-Leistung	81
3.3.5.7	Leistung in der prospektiven Gedächtnis- und Reaktionszeitaufgabe.....	81
3.3.5.8	Vertrauen in Automation und Selbstvertrauen.....	81
3.3.6	<i>Statistische Auswertung</i>	82
3.4	ERGEBNISSE EXPERIMENT I	82
3.4.1	<i>Fehlerdiagnosezeit</i>	82
3.4.2	<i>Überprüfung der Diagnosefunktion</i>	83
3.4.2.1	Informationssuchverhalten allgemein in Fehlerdiagnosephasen	83
3.4.2.2	Complacency	84
3.4.3	<i>Überwachung der Alarmfunktion: Informationssuchverhalten in den fehlerfreien Phasen</i> ...	85
3.4.4	<i>Commission Fehler</i>	86
3.4.5	„Return-to-manual“-Leistung	89
3.4.6	<i>Leistung in der prospektiven Gedächtnis- und Reaktionszeitaufgabe</i>	90
3.4.7	<i>Vertrauen in Automation und Selbstvertrauen</i>	92
3.5	DISKUSSION EXPERIMENT I	93
4	EXPERIMENT II: ZUM EINFLUSS VON SYSTEMAUSFÄLLEN	100
4.1	VORBEMERKUNG.....	100
4.2	EINLEITUNG.....	100
4.3	HYPOTHESEN EXPERIMENT II	102
4.4	METHODE EXPERIMENT II.....	103
4.4.1	<i>Stichprobe</i>	103

4.4.2	<i>Versuchsplan</i>	104
4.4.3	<i>Anpassungen der Experimentalumgebung AutoCAMS</i>	105
4.4.4	<i>Material</i>	107
4.4.4.1	Fragebogen zur Zuverlässigkeit des Assistenzsystems.....	107
4.4.4.2	Checkliste für Fehlerdetektion, -diagnose und -behebung.....	109
4.4.4.3	Handout zur Fehlerdiagnose und -behebung.....	110
4.4.5	<i>Durchführung</i>	110
4.4.5.1	Erster Untersuchungstag.....	110
4.4.5.2	Zweiter Untersuchungstag.....	111
4.4.6	<i>Abhängige Variablen</i>	114
4.4.6.1	Omission Fehler.....	114
4.4.6.2	Commission Fehler.....	114
4.4.6.3	Informationssuchverhalten in fehlerfreien Phasen.....	114
4.4.6.4	Fehlerdiagnosezeit.....	115
4.4.6.5	Informationssuchverhalten allgemein in Fehlerdiagnosephasen.....	115
4.4.6.6	Complacency bezüglich der Diagnosefunktion.....	115
4.4.6.7	„Return-to-manual“-Leistung.....	115
4.4.6.8	Leistung in der prospektiven Gedächtnis- und Reaktionszeitaufgabe.....	116
4.4.6.9	Wahrgenommene Zuverlässigkeit des Assistenzsystems und Einschätzung der eigenen Leistung.....	116
4.5	ERGEBNISSE EXPERIMENT II.....	117
4.5.1	<i>Omission Fehler</i>	117
4.5.2	<i>Commission Fehler</i>	118
4.5.3	<i>Überwachung der Alarmfunktion: Informationssuchverhalten in fehlerfreien Phasen</i>	119
4.5.4	<i>Fehlerdiagnosezeit</i>	120
4.5.5	<i>Überprüfung der Diagnosefunktion</i>	121
4.5.5.1	Informationssuchverhalten allgemein in Fehlerdiagnosephasen.....	121
4.5.5.2	Complacency bezüglich der Diagnosefunktion.....	122
4.5.6	„Return-to-manual“-Leistung.....	124
4.5.7	Leistung in der prospektiven Gedächtnis- und Reaktionszeitaufgabe.....	125
4.5.8	Wahrgenommene Zuverlässigkeit des Assistenzsystems und Einschätzung der eigenen Leistung.....	128
4.5.8.1	Informationsgruppe vs. Erfahrungsgruppe.....	128
4.5.8.2	Omission Fehler ja vs. nein.....	132
4.5.8.3	Commission Fehler ja versus nein.....	132
4.6	DISKUSSION EXPERIMENT II.....	135
5	GESAMTDISKUSSION	142
	LITERATURVERZEICHNIS	150
	ANHANG	159

ABBILDUNGSVERZEICHNIS

Abb. 1: Die Beziehung zwischen Vertrauen und Leistungsvermögen der Automation (aus Lee & See, 2004, S. 55).....	21
Abb. 2: Benutzeroberfläche der <i>multi-attribute task</i> -Batterie (zur Verfügung gestellt von der Technischen Universität Berlin, Fachgebiet Arbeits-, Ingenieur- und Organisationspsychologie) 31	
Abb. 3: Ergebnisse der Studie von Parasuraman, Molloy und Singh (1993, S. 11). Einfluss von Reliabilitätshöhe und Konstanz der Reliabilität auf die Wahrscheinlichkeit, Fehler der Automation zu entdecken.	32
Abb. 4: Rahmenmodell des Konzepts <i>complacency</i> (aus Manzey & Bahner, 2005, S. 103).....	36
Abb. 5: Complacency im Kontext von Überwachung und Assistenzsystemnutzung.....	51
Abb. 6: Integration der Konzepte <i>complacency</i> und <i>automation bias</i>	52
Abb. 7: Benutzeroberfläche von AutoCAMS: a) schematische Darstellung des Sauerstoff- und Stickstoffsystems, b) Bildschirmbereich für Sekundäraufgaben, c) manueller Steuerungsbereich, d) Masteralarm, e) Assistenzsystem AFIRA, f) Verlaufsanzeigenbereich.....	61
Abb. 8: Schematische Darstellung des Sauerstoff- und Stickstoffsystems	62
Abb. 9: Steuerungsmenü und Standards Gasströme	63
Abb. 10: Menü zur Auswahl und Aktivierung eines Reparaturauftrags.....	64
Abb. 11: AFIRA-Rückmeldungen nach einer erfolgreichen bzw. nicht erfolgreichen Fehlerreparatur.....	65
Abb. 12: Exp. 1: Fehlerdiagnosezeit (experimentelle Gruppen)	83
Abb. 13: Exp. I: Anteil relevanter Informationsabrufe in Fehlerdiagnosephasen (experimentelle Gruppen).....	84
Abb. 14: Exp. I: <i>Complacency</i> (experimentelle Gruppen)	84
Abb. 15: Exp. I: Anteil relevanter Informationsabrufe in fehlerfreien Phasen (experimentelle Gruppen).....	85
Abb. 16: Exp. I: <i>Complacency</i> (<i>commission</i> Fehler ja/nein)	86
Abb. 17: Exp. I: Fehlerdiagnosezeit (<i>commission</i> Fehler ja/nein)	87
Abb. 18: Exp. I: Informationsabrufe in Fehlerdiagnosephasen (<i>commission</i> Fehler ja/nein)	88
Abb. 19: Exp. I: Informationsabrufe in fehlerfreien Phasen (<i>commission</i> Fehler ja/nein)	88
Abb. 20: Exp. I: Fehlerdiagnosezeit beim Systemausfall im Vergleich zum Training (experimentelle Gruppen).....	89
Abb. 21: Exp. I: Leistung in der Reaktionszeitaufgabe in fehlerfreien Phasen und Fehlerphasen (<i>commission</i> Fehler ja/nein).....	91
Abb. 22: Exp. I: Vertrauen in Automation (experimentelle Gruppen)	92
Abb. 23: Exp. I: Selbstvertrauen (experimentelle Gruppen)	93
Abb. 24: AutoCAMS mit Fehlermodus-Button (links: Fehlermodus aktiviert – „Fehler“, rechts: Fehlermodus deaktiviert – „OK“).....	106
Abb. 25: Exp. II: <i>Omission</i> Fehler beim ersten und zweiten Systemausfall (experimentelle Gruppen).....	117
Abb. 26: Exp. II: Überprüfung der Fehldiagnose (Fehler 14) und <i>commission</i> Fehler	118
Abb. 27: Exp. II: Informationsabrufe in fehlerfreien Phasen (experimentelle Gruppen)	119

Abb. 28: Exp. II: Anteil relevanter Informationsabrufe in fehlerfreien Phasen (<i>commission</i> Fehler ja/nein)	120
Abb. 29: Exp. II: Informationsabrufe in Fehlerdiagnosephasen (<i>commission</i> Fehler ja/nein).....	122
Abb. 30: Exp. II: <i>Complacency</i> (experimentelle Gruppen)	123
Abb. 31: Exp. II: <i>Complacency</i> (<i>commission</i> Fehler ja/nein).....	123
Abb. 32: Exp. II: Return-to-manual-Leistung (experimentelle Gruppen).....	124
Abb. 33: Exp. II: Leistung in der prospektiven Gedächtnisaufgabe (fehlerfreie vs. Fehlerphasen)	126
Abb. 34: Exp. II: Leistung in der prospektiven Gedächtnisaufgabe für fehlerfreie und Fehlerphasen (<i>commission</i> Fehler ja/nein)	126
Abb. 35: Exp. II Leistung in der Reaktionszeitaufgabe in fehlerfreien vs. Fehlerphasen (<i>omission</i> Fehler ja/nein)	127
Abb. 36: Exp. II: Leistung in der Reaktionszeitaufgabe (<i>commission</i> Fehler ja/nein)	128
Abb. 37: Exp. II: Wahrgenommene Zuverlässigkeit der Alarmfunktion von AFIRA (experimentelle Gruppen).....	129
Abb. 38: Exp. II: Wahrgenommene Zuverlässigkeit der Diagnosefunktion von AFIRA (experimentelle Gruppen).....	130
Abb. 39: Exp. II: Wahrgenommene Zuverlässigkeit der Fehlermanagementempfehlungen von AFIRA (experimentelle Gruppen).	130
Abb. 40: Exp. II: Selbsteinschätzung der Fehlerdetektionsleistung (experimentelle Gruppen)...	131
Abb. 41: Exp. II: Selbsteinschätzung der Fehlerdiagnoseleistung (experimentelle Gruppen).....	131
Abb. 42: Exp. II: Selbsteinschätzung der Fehlermanagementleistung (experimentelle Gruppen)	132
Abb. 43: Exp. II: Wahrgenommene Zuverlässigkeit der Alarmfunktion von AFIRA (<i>commission</i> Fehler ja/nein)	133
Abb. 44: Exp. II: Selbsteinschätzung der Fehlerdiagnoseleistung (<i>commission</i> Fehler ja/nein).....	134
Abb. 45: Exp. II: Selbsteinschätzung der Fehlermanagementleistung (<i>commission</i> Fehler ja/nein)	134

TABELLENVERZEICHNIS

Tab. 1: Exp. I: CAMS- und AutoCAMS-Fragebogen.....	72
Tab. 2: Exp. I: Ablauf des CAMS-Trainings mit Blick auf auftretende CAMS-Fehlfunktionen ...	73
Tab. 3: Exp. I: Ablauf des AutoCAMS-Trainings mit Blick auf auftretende CAMS-Fehlfunktionen	75
Tab. 4: Exp. I: Ablauf der Testphase mit Blick auf auftretende CAMS-Fehlfunktionen	76
Tab. 5: Zur Diagnoseüberprüfung notwendige Prüfelemente	78
Tab. 6: Exp. II: AutoCAMS-Fragebogen.....	108
Tab. 7: Exp. II: Ablauf des CAMS-Trainings mit Blick auf auftretende CAMS-Fehlfunktionen	111
Tab. 8: Exp. III: Ablauf des AutoCAMS-Trainings mit Blick auf auftretende CAMS-Fehlfunktionen	112
Tab. 9: Exp. II: Ablauf der Testphase mit Blick auf auftretende CAMS-Fehlfunktionen.....	113
Tab. 10: Exp. II: Häufigkeitsverteilung von <i>commission</i> und <i>omission</i> Fehlern.....	118

1 THEORETISCHER HINTERGRUND

1.1 VORBEMERKUNG

Das thematische Feld dieser Arbeit, ein übersteigertes Vertrauen in automatisierte Systeme und seine Auswirkungen, wurde bislang fast ausschließlich im angloamerikanischen Raum wissenschaftlich untersucht. Damit verbunden ist, dass für viele zentrale Fachtermini im Deutschen keine äquivalent konnotierten Entsprechungen existieren. Aus diesem Grund werden in der vorliegenden Arbeit – soweit eine direkte Übersetzung nicht ohne Bedeutungsverlust möglich ist – die englischen Fachbegriffe verwendet.

1.2 EINLEITUNG

In modernen Mensch-Maschine-Systemen, wie man sie beispielsweise in Leitwarten der chemischen Industrie, in Flugzeugcockpits oder auch im Operationssaal vorfindet, wurden über die letzten zwei Jahrzehnte immer mehr Funktionen, die früher ausschließlich vom Menschen ausgeführt werden konnten, automatisiert. Von großer Vielfalt sind dabei die Aufgaben, die an die Maschine delegiert werden. Entsprechend lassen sich verschiedene Typen und Stufen automatisierter Systeme unterscheiden, die von nahezu manueller Steuerung bis hin zur vollautomatisierten Entscheidungsfindung und Ausführung von Aufgaben reichen (Parasuraman, Sheridan & Wickens, 2000; Wandke, 2005). Die Weiterentwicklung der technischen Möglichkeiten spiegelt sich dabei auch in der Gestaltung automatisierter Systeme wider: Während Automatisierungsbestrebungen früher vor allem auf die Übertragung manueller Regelungs- und Steuerungsfunktionen abzielten, werden heute vermehrt auch kognitive Funktionen der Urteils- und Entscheidungsfindung dem technischen System übertragen. Mitbegründet ist dies in der zunehmenden Leistungsfähigkeit und Komplexität technischer Systeme. Um diese für den Menschen überhaupt noch handhabbar zu machen, ist oftmals die Vorverarbeitung und Integration verschiedenster Daten aus unterschiedlichen Quellen Voraussetzung.

Mit der zunehmenden Automatisierung verbindet sich die Hoffnung, zum einen die Beanspruchung des Operators und den erforderlichen Trainingsaufwand zu reduzieren, zum anderen aber auch die Verlässlichkeit des Mensch-Maschine-Systems sowie dessen Wirtschaftlichkeit zu steigern. Oftmals wird davon ausgegangen, dass es möglich wäre, nahezu vollständig autonome Systeme zu schaffen, die keine oder allenfalls nur sehr wenige Eingriffe durch den Menschen erfordern – und folglich der Mensch als potenzielle Fehlerquelle ausgeschaltet werden könne.

Zugrunde liegt dabei die Annahme, dass eine früher vom Menschen ausgeführte Funktion durch eine Automation ersetzt werden könne ohne dadurch das System, innerhalb dessen die verschiedenen Funktionen – teils vom Menschen, teils durch Automation – auszurichten sind, nennenswert zu beeinträchtigen (Sarter, Woods & Billings, 1997). Eine solche Zerlegung eines komplexen Systems in voneinander unabhängige Teilfunktionen greift jedoch zu kurz. Die angestrebten Ziele – Entlastung der Operateure, höhere Zuverlässigkeit und geringere Kosten – werden nicht immer erreicht (Manzey, in Druck). Tragischer Beleg hierfür sind Unfälle und Katastrophen im Kontext technischer Systeme, bei denen als eine Ursache immer wieder eine nicht geglückte Interaktion zwischen Mensch und automatisiertem System verantwortlich gemacht wird.

Die vorliegende Arbeit soll einen Beitrag zur Aufklärung von Problemen dieser Interaktion leisten und damit auch mögliche Ansatzpunkte für ein verbessertes Zusammenspiel von Menschen und automatisierten Systemen liefern. Doch welche Probleme treten bei der Nutzung automatisierter Systemen im Detail auf? Um diese Frage beantworten zu können, ist zunächst zu klären, was unter einem automatisierten System bzw. dem als Synonym verwendeten Begriff Automation zu verstehen ist. Prinzipiell stellt Automation das Resultat eines Automatisierungsprozesses dar (Hauß & Timpe, 2000), der darin besteht, bestimmte Funktionen oder Aufgaben auf die Maschine zu verlagern. Aus psychologischer Perspektive relevant sind dabei solche Funktionen oder Aufgaben, die prinzipiell auch vom Menschen ausgeführt werden können. Wenn eine Rückdelegation der Funktion zurück auf den Menschen nicht möglich ist sondern permanent bei der Maschine verbleibt, wie etwa beim „automatischen“ Anlasser im Kraftfahrzeug oder auch bei Personenaufzügen, ist dies nach dem hier zugrunde gelegten Begriffsverständnis nicht als Automation zu werten (Parasuraman & Riley, 1997). In der vorliegenden Arbeit soll Automation stattdessen als „...any sensing, detection, informationprocessing, decision-making, or control action that could be performed by humans but is actually performed by machine“ (Moray, Inagaki & Itoh, 2000, S. 44) verstanden werden. Legt man diese Definition zugrunde, ist die Einführung von Automation zwangsläufig mit einer Anforderungsverschiebung für den Operator hin zu mehr Überwachungs- und Diagnosetätigkeiten verbunden. Die primäre Aufgabe des Operators besteht nunmehr in der „leitenden Kontrolle“ (*supervisory control*, Sheridan, 1987) des Systems und damit darin, nur dann einzugreifen und die eigentlich automatisierten Funktionen wieder selbst zu übernehmen, wenn das System fehlerhaft arbeitet oder ausfällt. Der Operator ist vom aktiven Reglungsprozess ausgeschlossen, bleibt also bezüglich des zugrunde liegenden Prozesses außerhalb des Rückkopplungskreislaufs („*out-of-the-loop*“, Wickens & Hollands, 2000).

Damit zeichnet sich ein erster Problembereich der Mensch-Automation-Interaktion ab. Wickens und Hollands (2000) nennen folgende mögliche negative Folgen der Abkopplung des Ope-

rateurs von der direkten Steuerung und Kontrolle und fassen diese mit dem Begriff *out-of-the-loop-unfamiliarity* (OOTLUF) zusammen: Als erste mögliche negative Folge verweisen die Autoren auf die Abnahme des Situationsbewusstseins. Diese steht in direktem Zusammenhang mit der Veränderung der Rolle des Menschen hin zum passiven Informationsempfänger. Situationsbewusstsein wird von Endsley (1995) als dreistufiger, hierarchisch aufgebauter Prozess aus erstens der Wahrnehmung relevanter Informationen, zweitens deren Verständnis und richtiger Interpretation und drittens der darauf aufbauenden Antizipation künftigen Geschehens definiert. Für den Aufbau eines adäquaten Situationsbewusstseins ist ein in seiner Frequenz der Dynamik des Systems entsprechender Abgleich des aktuellen Systemzustandes mit dem mentalen Modell des Operateurs erforderlich. Erfolgt dieser nicht, beispielsweise weil der Operateur dem System zu sehr vertraut oder aber aufgrund von Vigilanzproblemen, führt dies unweigerlich zu einer Abnahme des Situationsbewusstseins. Doch auch unzureichende Rückmeldungen und Transparenz der Automation sowie ein mangelndes Systemverständnis können sich negativ auf die Aufrechterhaltung des Situationsbewusstseins auswirken (Endsley, Bolté & Jones, 2003).

Als zweite negative Folge nennen Wickens und Hollands (2000) den Verlust von Fertigkeiten, der sich daraus ergibt, dass die Veränderung der Rolle des Menschen vom aktiv Ausführenden zum passiv Beobachtenden immer auch einen Verlust des kontinuierlichen manuellen Trainings der entsprechenden Fertigkeiten bedeutet. Wenig überraschend und eine natürliche Folge ist, dass die Fertigkeit zur Ausübung der entsprechenden Funktion mit der Zeit nachlässt und sich Effektivitäts- und Effizienzeinbußen ergeben. Solange die Automation fehlerfrei funktioniert, mag das wenig problematisch sein. Fatale Folgen kann dies dagegen nach sich ziehen, wenn der Mensch bei einem Ausfall oder Fehlfunktionen der Automation plötzlich gezwungen ist, eigentlich automatisierte Funktionen wieder selbst zu übernehmen. Hierin offenbart sich eine der Ironien von Automatisierung (Bainbridge, 1983): Während Automation im Normalfall die manuellen Fertigkeiten des Menschen überflüssig macht, und auch gerade deshalb zum Einsatz gelangt, wird der Mensch im Notfall zur alles entscheidenden Instanz.

Als dritte negative Folge wird ein zu starkes Verlassen auf das automatisierte System angeführt, was in der englischsprachigen Literatur unter dem Schlagwort *complacency* gefasst wird. Dabei handelt es sich um ein Phänomen, das eng mit einem übersteigerten Vertrauen in die Funktions- und Leistungsfähigkeit der Automation verbunden ist und sich in einer unzureichenden Überwachung automatisierter Systeme äußert.

Neben der Bündelung von Automationsproblemen unter dem Stichwort OOTLUF gibt es auch eine ganze Reihe anderer Versuche, mögliche unerwünschte Effekte von Automation zu strukturieren. Sheridan und Parasuraman (2006) orientieren sich etwa an den Hauptursachen für

automationsbezogene Vorfälle und Unfälle (vgl. Funk et al., 1999) und heben als zentrale Problembereiche Rückmeldungen über den Systemstatus, mangelndes Verständnis bezüglich der Funktionsweise der Automation und *overreliance* hervor. Der zuletzt genannte Problembereich bezieht sich dabei wiederum darauf, dass der Operateur sich zu sehr auf die Automation verlässt.

Parasuraman und Riley (1997) beschreiben etwas allgemeiner die drei Phänomene *automation abuse*, *misuse* und *disuse* als problematische Umgangsweisen mit Automation. Während *abuse* als unangemessener Einsatz von Automation durch Entwickler oder Manager definiert wird, beziehen sich *disuse* und *misuse* auf die Nutzung des Systems durch den Anwender oder Operateur. *Disuse* wird als „underutilization of automation“ (z.B. Ignorieren oder Ausschalten von automatisierten Alarm- und Sicherheitssystemen) beschrieben und *misuse* als „overreliance on automation“ (z.B. Nutzung des Systems, wenn es nicht genutzt werden sollte oder unzureichende Überwachung des Systems) definiert.

Sarter et al. (1997) beschreiben sieben, bei der Entwicklung der Systeme meist nicht bedachte Probleme der Mensch-Automation-Interaktion. Diese sind eine ungleiche Verteilung der Beanspruchung statt einer Reduktion derselben, neue Aufmerksamkeits-, Wissens- sowie Koordinationsanforderungen, mangelnde *mode awareness*, die Entstehung neuer Fehlerquellen statt einer prinzipiellen Reduktion derselben und die Notwendigkeit neuer Trainingsansätze sowie schließlich, auch hier, *complacency* als nicht erwartete negative Automationsfolge.

Deutlich wird, dass in all diesen Strukturierungsversuchen ein Problembereich immer wieder genannt, wenn auch mit unterschiedlichen Begriffen belegt wird. Sei es *complacency*, *overreliance* oder *misuse*: Ein wesentlicher Problembereich der Mensch-Automation-Interaktion wird in einem unangemessen hohen Vertrauen in die Automation und den damit verbundenen Verhaltenskonsequenzen gesehen.

Ein solches übersteigertes Vertrauen in Automation ist Gegenstand der vorliegenden Arbeit, wobei insbesondere zwei damit verbundene Phänomene, *complacency* und *automation bias*, näher betrachtet werden sollen. *Complacency* kann sich, wie bereits erwähnt, in einer unzureichenden Überwachung automatisierter Systeme äußern. Unter dem Begriff *automation bias* werden zwei Fehlertypen zusammengefasst, die ebenfalls mit einer unzureichenden Überwachung bzw. Überprüfung zusammenhängen. Allerdings ist der Anwendungskontext hier auf Entscheidungsassistenzsysteme eingegrenzt. So genannte *commission* Fehler bestehen darin, dass Operateure einer falschen Direktive eines solchen Assistenzsystems trotz prinzipieller Verfügbarkeit widersprechender Informationen folgen. Der zweite Fehlertyp, als *omission* Fehler bezeichnet, besteht darin, dass ein Operateur eine tatsächlich vorliegende kritische Situation übersieht, sofern diese vom Assistenzsystem nicht angezeigt wird. In der vorliegenden Arbeit sollen diese Phänomene bei der

Nutzung eines Assistenzsystems im Anwendungskontext der Prozesssteuerung untersucht werden. Zentrale Ziele der Arbeit sind dabei zum einen eine empirisch fundierte Klärung der Phänomene sowie ihres Bezugs untereinander und zum anderen die Untersuchung möglicher Gegenmaßnahmen.

Im Folgenden wird zur Einbettung der Thematik zunächst generell auf Vertrauen in Automation eingegangen (Kap. 1.3). Unterschieden werden dabei ein mangelndes und ein übersteigertes Vertrauen in Automation. Beide Seiten werden vorgestellt, wobei der Schwerpunkt, dem Gegenstand der Arbeit entsprechend, auf theoretischen Überlegungen und empirischen Studien zum übersteigerten Vertrauen liegen wird. Sodann werden die Konzepte *complacency* (Kap. 1.4) und *automation bias* (Kap. 1.5) als mögliche Verhaltenskorrelate eines übersteigerten Vertrauens vorgestellt und definiert sowie der diesbezügliche Stand der Forschung berichtet. Auch wird auf methodische Schwächen bisheriger empirischer Untersuchungen der Phänomene eingegangen (Kap. 1.4.4 bzw. 1.5.3). In Kapitel 1.6 wird der Versuch einer Integration der beiden Konzepte gewagt und ein integratives Rahmenmodell der Konzepte vorgestellt. Zum Abschluss des Theorieteils werden Frage- und Zielstellung konkretisiert (Kap. 1.7)

1.3 VERTRAUEN IN AUTOMATION

1.3.1 Begriffliche Einordnung und Grundlagen

Die Anfänge der empirischen Untersuchung des Vertrauens in Automation gehen auf Muir (1989, berichtet in Muir & Moray, 1996) zurück. Seitdem haben sich eine Vielzahl von Studien und mittlerweile auch Überblicksarbeiten der Thematik gewidmet (Muir, 1994; Kelly, Boardman, Goillau & Jeannot, 2003; Lee & See, 2004; Manzey & Bahner, 2005; Madhavan & Wiegmann, 2007a). Dennoch lässt sich auch heute noch kein geteiltes Verständnis bezüglich der Definition und der Entstehung von Vertrauen in Automation feststellen.

Ein Grund hierfür mag darin liegen, dass der Begriff des Vertrauens, auch losgelöst vom Vertrauensgegenstand, kein einheitlich verwandtes Konzept darstellt. Luhmann (2000) beschreibt aus soziologischer Perspektive Vertrauen als Mechanismus zur Reduktion (sozialer) Komplexität und liefert damit einen Hinweis darauf, warum Vertrauen gerade in immer komplexer werdenden Mensch-Maschine-Systemen eine so bedeutende Rolle zukommt.

Aus psychologischer Perspektive finden sich vermehrt Definitionsansätze, die Vertrauen als Persönlichkeitseigenschaft bzw. generalisierte Erwartung beschreiben (Rotter, 1967; Rempel, Holmes & Zanna; 1985, Barber, 1983). Gemein ist diesen Ansätzen, dass Vertrauen als Einstel-

lung oder Erwartung bezüglich der Wahrscheinlichkeit erwünschter Reaktionen gesehen wird. Ein zweiter Ansatz besteht darin, Vertrauen als Intention oder Handlungsbereitschaft zu konzeptionalisieren (Johns, 1996; Moorman, Deshpande & Zaltman, 1993). Die meist zitierte Definition, die hier zuzuordnen ist, stammt von Mayer, Davis und Schoorman (1995). Vertrauen wird dort als Bereitschaft eine Interaktion einzugehen verstanden, in der die Handlungen des Interaktionspartners potenziell zum Schaden des Vertrauenden gereichen könnten, und zwar unabhängig von der Möglichkeit den Interaktionspartner kontrollieren zu können.

Eine dritte Gruppe von Autoren (vgl. auch Deutsch, 1960; Kramer, 1999) sieht Vertrauen nicht mehr als Einstellung oder Intention sondern siedelt dieses bereits als Verhaltensergebnis an, wie etwa auch Meyer (2001), der *reliance* und *compliance* als zwei verschiedene Aspekte von Vertrauen im Kontext des Umgangs mit Alarmen unterscheidet. *Reliance* bezieht sich dabei auf das Verlassen darauf, dass in Abwesenheit eines Alarms tatsächlich kein kritischer, handlungsbedürftiger Zustand vorliegt. Demgegenüber beschreibt *compliance* inwiefern sich ein Operateur bei Vorliegen eines Alarms darauf verlässt, dass auch faktisch ein kritischer Zustand vorliegt und dem in seiner unmittelbaren Verhaltensreaktion auf den Alarm Ausdruck verleiht.

Allerdings macht gerade die zuletzt genannte Konzeptionalisierung auch das Problem der Abgrenzung von Vertrauen in Automation gegenüber verwandten Konzepten wie *automation reliance* oder auch *automation use* deutlich. Gegen eine Gleichsetzung spricht, dass das Verlassen auf ein technisches System und daraus folgende beobachtbare Handlungen, im Sinne einer Aktivierung (bzw. Nutzung) oder Deaktivierung (bzw. Nicht-Nutzung) der Automation, nicht ausschließlich durch das Vertrauen in das System determiniert sind, sondern auch anderen Einflussfaktoren wie z.B. Arbeitsbelastung, Situationsbewusstsein oder Selbstvertrauen unterliegen. Vor diesem Hintergrund scheint es plausibel, dem Ansatz von Lee und See (2004) zu folgen. Sie schlagen unter Zuhilfenahme der Theorie des überlegten Handelns von Ajzen und Fishbein (1980) vor, Vertrauen noch vor der Intention zu Handeln anzusiedeln und definieren Vertrauen allgemein als „the attitude that an agent will help achieve an individual’s goals in a situation characterized by uncertainty and vulnerability“ (Lee & See, 2004, S. 54). Bezogen auf den Automationskontext beeinflusst dieses Vertrauen – zusammen mit anderen Faktoren wie z.B. der Beanspruchung oder dem wahrgenommenen Risiko – die Intentionsbildung und diese wiederum zusammen mit Leistungseigenschaften der Automation und zeitlichen Rahmenbedingungen das tatsächliche Verhalten (*reliance action*). Vertrauen leitet also das Verhalten, determiniert es aber nicht vollständig.

Neben der Frage nach einer angemessenen Definition von Vertrauen im Kontext der Mensch-Automation-Interaktion stellt sich auch die Frage nach der Entwicklung bzw. der

Grundlage von Vertrauen. Als frühe und auch heute noch oft zitierte Ansätze zur Beantwortung dieser Frage sind die Vorschläge von Sheridan (1988) und Rempel et al. (1985) zu nennen. Sheridan (1988) vermutet als mögliche, das Vertrauen bestimmende Aspekte, die sieben Faktoren *reliability*, *robustness*, *familiarity*, *understandability*, *explication of intention*, *usefulness* und *dependence*.

Rempel et al. (1985) beschreiben die (interpersonale) Vertrauensentstehung als Prozess, der in drei Stufen verläuft. An erster Stelle steht *predictability*, gefolgt von *dependability* und schließlich *faith*. Dabei wird davon ausgegangen, dass *predictability* als das Ausmaß in dem künftiges Verhalten antizipiert oder vorhergesagt werden kann, die Basis des Vertrauens in der noch jungen interpersonalen Beziehung bildet. *Dependability* folgt als nächstes und bezieht sich auf das Ausmaß in dem das Verhalten des Gegenübers als konsistent erlebt wird. Im weiteren Verlauf der Beziehung verlagert sich die Vertrauensbasis dann auf *faith*, was dem etwas weniger greifbaren Gesamturteil, dass man sich auf den Interaktionspartner verlassen kann, entspricht und in stärkerem Maße eine zukünftige Perspektive einschließt. Spätere Übertragungsversuche auf die Mensch-Automation-Interaktion bedienten sich meist dieses Modells als Anleihe, wobei jedoch der Gedanke an in dieser Reihung aufeinander folgende Phasen meist aufgegeben wurde (Muir, 1994; Muir & Moray, 1996; Lerch, Prietula & Kulik, 1997; Lee & Moray, 1992; Lee & See, 2004; Madhavan & Wiegmann, 2007a).

Zum Teil wurde auch der Versuch unternommen, allgemeinere und besser zum Automationskontext passende Begrifflichkeiten zu finden. So benennen Lee und Moray (1992; auch Lee & See 2004) *performance*, *process* und *purpose* als essentielle Vertrauensdimensionen.

Performance ist die Antwort auf die Frage, „was“ die Automation macht. Im Fokus steht daher die bisherige und gegenwärtige Leistung der Automation, wie etwa *robustness* und *reliability* (Sheridan, 1988) oder auch *predictability* (Rempel et al., 1985).

Process bezieht sich darauf „wie“ die Automation arbeitet und beschreibt das Ausmaß, in dem die Algorithmen der Automation der Situation angemessen und geeignet sind, die Ziele des Operators zu erreichen. Hierin zeigt sich eine Parallele zu *understandability*, *familiarity* und *dependability* (Sheridan, 1988, bzw. Rempel et al., 1985).

Purpose beschreibt, „warum“ die Automation entwickelt wurde und ist dadurch gekennzeichnet, inwieweit die Automation so, wie vom Entwickler intendiert, genutzt wird bzw. genutzt werden kann. Hier zeigt sich eine konzeptuelle Nähe zu den zwei von Sheridan (1988) genannten Faktoren *explication of intention* und *usefulness* als auch, etwas weiter gefasst, zu *faith* im Sinne von Rempel et al. (1985).

Der letzt genannte Aspekt (*purpose*) zeigt möglicherweise am deutlichsten das grundlegende Problem dieser und vergleichbarer Versuche, Grundfaktoren des Vertrauens in Automation zu identifizieren. Nahe liegender Weise wird dabei vom Vertrauen in interpersonalen Beziehungen ausgegangen und versucht, dieses auf die Mensch-Automation-Interaktion zu übertragen. Während bei der „Mensch-Mensch-Interaktion“ die dem Gegenüber unterstellte Motivation seines Handelns zweifelsohne eine zentrale Rolle spielt und in der Regel auch durch mehr oder weniger offensichtliche Hinweisreize eingeschätzt werden kann, ist dies in der Mensch-Automation-Interaktion nur sehr bedingt der Fall. Der Vorschlag von Lee (Lee & Moray 1992; Lee & See, 2004), diese motivationale oder intentionale Komponente nicht direkt beim (technischen) Interaktionspartner, sondern beim (menschlichen) Entwickler anzusiedeln, leuchtet ein. Allerdings schließt sich hier die Frage an, inwiefern die Intentionen des Entwicklers bezüglich der zu entwickelnden Automation ähnlich vielschichtig und fein abgestuft sein mögen, wie die Motivationslage eines Gegenübers in interpersonalen Beziehungen. Ein Entwickler wird in aller Regel wenig motiviert sein, eine Automation zu schaffen, die dem Operateur nicht nutzt oder gar schaden kann – während eine solche Motivation in interpersonalen Situationen durchaus auftreten kann. Unterschiedliche Entwicklungsintentionen werden sich somit allenfalls in einer variierenden Eignung des Entwicklungsprodukts für die jeweilige Einsatzsituation widerspiegeln. Liegen dem Operateur nun Anzeichen dafür vor, dass die Automation für die jeweilige Situation nicht ausgelegt ist, wird das sicherlich seine Bereitschaft, diese zu nutzen, negativ beeinflussen. Ein Beispiel hierfür findet sich etwa bei Hockey und Maule (1995). Sie berichten von Operateuren, die die eigentlich zuverlässig arbeitende Prozesssteuerung einer chemischen Anlage immer wieder abschalteten und die Funktion manuell ausführten obwohl dies nicht vorgesehen war. Hintergrund hierfür war, dass die Operateure zum einen die Geschwindigkeit der automatischen Steuerung für zu gering erachteten und zum anderen die Funktionalität der automatischen Steuerungen nicht zu den für sie gewohnten Strategien passte. Offenbar nahmen die Operateure die manuelle Steuerung als die bessere Alternative gegenüber der als suboptimal wahrgenommenen automatischen Prozesssteuerung wahr – trotz der aus der manuellen Steuerung resultierenden Erhöhung von Fehlerrisiko und Beanspruchung. Ob eine solche Nicht-Nutzung der Automation allerdings vermittelt über eine Vertrauensminderung geschieht oder aber eine eingeschränkte Eignung der Automation direkt und unabhängig vom Vertrauen auf die Nutzungsbereitschaft wirkt, ist zumindest fraglich. Möglicherweise sind am Vertrauen in interpersonalen Beziehungen auch in sehr viel stärkerem Maß emotionale Prozesse beteiligt als dies in Bezug auf automatisierte Interaktionspartner der Fall ist.

Neben diesen theoretischen Überlegungen zur Frage der Übertragbarkeit des Vertrauens in interpersonalen Beziehungen auf das Vertrauen in Automation gibt es auch einige empirische

Befunde, die zumindest eine direkte Übertragung als unzulässig erscheinen lassen. So fanden Dzindolet, Pierce, Beck und Dawe (2002) in einer Experimentalserie, deren Gegenstand eine visuelle Detektionsaufgabe war, bei der die Probanden vermeintlich entweder durch eine Automation oder einen Kollegen unterstützt wurden, folgende Unterschiede: Einerseits schätzten die Probanden die Leistung der Automation von vornherein höher ein als die menschliche Leistung, vertrauten dieser also mehr. Andererseits entschieden sich die Probanden später eher gegen die Nutzung der Automation als gegen die Unterstützung durch einen menschlichen Kooperationspartner. Dzindolet et al. (2002) erklären dies damit, dass die Probanden möglicherweise das Schema einer perfekt arbeitenden Automation verinnerlicht hatten (*perfect automation scheme*). Diese hohe Erwartungshaltung an die Automation scheint dazu zu führen, dass spätere Fehler derselben sich umso stärker auswirken, während menschlichen Interaktionspartnern Fehler offenbar eher nachgesehen werden. Die Existenz eines solchen *positive bias* gegenüber automatisierten Systemen wird durch Befunde von Dijkstra, Liebrand und Timminga (1998), Dzindolet, Peterson, Pomranky, Pierce und Beck (2003) sowie Madhavan und Wiegmann (2007b) gestützt. In eine ähnliche Richtung weisen auch Ergebnisse von Jian, Bisantz und Drury (2000), denen zufolge die Vertrauensurteile von Menschen gegenüber Maschinen extremer ausfallen als jene gegenüber Menschen, wobei dies insbesondere beim Misstrauen der Fall zu sein scheint. Offenbar ist es leichter, einer Maschine als einem Menschen sein Vertrauen zu entziehen.

Einen weiteren Unterschied zwischen dem Vertrauen in Automation und menschliche Interaktionspartner zeigten Lewandowsky, Mundy und Tan (2000). Sie fanden, dass die Delegation einer Aufgabe (Steuerung der Pumpen in einer Pasteurisierungsfabrik) an einen menschlichen Kooperationspartner in stärkerem Maße von der Einschätzung der eigenen Vertrauenswürdigkeit der Probanden bestimmt wird als die Delegation an die Automation. In derselben Studie zeigte sich auch, dass der Zusammenhang zwischen Vertrauen in den (entweder menschlichen oder automatisierten) Kooperationspartner und der Delegationsentscheidung für die Automationsbedingung höher ausfiel. Dies könnte ein Hinweis darauf sein, dass dem Vertrauen in der Mensch-Automation-Interaktion sogar eine vergleichsweise höhere Bedeutung für tatsächliches Handeln zukommt, es sich aber zugleich auch weniger komplex zusammensetzt. Für eine geringere Komplexität der Vertrauensentstehung im Automationskontext sprechen auch die Ergebnisse von Lerch et al. (1997): Während die Probanden die Vertrauenswürdigkeit einer Automation in erster Linie auf Grundlage objektiver Leistungsdaten beurteilten, fanden bei der Beurteilung des menschlichen Kooperationspartners auch andere Aspekte, wie vermutete Anstrengung und Expertise, in stärkerem Maße Berücksichtigung.

Zusammenfassend lässt sich festhalten, dass es aus theoretischer Perspektive zwar nahe liegt, von einer Parallelität des Vertrauens in Menschen und des Vertrauens in Automation auszugehen, die berichteten empirischen Befunde lassen an einer direkten Übertragbarkeit jedoch begründete Zweifel aufkommen (Madhavan & Wiegmann, 2007a): So scheint das Schema einer perfekt arbeitenden Automation in einem sehr hohen Anfangsvertrauen zu resultieren. Im Zuge einer leistungsbezogenen Beurteilung der Automation werden eventuelle Fehler derselben besonders stark gewichtet und möglicherweise sogar mit einem kompletten Vertrauensentzug quittiert. Demgegenüber scheint die interpersonale Vertrauensbildung gekennzeichnet durch das Zugeständnis von Unvollkommenheit, eine nicht rein leistungsbezogene Beurteilung bei der auch andere Aspekte Berücksichtigung finden und eine nachsichtige Bewertung von Fehlern. Möglicherweise ist eben doch der soziale Charakter der Interaktionssituation für eine höhere Komplexität derselben verantwortlich. Die Kehrseite hiervon ist, dass Vertrauen in Automation zwar vielleicht einfacher konstituiert ist, dafür aber möglicherweise in besonders starkem Maße das Verhalten gegenüber der Automation beeinflusst.

Dies legt die Frage nahe, wie das Vertrauen in Automation ausgeprägt sein sollte, um das Verhalten in erwünschter Weise zu beeinflussen. Die zentrale Größe für die Beantwortung dieser Frage liegt in den Leistungsmerkmalen der Automation. Idealerweise sollte das Vertrauen in eine Automation den sie kennzeichnenden Merkmalen (z.B. Zuverlässigkeit, Verständlichkeit, Nützlichkeit) entsprechen und diesen angemessen sein. Sowohl ein geringes Vertrauen in eine hoch zuverlässige Automation als auch ein hohes Vertrauen in eine vergleichsweise unzuverlässige Automation gilt es zu vermeiden (Gempeler & Wickens, 1998). Lee und See (2004) differenzieren dieses Basiskonzept etwas stärker aus und unterscheiden drei Parameter, die für ein angemessenes Vertrauen in Automation von Bedeutung sind: Kalibrierung, Spezifität und Auflösung. Letzteres bezieht sich darauf, wie präzise ein Vertrauensurteil unterschiedliche Zuverlässigkeitsgrade der Automation zu differenzieren vermag. Der zweite Parameter, Spezifität, spiegelt das Ausmaß wider, indem verschiedenen Systemkomponenten einer Automation unterschiedliches Vertrauen im Hinblick auf ihr Leistungsvermögen entgegengebracht wird. Und schließlich der dritte Parameter, Kalibrierung, bezieht sich auf die Entsprechung von subjektivem Vertrauen und objektivem Leistungsvermögen der Automation – eine schlechte Kalibrierung äußert sich dementsprechend entweder in einem übersteigerten (*overtrust*) oder einem mangelnden Vertrauen (*distrust*). Als mögliche Folgen werden *misuse* bzw. *disuse* angenommen (vgl. Kap. 1.2). Abbildung 1 veranschaulicht eine angemessene Vertrauenskalibrierung sowie eine „gute“ und eine „schlechte“ Auflösung.

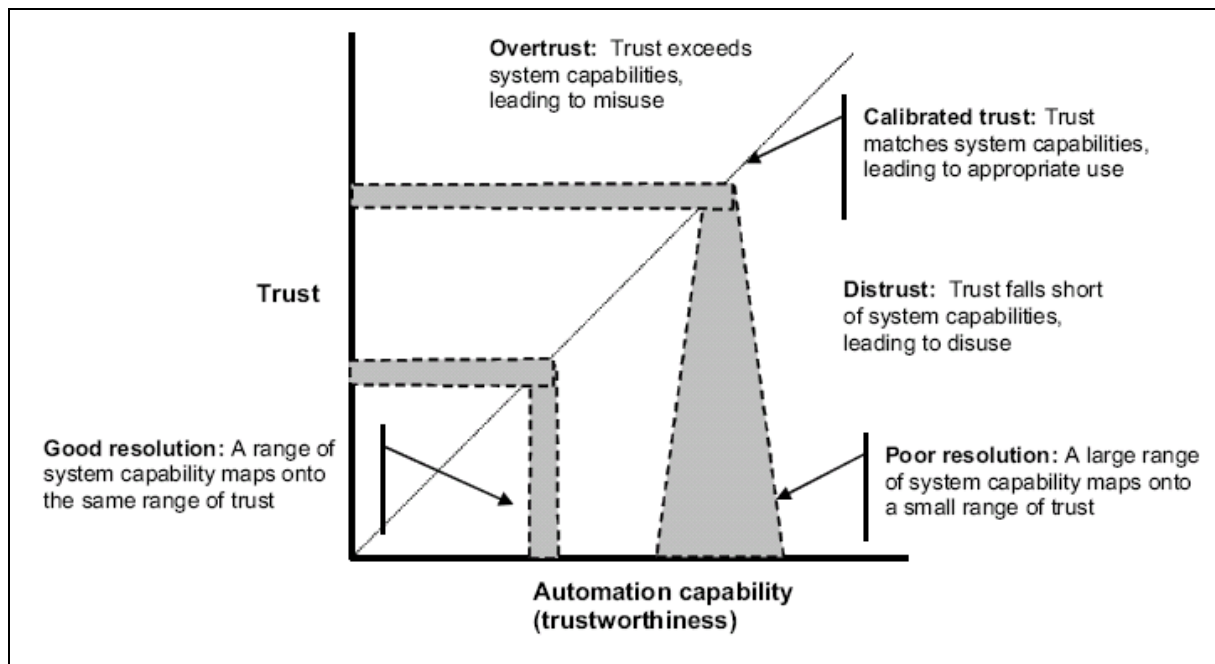


Abb. 1: Die Beziehung zwischen Vertrauen und Leistungsvermögen der Automation (aus Lee & See, 2004, S. 55).

Zusammenfassend lässt sich festhalten, dass das Konzept Vertrauen in Automation zwar oft Gegenstand theoretischer Klärungsversuche war, bislang aber kein geteiltes Begriffsverständnis auszumachen ist, und auch die Frage der zugrunde liegenden Vertrauensdimensionen nicht abschließend geklärt ist. Für die vorliegende Arbeit soll hilfsweise die relativ allgemeine Definition von Lee und See (2004) zugrunde gelegt und Vertrauen als die Überzeugung einer Person, dass ihr Interaktionspartner (Mensch oder Automation) sie bei der Erreichung ihrer Ziele in einer durch Unsicherheit gekennzeichneten Situation unterstützen wird, begriffen werden. Mit Blick auf die der Vertrauensentstehung zugrunde liegenden Dimensionen lässt der bisherige Forschungsstand keine Festlegung zu. Wichtig erscheint aber die Beobachtung, dass gegenüber dem Vertrauen in sozialen Interaktionssituationen, das Vertrauen in Automation sich einerseits durch eine geringere Komplexität auszeichnet, andererseits aber eine höhere Verhaltensrelevanz zu besitzen scheint. Je nach dem wie stark das Vertrauen und die Vertrauenswürdigkeit der Automation ausgeprägt sind, lässt sich das subjektive Vertrauen einer Person in ein automatisiertes System als mangelnd, angemessen oder übersteigert beschreiben.

Während in diesem Kapitel der Schwerpunkt auf einer theoretischen Auseinandersetzung mit dem Konzept Vertrauen lag, soll in den beiden nachfolgenden Kapiteln durch die Aufbereitung der empirischen Befundlage das bisher skizzierte Bild von Vertrauen in Automation geschärft werden. Die Strukturierung der empirischen Befunde orientiert sich dabei an den beiden unangemessenen Vertrauensausprägungen eines übersteigerten und mangelnden Vertrauens. Ziel

der Ausführungen ist es, aufzuzeigen, welche Aspekte bei der Entstehung eines übersteigerten bzw. mangelnden Vertrauens eine Rolle spielen.

1.3.2 *Mangelndes Vertrauen in Automation*

Probleme eines mangelnden Vertrauens in Automation und einer daraus folgenden mangelnden Nutzung (*automation disuse*, Parasuraman & Riley, 1997) automatisierter Systeme wurden insbesondere für Situationen untersucht, in denen dem Menschen die Entscheidung über Nutzung oder Nichtnutzung eines automatisierten Systems obliegt. Beispiele für solche Systeme sind insbesondere adaptierbare automatisierte Systeme, bei denen der Mensch über die Funktionsallokation flexibel entscheiden kann (z.B. Autopilot, Navigationssysteme, Systeme zur automatischen Mustererkennung), oder auch automatisierte Warn- und Alarmsysteme, die entsprechend beachtet oder ignoriert werden können.

Betrachtet man zunächst die Nutzung adaptierbarer automatisierter Systeme, so zeigt sich in der Regel ein enger korrelativer Zusammenhang zwischen der Zuverlässigkeit dieser Systeme auf der einen und dem Vertrauen, das diesen Systemen entgegengebracht wird bzw. ihrer Nutzung auf der anderen Seite (Muir & Moray, 1996). Allerdings wurden in der Interaktion mit objektiv zuverlässig (wenn auch nicht perfekt) arbeitenden Systemen immer wieder Situationen beobachtet, in denen Personen sich gegen die Nutzung der Automation entscheiden und die infrage stehende Funktion manuell ausüben (Beck, Dzindolet & Pierce, 2007; Dzindolet et al., 2002; Hockey & Maule, 1995; Riley, 1996).

Für diesen Effekt können zwei verschiedene Ursachen geltend gemacht werden (vgl. Manzey & Bahner, 2005): Erstens kann die Wahrnehmung von einzelnen Automationsfehlern zu einer Unterschätzung des objektiven Leistungsvermögens der Automation und damit zu einem Vertrauenseinbruch führen. Mehrere Untersuchungen zeigen, dass das Vertrauen in eine Automation sehr anfällig ist und selbst durch vergleichsweise unbedeutende Fehler der Automation zusammenbrechen kann. So stellten Muir und Moray (1996) fest, dass das Vertrauen in die Automation (hier automatisierte Steuerung von Pumpen in einer Pasteurisierungsfabrik) bereits durch kleine Fehler, die sich nicht auf die allgemeine Leistung des Systems auswirkten, signifikant gemindert wurde und dies dann in einer Ablehnung der Automation und manueller Steuerung resultierte. Ein solcher Effekt wurde auch in anderen Studien beobachtet (Lee & Moray, 1992, 1994; Dzindolet et al., 2003; Dzindolet, Pierce, Beck & Dawe, 1999; Wiegmann, Rich & Zhang, 2001, De Vries, Midden & Bouwhuis, 2003). Offenbar können die durch die Automationsfehler hervorgerufenen Vertrauenseinbußen größer sein als es dem objektiven Leistungsvermögen der Automation entspricht und damit ein anfangs möglicherweise übersteigertes Vertrauen (*positive bias*, vgl.

Kap. 1.3.1) in ein mangelndes Vertrauen umkehren. Zudem zeigen bisherige Befunde, dass sich mangelndes Vertrauen in eine Systemkomponente eines automatisierten Systems auch auf andere Systemkomponenten generalisiert (Muir & Moray, 1996) und damit eine hohe funktionale Spezifität des Vertrauens, wie von Lee und See (2004) als ein Aspekt eines angemessenen Vertrauens gefordert, nicht zu erwarten ist.

Verschiedene Untersuchungen haben sich der Frage gewidmet, unter welcher Voraussetzung es zu so einem extremen Vertrauensabfall kommt. In einer Studie von Masalonis et al. (1998), in der das Vertrauen von Fluglotsen in ein Assistenzsystem zur automatischen Erkennung von Konflikten zwischen Luftfahrzeugen untersucht wurde, zeigte sich, dass Fehler des Systems, die mit hohen potenziellen Kosten (Übersehen tatsächlich vorliegender Konflikte) verbunden waren, das Vertrauen in das System signifikant stärker minderten als Fehler mit geringen Kosten (falsche Alarmer). Offenbar sind Vertrauensverluste insbesondere dann stark ausgeprägt, wenn die mit einem Automationsfehler einhergehenden Kosten als groß wahrgenommen werden. Ähnlich wie Fehler mit vergleichsweise harmlosen Folgen wirken sich auch Fehler, die vom Operateur kompensiert werden können, weniger stark vertrauensmindernd aus (Muir & Moray, 1996; Lee & Moray, 1992).

Neben den wahrgenommenen Kosten und der Verfügbarkeit von Kompensationsmöglichkeiten scheint auch die Nachvollziehbarkeit bzw. Vorhersehbarkeit der Automationsfehler eine Rolle zu spielen. Je besser die Funktionen automatisierter Systeme nachvollzogen werden können und damit potenzielle Fehler der Automation erklärbar und vorhersehbar werden, desto geringer wirken sich jene Fehler offenbar auf das Vertrauen und die Nutzung der Systeme aus (Dzindolet et al., 2003). In eine ähnliche Richtung weisen auch die Befunde von Madhavan, Wiegmann und Lacson (2006). Sie zeigten, dass Automationsfehler vor allem dann zu Vertrauenseinbußen führen, wenn sie bei vergleichsweise einfachen Aufgaben auftreten und damit für den Operateur nicht nachvollziehbar ist, warum diese nicht vermieden werden konnten („*easy error hypothesis*“).

Eine zweite mögliche Ursache für *automation disuse* kann im Selbstvertrauen in die manuelle Ausführung der automatisierten Funktion gesehen werden. Nicht die absolute Höhe des Vertrauens, sondern die Differenz zwischen dem Vertrauen in Automation und dem Vertrauen in die eigene Leistungsfähigkeit ist dabei die entscheidende Größe, die das Ausmaß, in dem die Automation genutzt wird, zu determinieren scheint. Wenn das Vertrauen in die eigene Leistungsfähigkeit das Vertrauen in die Automation übersteigt, entscheiden sich Operateure offenbar für die manuelle Kontrolle (De Vries et al., 2003; Lee & Moray, 1992; 1994; Lewandowsky et al., 2000; Moray et al., 2000; Riley, 1996; Waern & Ramberg, 1996). Das kann dazu führen, dass auch eine weitgehend zuverlässig arbeitende Automation nicht oder nur unzureichend genutzt wird. Ver-

trauen in Automation stellt somit keine absolute Größe, sondern immer eine relative Größe dar. Erst durch den Vergleich mit der wahrgenommenen eigenen Leistungsfähigkeit bestimmt sich die Verhaltensauswirkung des Vertrauens in Automation.

Mangelndes Vertrauen spielt jedoch nicht nur im Kontext der Nutzung adaptierbarer Systeme eine Rolle. Auch im Umgang mit automatisierten Warn- und Alarmsystemen (z.B. *Ground Proximity Warning Systems* im Flugzeug, Kollisionswarnsysteme im Fahrzeug) werden oft Probleme beschrieben, die sich aus einem unzureichenden Vertrauen ergeben (vgl. auch Manzey & Bahner, 2005). Alarmsysteme werden meist dort eingesetzt, wo das Übersehen kritischer Ereignisse schwerwiegende Folgen nach sich ziehen kann. Dementsprechend wird in der Regel eine hohe Sensitivität angestrebt, damit alle potenziell kritischen Zustände zuverlässig entdeckt und der Operateur bzw. Nutzer rechtzeitig auf diese aufmerksam gemacht werden kann. Die Kehrseite der zuverlässigen Entdeckung möglichst aller kritischen Ereignisse liegt darin, dass dadurch die Rate falscher Alarme ansteigt. Die wiederholte Erfahrung falscher Alarme kann jedoch schnell dazu führen, dass dem System nicht mehr vertraut wird und Alarme mit vergleichsweise geringer Priorität behandelt oder gar ignoriert werden. Breznitz wies bereits 1983 auf diese Problematik hin und prägte hierfür den Begriff des *cry wolf* Syndroms.

Zusammenfassend lässt sich festhalten, dass verschiedene Systemeigenschaften identifiziert werden können, mit denen sich Probleme, die durch ein mangelndes Vertrauen in Automation entstehen, reduzieren lassen. Dazu gehören in erster Linie eine möglichst hohe Reliabilität und Transparenz der Funktionalität der Automation. Letztere ist insbesondere dann wichtig, wenn sich bestimmte Automationsfehler – wie z.B. falsche Alarme bei Warnsystemen – nicht vermeiden lassen. Wie gezeigt wurde, lassen sich mit Automationsfehlern verbundene Einbußen des Vertrauens in das System dann reduzieren, wenn die Fehler für den Operateur vorhersehbar und ihre Entstehung nachvollziehbar sowie ihre Folgen mit geringen Kosten verbunden bzw. gut kompensierbar sind. Darüber hinaus scheint es wichtig, die Funktionalität der Automatik möglichst dicht an den mentalen Modellen und vertrauten Bearbeitungsstrategien der Bediener auszurichten, um deren Vertrauen und Akzeptanz den Systemen gegenüber zu erhöhen und eine optimale Nutzung der Automationsmöglichkeiten zu gewährleisten.

Jedoch liegt ein wesentlicher Kosteneffekt einer derartigen Optimierung der Systemgestaltung darin, dass sich damit gleichzeitig das Risiko eines übersteigerten Vertrauens in die Automation erhöht, wie im Folgenden näher ausgeführt werden soll.

1.3.3 Übersteigertes Vertrauen in Automation

Davon ausgehend, dass Vertrauen eine kontinuierliche Variable mit den beiden Endpolen „kein Vertrauen“ (ggf. mangelndes Vertrauen) und „vollkommenes Vertrauen“ (ggf. übersteigertes Vertrauen) darstellt, liegt es nahe, dass gerade jene Automationsaspekte, die einem zu geringen Vertrauen entgegen wirken, zugleich die Basis für ein übersteigertes Vertrauen bilden. An erster Stelle zu nennen ist hier die Zuverlässigkeit der Automation – während Automationsfehler das Vertrauen mindern, führt die wiederholte Erfahrung, dass ein System fehlerfrei und wie gewünscht arbeitet, zu einer Steigerung des Vertrauens. Doch selbst wenn Fehler auftreten, muss dies, wie im vorangegangenen Kapitel (1.3.2) beschrieben, nicht zwingend in einem Vertrauensverlust resultieren. So etwa, wenn die Fehler vorhersehbar, nachvollziehbar, kompensierbar und/oder mit vergleichsweise geringen Kosten verbunden sind. Diese Eigenschaften stellen Grundvoraussetzungen für ein angemessenes Vertrauen dar, können aber damit zugleich auch das Auftreten eines übersteigerten Vertrauens fördern: Je leistungsfähiger, zuverlässiger und „besser“ die Automation gestaltet ist, desto eher kann es auch zu unerwünschten Folgen, wie hier einem übersteigerten Vertrauen in die Automation und damit verbundenen riskanten Verhaltensweisen, kommen. Gegenüber einem mangelnden Vertrauen, das sich insbesondere in einer mangelnden Nutzung automatisierter Systeme und darüber in einer Erhöhung von Fehlerrisiken und Beanspruchung niederschlägt, wirkt sich ein übersteigertes Vertrauen in die Automation oft sehr viel unmittelbarer aus. Dies mag das folgende Beispiel verdeutlichen.

Exkurs: die Havarie der „Royal Majesty“

Am 10. Juni 1995 lief die „Royal Majesty“, ein Kreuzfahrtschiff mit über 1.500 Personen an Bord, in der Nähe von Nantucket Island (Massachusetts) auf Grund. Das Schiff war auf dem Weg von den Bermudas nach Boston – bis zuletzt war die Crew der festen Überzeugung auf dem richtigen Kurs zu sein. Erst nach der völlig überraschenden Havarie registrierte die Crew, dass sie 17 Meilen vom vermeintlichen Kurs abgekommen war.

Was war passiert? Bereits kurz nach dem Auslaufen der „Royal Majesty“ löste sich das Wind und Wetter ausgesetzte Antennenkabel des *Global Positioning System* (GPS). Für etwa eine Sekunde ertönte daraufhin ein Alarmton; auf der Brücke nahm dies jedoch niemand wahr. Nachdem keine Satellitendaten mehr empfangen werden konnten, wechselte das GPS in den *dead reckoning* Modus, der darin besteht, die Position des Schiffs mittels des Kurses und der Geschwindigkeit zu berechnen. Die Abdrift durch Strömungen, Wind oder auch minimale Steuerfehler bleiben dabei jedoch unberücksichtigt. Davon ausgehend, dass das GPS wie auch das *Integrated Bridge System* problemlos funktionierten, wurde kurz nachdem das Schiff ausgelaufen war bedenkenlos,

und wie es der gängigen Praxis auf dieser Strecke entsprach, die Navigation des Schiffes dem Auto-Piloten übergeben. Da dieser aber die Positionsdaten des Schiffs vom GPS erhielt, wurde das Schiff auf der Grundlage fehlerhafter Positionsdaten manövriert. Doch offenbar fiel für mehr als 34 Stunden niemandem die stetig zunehmende Kursabweichung auf: Stündlich waren die Positionsdaten des Schiffs zu protokollieren. Jedes Mal wurden genau jene Positionsdaten eingetragen, die das GPS fehlerhaft anzeigte. Eine Verifikation dieser Daten anhand des korrekt funktionierenden LORAN-C (berechnet auf Basis von Radiosignalen von küstennahen Sendern die Position des eigenen Schiffes) erfolgte nicht. Mit zunehmender Abweichung vom geplanten Kurs mehrten sich die Anzeichen, dass die angenommene Position nicht der tatsächlichen entsprechen konnte: Aufgrund der unmittelbaren Nähe zu gefährlichen Untiefen ist die Anfahrt des Hafens von Boston über die *Boston Traffic Lanes* gesichert und durch sechs Tonnen markiert. Der diensthabende Offizier attestierte, dass er die erste, den Eingang zu den *Boston Traffic Lanes* markierende Tonne ordnungsgemäß und zum erwarteten Zeitpunkt passiert habe. Allerdings kommt er zu diesem Schluss ausschließlich aufgrund der Radarkarte, die neben den Tonnen auch den geplanten Kurs und die eigene Position anzeigt; letztere ebenfalls auf Basis der längst fehlerhaften GPS-Daten. Wie sich später herausstellt, passierte die „Royal Majesty“ nicht wie gedacht die Eingangstonne zu den *Boston Traffic Lanes*, sondern eine 15 Meilen weiter westlich gelegene Tonne, die die Nantucket-Untiefen markiert. Eine direkte visuelle Überprüfung der Tonne anhand ihrer spezifischen Leuchtcharakteristik erfolgte nicht. Etwas später berichteten mehrere Wachen dem diensthabenden Offizier auf der Brücke von blinkenden gelben und roten Lichtern auf Backbord – diese sind normalerweise von den *Boston Traffic Lanes* aus nicht zu sehen und markierten offenbar die Küstenlinie von Nantucket Island. Der Offizier nahm die Information zur Kenntnis, reagierte jedoch nicht. Derselbe Offizier berichtete wenige Minuten später dem Kapitän auf Nachfrage, auch die zweite Tonne der *Boston Traffic Lanes* passiert zu haben. Im Nachhinein gibt er an, dass er die Tonne zwar weder direkt noch auf dem Radar gesehen habe, es aber möglich sei, dass das Radar eine Tonne nicht reflektiere. Daher habe er die GPS-Daten als Beleg für den richtigen Kurs gewertet. Kurz nachdem der Kapitän beruhigt die Brücke wieder verlassen hat, läuft die „Royal Majesty“ für die gesamte Besatzung völlig überraschend auf Grund.

Die später erfolgte Unfallanalyse (National Transportation Safety Board, 1997; vgl. auch Degani, 2003) macht folgende Ursachen für die Havarie verantwortlich: *Overreliance* aller diensthabenden Offiziere, das Versäumnis der Reederei, die Crew angemessen bezüglich der automatisierten Systeme zu trainieren, Mängel im Design und der Implementierung des *Integrated Bridge System* als auch bezüglich der Betriebsabläufe, und schließlich das Versäumnis des zuletzt diensthabenden Offiziers, korrektive Maßnahmen einzuleiten, obwohl mehrere Hinweise nur als Abkommen vom Kurs interpretiert werden konnten.

Dieses Beispiel schildert auf eindrückliche Weise ein oft beschriebenes Phänomen: Ein übersteigertes Vertrauen in die Automation führt dazu, dass letztendlich niemand mehr die Funktions- und Leistungsfähigkeit derselben hinterfragt; eine Überwachung der Automation im Sinne von Kontrollen und Verifikationen anhand anderer verfügbaren Quellen (z.B. LORAN-C, Sichtprüfung der Tonnen) unterbleibt. Mehr noch: Informationen, die Zweifel an den automatisch generierten Positionsdaten aufkommen lassen müssten (Lichter werden wahrgenommen, die nicht da sein dürften, wichtige Markierungstonnen werden weder vom Radar angezeigt noch direkt gesehen), werden zwar registriert, aber offenbar nicht als Falsifikation der automatisch generierten Informationen gewertet. Trotz widersprechender Informationen wird der Automation blind gefolgt.

Glücklicherweise wurde bei der Havarie der „Royal Majesty“ niemand verletzt. Andere Unfälle mit vergleichbaren Unfallursachen haben jedoch weitaus tragischere Folgen nach sich gezogen, wie etwa der Absturz der Boeing 757 nahe Cali, Kolumbien, 1995 (Aeronautica Civil of the Republic of Columbia, 1996), bei dem die gesamte Besatzung und 151 Passagiere ums Leben kamen.

Die beobachtbaren (und zu vermeidenden) Verhaltensfolgen eines übersteigerten Vertrauens lassen sich unter dem Begriff *automation misuse* zusammenfassen (Parasuraman & Riley, 1997; Beck, Dzindolet, & Pierce, 2002). Untersucht wurden derartige Probleme bislang sowohl im Hinblick auf die überwachende Kontrolle automatisierter Systeme als auch mit Blick auf die Nutzung automatisierter Assistenzsysteme. Während *complacency* (Parasuraman, Molloy & Singh, 1993) ein Phänomen aus dem Kontext der überwachenden Kontrolle darstellt, bezeichnet *automation bias* (Mosier & Skitka, 1996) ein verwandtes Problem, das bei der Nutzung von Assistenzsystemen auftreten kann. In den folgenden zwei Kapiteln soll auf beide Aspekte getrennt voneinander eingegangen werden. Dabei sollen jeweils der Versuch einer Begriffsklärung vorgenommen und bisherige empirische Befunde mit Blick auf begünstigende bzw. entgegenwirkende Faktoren dargestellt werden.

1.4 COMPLACENCY

1.4.1 *Complacency – Begriffsklärung*

Geprägt wurde der Begriff *complacency* in Zusammenhang mit psychologischen Problemen der zunehmenden Automatisierung im Luftfahrtbereich. Beschrieben werden damit allgemein Defizite bei der Überwachung automatisierter Systeme im Flugzeugcockpit und daraus resultierende Risiken für die Flugsicherheit (Wiener, 1985). Funk und Lyall (2000) sehen darin auch 15

Jahre später noch eines der drei wichtigsten Problemfelder im Hinblick auf mögliche negative Auswirkungen der Cockpit-Automation. Das oben geschilderte Beispiel der „Royal Majesty“ zeigt jedoch, dass das Phänomen nicht auf den Luftfahrtkontext beschränkt zu bleiben scheint: Auch bei diesem Unglück spielte *complacency* im Sinne einer unzureichenden Überwachung des GPS eine zentrale Rolle. In allen Bereichen, die durch eine zunehmende Automatisierung gekennzeichnet sind, wie zum Beispiel Prozesskontrolle, Kraftfahrzeugführung oder Bahn- und Schiffsverkehr, wird eine Auseinandersetzung mit *complacency*-Effekten und vergleichbaren Phänomenen nicht ausbleiben können.

In deutlichem Kontrast zu der Bedeutung, die dem *complacency*-Konzept in der Praxis auf der Basis anekdotischer Berichte oder im Rahmen von Unfallanalysen zugeschrieben wird, steht allerdings, dass sich bislang kein geteiltes Begriffsverständnis feststellen lässt. Ähnlich wie die Konzepte „Situationsbewusstsein“ oder „mentale Beanspruchung“ hat sich in den letzten Jahren „*complacency*“ als ein Schlagwort im *human factors*-Bereich eingebürgert. Intuitiv scheint jeder etwas damit zu verbinden – was jedoch, scheint höchst unterschiedlich zu sein (Dekker & Hollnagel, 2004).

Ein Teil der Definitionsansätze betont den Zusammenhang von *complacency* mit einem bestimmten kognitiven Zustand wie etwa Wiener (1981), der das Konzept als „psychological state characterized by a low index of suspicion“ (S. 117) fasst. Billings, Lauber, Funkhouser, Lyman und Huff (1976) stellen hingegen den Bezug zu Aufmerksamkeitsprozessen und einem übersteigerten Vertrauen in das System her und beschreiben *complacency* in einer sehr häufig zitierten Definition als „self-satisfaction which may result in non-vigilance based on an unjustified assumption of satisfactory system state“ (S. 23). Auch Moray (2003; Moray & Inagaki, 2000) betont den Bezug von *complacency* zu Aufmerksamkeitsprozessen. Danach spiegelt sich *complacency* in einem unangemessenen Überwachungsverhalten in dem Sinne wider, dass eine Automation seltener überwacht wird als es aufgrund ihrer Leistungsfähigkeit (insbesondere Zuverlässigkeit) notwendig wäre.

Deutlich wird, dass sich in der Literatur kein übereinstimmender semantischer Gehalt des Konzepts ausmachen lässt. Dekker und Hollnagel (2004) gehen sogar soweit, *complacency* als wissenschaftlich unbrauchbares Alltagskonzept (*folk model*) einzustufen: Erklärung durch Substitution, Immunisierung gegen Falsifikation und Übergeneralisierung seien die wenig rühmlichen Eigenschaften des Konzepts. Um dieser in Teilen sicher berechtigten Kritik zu begegnen und um eine wissenschaftliche Auseinandersetzung mit *complacency* potenziell fruchtbar zu machen, ist eine Konkretisierung des Begriffs zwingend notwendig. Hierzu soll das Konstrukt zunächst von ver-

wandten Konzepten wie OOTLUF, Situationsbewusstsein, *overreliance* und *overtrust* abgegrenzt werden.

Complacency wird, wie bereits einleitend erwähnt (vgl. Kap. 1.2), im Kontext von OOTLUF genannt. Einige Ansätze ziehen *complacency* als Ursache von OOTLUF heran (z.B. O'Hare & Roscoe, 1990). In den meisten Konzeptionalisierungsversuchen wird *complacency* aber als eine mögliche Ausprägung von OOTLUF angeführt (Kaber & Endsley, 1997; Wickens & Hollands, 2000). Vor diesem Hintergrund stellt sich die Frage, wie das Konzept von einem Verlust von Situationsbewusstsein abzugrenzen ist. Kaber und Endsley (1997) nennen die Abnahme von Situationsbewusstsein und *complacency* als gleichberechtigt nebeneinander stehende Konzepte und vermitteln so implizit das Bild zweier unabhängiger Phänomene. Dagegen sieht Kern (1998) einen kausalen Zusammenhang, wobei *complacency* einen Verlust an Situationsbewusstsein verursacht. Das Konzept als einzige Determinante eines Verlusts von Situationsbewusstsein zu begreifen, scheint jedoch wenig schlüssig: Situationsbewusstsein ist als Konzept breiter angelegt und wird für deutlich mehr Situationen herangezogen als das für das Konzept *complacency* der Fall ist. Dies ergibt sich auch aus den Merkmalen, die einen Verlust von Situationsbewusstsein kennzeichnen: Defizite mit Blick auf die Wahrnehmung, das Verstehen und die Vorhersage des Systemzustands (Endsley, 1995) können auch durch andere Ursachen wie z. B. Vigilanzdefizite oder eine ungünstige Gestaltung der Mensch-Maschine-Schnittstelle bedingt sein. Es scheint demnach plausibler, *complacency* als eine von mehreren möglichen Determinanten des Verlusts von Situationsbewusstsein zu fassen.

Complacency wird oftmals in die Nähe der Konzepte *reliance* und *overconfidence* bzw. Vertrauen gerückt (Stokes & Kite, 1994; Kern, 1998), sodass auch hier eine begriffliche Abgrenzung versucht werden soll. Nach Riley (1996) repräsentiert *reliance* „the probability that an operator will use automation and is influenced by the operator's self-confidence and level of trust in the automation“ (S. 21). In ähnlicher Weise fassen Parasuraman und Riley (1997) unter *overreliance* die Nutzung einer Automatisierung, obwohl sie nicht genutzt werden sollte bzw. die fehlerhafte Überwachung der Automatisierung. Während *overreliance* demnach Situationen umfasst, in denen dem Nutzer die a priori Entscheidung über Nutzung oder Nichtnutzung der Automatisierung obliegt, wird *complacency* meist ausschließlich auf die Überwachung automatisierter Systeme bezogen. Somit stellt *complacency* eine Sonderform von *overreliance* dar. Oftmals gleichgesetzt wird *complacency* jedoch auch mit einem übersteigerten Vertrauen in Automation. Wie schon in Kapitel 1.3.1 beschrieben, scheint es jedoch angemessener, Vertrauen in Automation als Einstellung und *reliance* als das korrespondierende Verhalten zu begreifen (Lee & See, 2004) und davon auszugehen, dass Vertrauen nur mittelbar über Intentionen Einfluss auf das Verhalten nimmt, wobei

verschiedene äußere Faktoren (z. B. Zeitdruck) die Intentionsbildung und deren Umsetzung in konkretes Verhalten zusätzlich beeinflussen können. Ein hohes Vertrauen muss sich folglich nicht zwingend im Verhalten widerspiegeln.

Auf Grundlage der bisherigen Ausführungen und in Übereinstimmung mit Manzey und Bahner (2005) soll in der vorliegenden Arbeit *complacency* als ein Merkmal der Mensch-Maschine-Interaktion verstanden werden, das sich auf der Verhaltensebene in einer unzureichenden Überwachung automatisierter Systeme durch den Menschen widerspiegelt, mit einem übersteigerten Vertrauen in die Funktions- und Leistungsfähigkeit der Automation zusammenhängt und als Folge zu einem Übersehen kritischer Systemzustände führen kann. Indem damit die Folgen eines unzureichenden Überwachungsverhaltens in die Definition eingeschlossen werden, wird der in den meisten Untersuchungen gewählten Operationalisierung des Konzeptes Rechnung getragen (z.B. Parasuraman et al., 1993). Dabei muss allerdings berücksichtigt werden, dass das Übersehen von Automationsfehlern allenfalls ein mögliches, sicherlich aber nicht hinreichendes Anzeichen für das Vorliegen von *complacency* darstellt, wie in Kapitel 1.4.4 noch näher ausgeführt wird. Zunächst soll jedoch die empirische Befundlage mit Blick auf mögliche Entstehungsbedingungen von *complacency* berichtet werden.

1.4.2 Entstehungsbedingungen von *complacency*

Bisherige empirische Studien legen es nahe, drei verschiedene Faktoren zu unterscheiden, die zu der Entstehung von *complacency* im Umgang mit Automation beitragen: Merkmale der Automation, der Person und des situativen Kontextes (vgl. auch Bahner & Manzey, 2004; Manzey & Bahner, 2005).

1.4.2.1 Einflüsse von Merkmalen der Automation

Bereits die erste empirische Untersuchung des Konzepts *complacency* (Parasuraman et al., 1993), die auch heute noch sehr häufig zitiert wird, adressierte Merkmale der Automation. Als Versuchsumgebung diente hier, wie auch in den meisten nachfolgenden Studien, die *multi-attribute task battery* (MAT-Batterie, Comstock & Arnegard, 1992), eine Flugsimulations-Mikrowelt mit verschiedenen Subsystemen. Abbildung 2 zeigt die Benutzeroberfläche der MAT-Batterie. Die Probanden von Parasuraman et al. (1993) bearbeiteten parallel drei Aufgaben. Eine Tracking-Aufgabe (oben Mitte), eine Ressourcen-Management-Aufgabe (unten) und eine Systemüberwachungsaufgabe (oben links). Bei letzterer variierten im Normalfall die vier Parameterwerte um die Mittelpunkte der vier Vertikalanzeigen. Bisweilen traten jedoch Abweichungen dahingehend auf, dass der Mittelwert eines Parameters nach oben oder unten verschoben war und somit nicht

mehr mit der Mitte der Vertikalanzeige übereinstimmte. Derartige Abweichungen wurden im Regelfall automatisch korrigiert, wobei ein Eingreifen der Automation durch Aufblinken des „Warning“-Feldes angezeigt wurde.

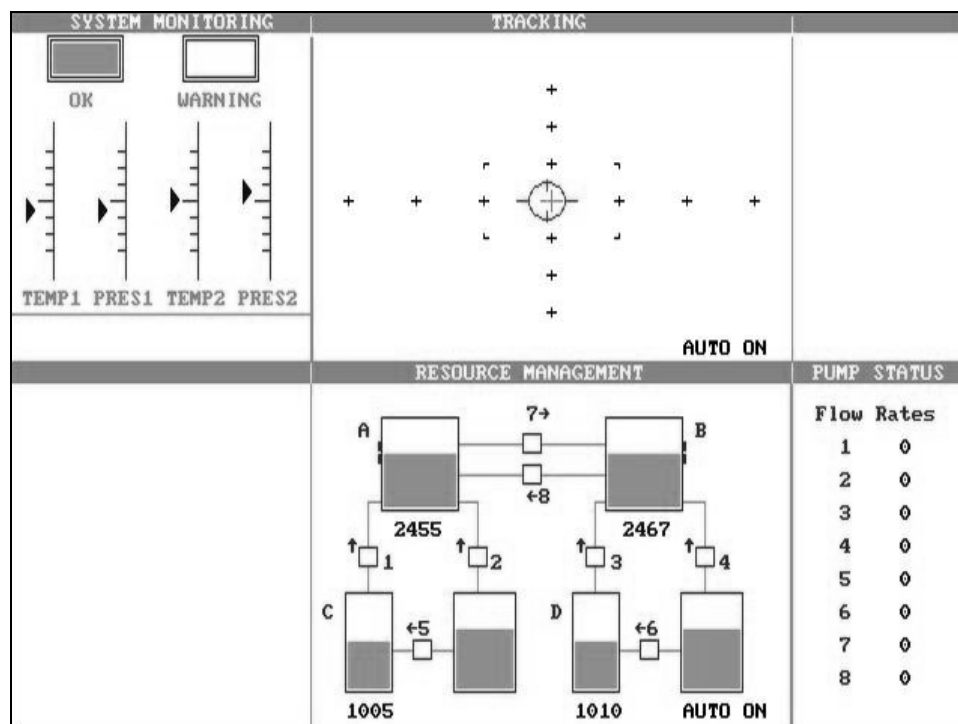


Abb. 2: Benutzeroberfläche der *multi-attribute task*-Batterie (zur Verfügung gestellt von der Technischen Universität Berlin, Fachgebiet Arbeits-, Ingenieur- und Organisationspsychologie)

Fehler der Automation äußerten sich darin, dass eine Abweichung der zu regelnden Parameter nicht automatisch korrigiert wurde, sondern stattdessen von den Probanden per Tastendruck behoben werden sollte. Der Prozentuale Anteil von Automationsfehlern, die von den Probanden innerhalb eines bestimmten Zeitfensters korrigiert wurden, diente der Berechnung der Fehlerdetektionsrate, die wiederum als *complacency*-Indikator (umgekehrt proportional) interpretiert wurde.

Experimentell variiert wurde zum einen die Höhe der Reliabilität der Automation (hoch 87.5 % versus niedrig 56.25 %; ergibt sich jeweils aus 100 abzüglich des prozentualen Anteils von Automationsfehlern an den auftretenden Paramaterabweichungen) und zum anderen die Konstanz der Reliabilität (konstant hoch bzw. konstant niedrig versus variabel). Unter der variablen Bedingung fluktuerte die Reliabilität alle zehn Minuten zwischen hoch und niedrig, wobei ein Teil der Probanden mit der hoch reliablen Bedingung begann (variabel: hoch-niedrig) und ein anderer Teil zunächst die Erfahrung einer niedrigen Systemreliabilität machte (variabel: niedrig-hoch). Parasuraman et al. (1993) fanden, dass der Konstanz der Reliabilität die entscheidende Rolle für die Entstehung von *complacency* zukommt. Unter der Bedingung konstanter Reliabilität entdeckten die Probanden signifikant weniger Automationsfehler als unter der variablen Reliabilitätsbedingung (siehe Abb. 3).

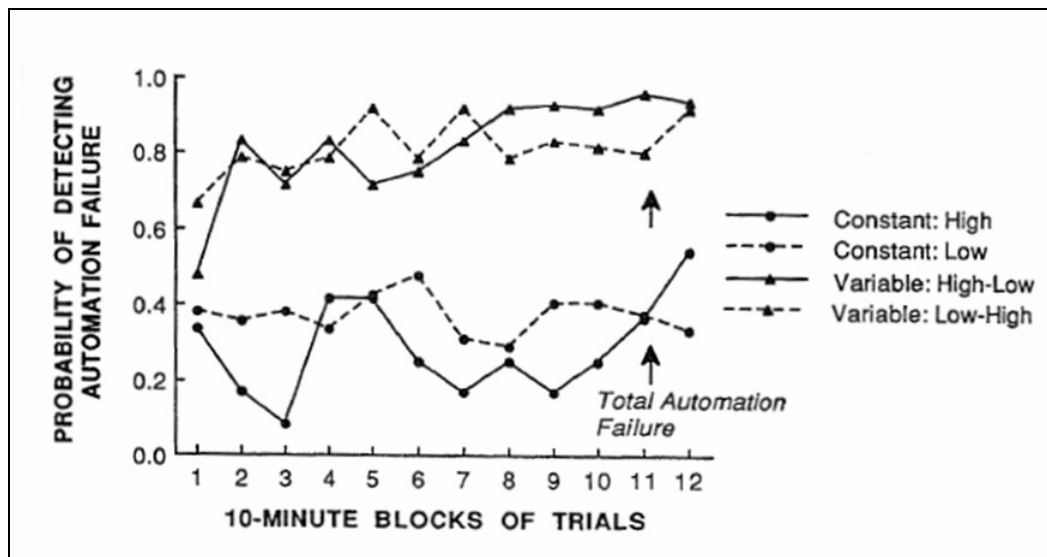


Abb. 3: Ergebnisse der Studie von Parasuraman, Molloy und Singh (1993, S. 11). Einfluss von Reliabilitätshöhe und Konstanz der Reliabilität auf die Wahrscheinlichkeit, Fehler der Automation zu entdecken.

Ein Haupteffekt für die Höhe der Reliabilität konnte überraschenderweise nicht gefunden werden, obwohl es zunächst nahe liegen würde, anzunehmen, dass zumindest in der konstant niedrig reliablen Bedingung die Entdeckungsleistung besser und damit *complacency* geringer ausgeprägt sein müsste als unter der hoch reliablen Bedingung.

Der Befund, dass *complacency* (jeweils operationalisiert über die Detektionsrate bezüglich der Automationsfehler) insbesondere bei konstant hoher Automationsreliabilität auftritt und durch variable Reliabilität reduziert werden kann, wurde unter Verwendung desselben Untersuchungsparadigmas repliziert (Singh, Molloy & Parasuraman, 1997; Prinz, DeVries, Freeman & Mikulka, 2001). Allerdings wurde in diesen Studien keine Bedingung mit konstant niedriger Reliabilität realisiert. In einer direkten Replikation der Studie von Parasuraman et al. (1993) fanden Bagheri und Jamieson (2004a) dagegen keinen Unterschied zwischen konstanter und variabler Reliabilität. Einzig die Gruppe mit konstant hoher Reliabilität zeigte gegenüber den anderen drei – sich nicht unterscheidenden – Gruppen (variabel: hoch-niedrig, variabel: niedrig-hoch und konstant: niedrig) eine signifikant schlechtere Fehlerdetektionsleistung. In einer weiteren Studie derselben Autoren (Bagheri & Jamieson, 2004b) wurde eine schlechtere Leistung für die Bedingung konstant hoher Reliabilität nur bei direkter Analyse des Überwachungsverhaltens über Blickbewegungsdaten beobachtet. Mit Blick auf die von Parasuraman et al. (1993) gewählte Operationalisierung von *complacency* über die Fehlerdetektionsrate zeigte sich in dieser Studie dagegen kein Einfluss der Reliabilitätsbedingung. Diese Befunde legen die Vermutung nahe, dass möglicherweise weniger der Konstanz der Reliabilität sondern vielmehr der Reliabilitätshöhe eine entscheidende Rolle für das Überwachungsverhalten zukommt. Bailey und Scerbo (2007) untersuchten den Einfluss von deutlich höheren – und damit ökologisch valideren – Reliabilitäten auf *complacency*: 87 % als nied-

rige (dies entspricht in etwa der hohen Reliabilität der übrigen Studien) und 97 % als hohe Reliabilität. Erwartungsgemäß entdeckten die Probanden signifikant weniger Automationsfehler in der hoch-reliablen Bedingung. Möglicherweise wurden von den Probanden in den früheren Untersuchungen die gewählten Reliabilitätsstufen von 56.25 und 87.5 % auch als insgesamt unzuverlässig oder der subjektive Reliabilitätsunterschied geringer wahrgenommen als es der objektiven Differenz von 31.25 % entspricht. In diese Richtung weist auch eine Studie von Manzey und Otte (2007) in der gezeigt wurde, dass erst bei einer Reliabilitätshöhe von über 90 % deutlich weniger überwacht wird, während Reliabilitäten von 56.25 % bis einschließlich 87.5 % zu einer vergleichbar hohen Überwachungsintensität führten. Doch auch das könnte die kontraintuitiven Befunde früherer Studien nur teilweise erklären. Auch wenn noch nicht abschließend geklärt ist, welche relative Bedeutsamkeit der Reliabilitätshöhe bzw. -konstanz zukommt, scheint eine hoch reliable Automation *complacency* zu begünstigen.

Darüber hinaus liegt es nahe, davon auszugehen, dass jene Automationsmerkmale, die einem übersteigerten Vertrauen in Automation Vorschub leisten, auch an der *complacency*-Entstehung beteiligt sind. Wie in Kapitel 1.3.3 ausgeführt, sind dies u. a. die Transparenz der Automation und die Nachvollziehbarkeit von ggf. auftretenden Automationsfehlern. Allerdings scheint die präzise Information darüber, wann Fehler in welchem Umfang auftreten werden, *complacency* auch entgegenwirken zu können. Von einer solchen Information profitierten in einer Studie von Bagheri und Jamieson (2004b) insbesondere die Probanden, die mit einer konstant hoch-reliablen Automation (die ja in der Regel *complacency* begünstigt) arbeiteten: Ihnen gelang es vor dem Hintergrund der Information, dass die Automation konstant knapp unter 100 % reliabel arbeite (dies wurde auch anhand von Graphiken veranschaulicht), ihr Überwachungsverhalten zu intensivieren. Unterschiede bezüglich der Fehlerdetektion zwischen jener Gruppe und den übrigen Gruppen mit niedrig-reliabler oder variabler Reliabilität verschwanden gänzlich.

In diesem Fall führte die verbesserte Vorhersehbarkeit und Transparenz offenbar zu einem angemessenen Vertrauen in die Automation. Dies macht deutlich, wie verkürzt es wäre, davon auszugehen, dass jede hoch und konstant reliable sowie transparent und nachvollziehbar gestaltete Automation zu *complacency* führt. Da diese Automationseigenschaften insbesondere als Grundvoraussetzungen für ein angemessenes Vertrauen zu verstehen sind, liegt es auf der Hand, dass zusätzliche Bedingungen gegeben sein müssen, damit es zu einem übersteigerten Vertrauen und in der Folge zu *complacency* kommt. Dabei könnten vor allem Wechselwirkungen mit individuellen Merkmalen der Benutzer von Bedeutung sein.

1.4.2.2 Einflüsse durch Merkmale der Person

Verschiedene Untersuchungen zeigen, dass sich Individuen in ihrer Anfälligkeit für *complacency* unterscheiden. Erste empirische Arbeiten dazu wurden von Singh, Molloy und Parasuraman (1993a, b) durchgeführt. Sie nehmen an, dass ein übersteigertes Vertrauen in Automation mit generalisierten technikbezogenen Einstellungen einer Person zusammenhängen. Zur Erfassung dieser Einstellungen entwickelten sie die *Complacency Potential Rating Scale* (CPRS, Singh et al., 1993a). Erste Untersuchungen mit dieser Skala zeigen, dass damit tatsächlich interindividuelle Unterschiede in technikbezogenen Einstellungen zuverlässig erfasst werden können. Untersuchungen zur Vorhersagevalidität des so diagnostizierten *complacency*-Potentials für das Überwachungsverhalten im Umgang mit Automation zeigen allerdings noch kein konsistentes Muster und deuten damit auf mögliche moderierende Einflüsse sowohl von Systemeigenschaften als auch situativen Rahmenbedingungen hin. So fanden Prinzel et al. (2001) nur dann den erwarteten Einfluss des *complacency*-Potentials auf das Überwachungsverhalten, wenn die Personen mit einem konstant zuverlässigen System interagierten. In diesem Fall wurden die seltenen Automationsfehler signifikant häufiger von Personen entdeckt, die nur eine geringe Ausprägung auf der CPRS aufwiesen. Neben rein technikbezogenen Einstellungen scheinen aber auch noch weitere allgemeine Persönlichkeitsmerkmale mit dem Auftreten von *complacency* in Zusammenhang zu stehen. Darauf weisen die Befunde, dass eine hohe Neigung zu Langeweile Überwachungsdefizite ebenso begünstigt wie eine hohe Neigung zu kognitiven Fehlern (Prinzel et al., 2001). Weiterhin scheinen Personen mit einer niedrigen Selbstwirksamkeitserwartung bezüglich der zu leistenden Aufgabe eher *complacency* zu zeigen (Prinzel, 2002). Das passt zu den bereits dargestellten Befunden (siehe Kap. 1.3.2), nach denen Personen sich vor allem dann auf die Automation verlassen und diese nutzen, wenn das Selbstvertrauen in die eigene Leistungsfähigkeit geringer ausgeprägt ist als das Vertrauen in die Automation. Auch gibt es erste Befunde, nach denen Sorglosigkeit/Optimismus und Risiko-/ Ungewissheitstoleranz die Entstehung von *complacency* begünstigen (Feuerberg, Bahner & Manzey, 2005).

1.4.2.3 Einflüsse durch Merkmale des situativen Kontextes

Auch wenn ein Individuum im Umgang mit einem bestimmten automatisierten System – aufgrund bestimmter persönlicher Merkmale und spezifischer Eigenschaften der Automation – eine Neigung zu übersteigertem Vertrauen in die Automation aufweist, muss es noch nicht unbedingt zu einer Verhaltensmanifestation von *complacency* kommen. Personenmerkmale und Merkmale der Automation stellen in diesem Sinne eine zwar notwendige, vermutlich aber nicht hinreichende Bedingung für ein nachlässiges Überwachungsverhalten dar. Eigenschaften der Situation

kann hier eine moderierende Rolle zugeschrieben werden. So wird als eine wesentliche situative Grundvoraussetzung immer wieder das Vorhandensein konkurrierender Zielstellungen, wie sie unter Mehrfacharbeitsbedingungen vorliegen, beschrieben (Parasuraman et al., 1993; Thackray & Touchstone, 1989). Ein Beispiel dafür liefert die bereits in Kapitel 1.4.2.1 zitierte Untersuchung von Parasuraman et al. (1993): Fehler der Automation wurden nur dann zuverlässig entdeckt, wenn die Überwachungsaufgabe die einzige Aufgabe war, die bearbeitet werden sollte. War die Aufgabe eingebettet in ein komplexes Szenario mit mehreren anderen, konkurrierenden Aufgaben nahm die Qualität der Überwachungsleistung – gemessen an den entdeckten Automationsfehlern – signifikant ab. Allerdings ist fraglich, ob sich darin tatsächlich auch die Folgen eines übersteigerten Vertrauens in die Automation widerspiegeln. Die Reduktion der Überwachungsleistung und das damit einhergehende Verlassen auf die Zuverlässigkeit der Automation könnte in diesem Zusammenhang auch als eine kompensatorische Strategie verstanden werden, um mit den erhöhten Anforderungen bei der Aufgabenbewältigung zurecht zu kommen und genügend Ressourcen für die Bearbeitung der konkurrierenden Aufgaben zur Verfügung stellen zu können. Jedoch weisen die Ergebnisse von Bailey und Scerbo (2007) darauf hin, dass das Vertrauen in die Automation auch hier eine Rolle gespielt haben könnte. In ihrer Studie nahm die Überwachungsleistung der Probanden mit zunehmender Komplexität des zu überwachenden Subsystems ab, was gut zu den bereits erwähnten Befunden passt. Die Überwachungsleistung war jedoch direkt mit dem Vertrauen in das jeweilige Subsystem verbunden (hohes Vertrauen – geringe Überwachungsleistung). In ähnlicher Weise wie die Komplexität der Überwachungsaufgabe selbst bzw. das Vorhandensein konkurrierender Aufgaben *complacency* zu beeinflussen vermag, vermuten Lee und See (2004) und Parasuraman et al. (1993), dass auch eine hohe mentale Beanspruchung bzw. Ermüdung *complacency* fördern.

Ein weiterer situativer Aspekt, der als möglicher moderierender Faktor diskutiert wird, ist die Höhe des in einer bestimmten Situation wahrgenommenen Risikos, welches mit einem Übersehen von Automationsfehlern verknüpft ist (*perceived risk*, Lee & See, 2004). Danach ist anzunehmen, dass eine übermäßige Reduktion der Überwachungsintensität vor allem dann vorgenommen wird, wenn das wahrgenommene Risiko als gering eingestuft wird. Vor dem Hintergrund der bereits in Kapitel 1.3.2 beschriebenen Befunde, denen zufolge Automationsfehler v. a. dann vertrauensmindernd wirken, wenn sie mit hohen Kosten verbunden und schlecht kompensierbar sind, scheint diese Annahme durchaus plausibel.

Schließlich könnte auch die Konstanz der Funktionsallokation über die Zeit für das Auftreten von *complacency* eine Rolle zu spielen. Befunde von Parasuraman, Mouloua und Molloy (1996) zeigen, dass sich die Leistung bei der Entdeckung von Automationsfehlern dann deutlich steigern

lässt, wenn die automatisierten Funktionen während der Bearbeitung für jeweils zeitlich begrenzte Abschnitte immer wieder an den Operateur „zurückgegeben“ werden, d.h. manuell ausgeführt werden müssen. Dabei lässt sich auf Basis der vorliegenden Daten allerdings nur schwer entscheiden, ob es sich bei diesem Effekt wirklich um eine Reduktion von *complacency*-Effekten handelt, oder ob nicht andere Faktoren – z.B. eine Vermeidung von Vigilanzeinbußen oder die bessere Aufrechterhaltung eines adäquaten mentalen Modells der Systemfunktionen – dafür verantwortlich sind.

1.4.3 Integratives Rahmenmodell des Konzepts *complacency*

Zusammenfassend lässt sich festhalten, dass offenbar drei verschiedene Einflussbereiche – Merkmale der Automation, personale Aspekte und der situative Kontext – bei der Entstehung von *complacency* zu berücksichtigen sind. Dabei ist das genaue Zusammenwirken dieser drei Faktoren bisher nur unzureichend untersucht worden. Einen ersten Versuch, die bisher vorliegenden Befunde zu einem integrativen Modell zusammenzufassen, haben Manzey und Bahner (2005; siehe auch Bahner & Manzey, 2004) vorgelegt (siehe Abb. 4).

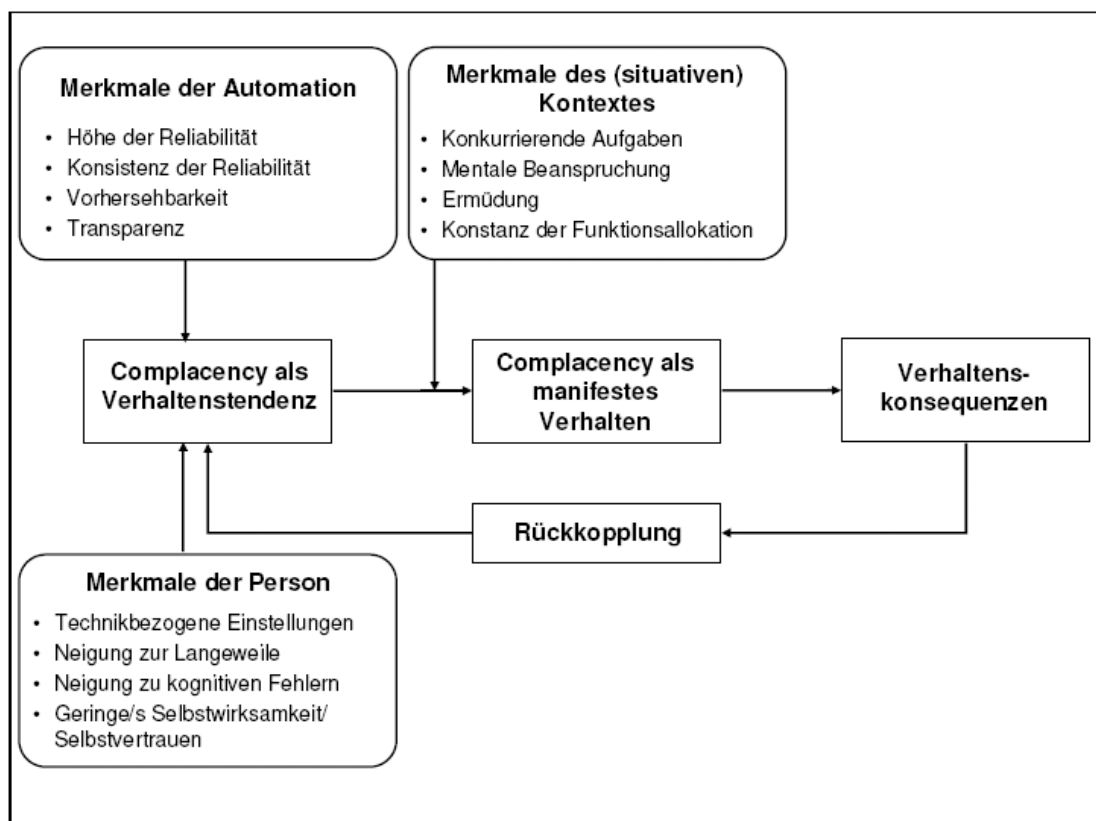


Abb. 4: Rahmenmodell des Konzepts *complacency* (aus Manzey & Bahner, 2005, S. 103)

Unterschieden werden zwei verschiedene Facetten von *complacency*, die als Verhaltenstendenz bzw. manifestes Verhalten bezeichnet werden. Die Verhaltenstendenz repräsentiert die grund-

sätzliche Neigung einer Person, die im Umgang mit einem bestimmten automatisierten System geforderten Überwachungsleistungen zu vernachlässigen. Im Sinne einer „Anfälligkeit“ für ein defizitäres Überwachungsverhalten wird sie maßgeblich von zwei Faktoren und deren Wechselwirkung bestimmt: Zum einen von Personenmerkmalen, die zeitlich überdauernd und situativ invariant sein (z.B. generalisierte technikbezogene Einstellungen), oder sich auf eine konkrete Aufgabe beziehen können (z.B. Selbstwirksamkeitserwartung), und zum anderen von den Merkmalen der Automation, insbesondere ihrer Reliabilität.

Als eine Verhaltenstendenz stellt sie allerdings nur eine notwendige, nicht jedoch hinreichende Bedingung für eine Manifestation von *complacency* im tatsächlichen Überwachungsverhalten dar. Diese zeigt sich erst dann, wenn zusätzlich bestimmte situative Faktoren, wie z.B. das Erfordernis, konkurrierende Ziele zu verfolgen, gegeben sind. In diesem Sinne moderieren situative Merkmale also den Zusammenhang zwischen der Verhaltenstendenz und dem beobachtbaren Verhalten. Zusätzlich wird in dem Modell angenommen, dass es zwischen den Verhaltenskonsequenzen und der Verhaltenstendenz eine Rückkopplung gibt. Gerade bei hoch reliablen Systemen, in denen Automationsfehler äußerst seltene Ereignisse darstellen, wird der Bediener wiederholt die Erfahrung machen, dass eine – im Vergleich zu einer optimalen Strategie – zu seltene Überwachung des Systems keine negativen Verhaltenskonsequenzen nach sich zieht. Dies kann im Sinne einer positiven Rückkopplung zu einem Lernprozess führen, der die Tendenz zu einem übersteigerten Vertrauen in die Automation und zu einem entsprechenden Verhalten verstärkt. Ein solches Muster konnte von Bailey und Scerbo (2007) auch empirisch gezeigt werden: Mit zunehmender Dauer der Überwachung zeigte sich ein deutlicher Einbruch der Überwachungsleistung und damit verbunden auch eine deutlich verminderte Wahrscheinlichkeit, einen Automationsfehler zu entdecken. Eine konzeptuelle Nähe besteht hier auch zum Konzept der gelernten Sorglosigkeit (Frey & Schulz-Hardt, 1997), das Lüdtke und Möbus (2002) bereits in einen ähnlichen Zusammenhang mit Problemen bei der Mensch-Maschine-Interaktion gebracht haben. Umgekehrt ist davon auszugehen, dass – in der Regel selten auftretende – negative Verhaltenskonsequenzen, z.B. das Übersehen von folgenschweren Automationsfehlern, das Vertrauen in die Automation abrupt mindern und damit die Wahrscheinlichkeit, *complacency* als Verhalten zu zeigen, stark reduzieren.

Insgesamt betrachtet ist jedoch die bisher vorliegende empirische Basis zum Phänomen *complacency* noch sehr begrenzt und auch das skizzierte Rahmenmodell von Manzey und Bahner (2005) hat in weiten Teilen einen hypothetischen Charakter. Eine dieser Begrenzungen betrifft die in dem Modell genannten einzelnen Automations-, Personen- und Kontextfaktoren, deren Wirkung auf die verschiedenen Aspekte von *complacency* und mögliche Interaktionen bisher nur

ansatzweise aufgeklärt sind. Das Modell kann damit keinen Anspruch auf Vollständigkeit erheben, sondern eher als heuristischer Ansatzpunkt für weitere Forschungsarbeiten in diesem Bereich genutzt werden. Eine zweite Begrenzung betrifft methodische Probleme, die die bisherige Forschung zum Phänomen *complacency* aufweist. Auf diese soll im Folgenden genauer eingegangen werden.

1.4.4 Methodische Probleme bisheriger Studien zum Phänomen *complacency*

Die nachfolgenden Ausführungen beziehen sich auf die Operationalisierung des Konstrukts, sodass noch einmal die hier verwendete Definition ins Gedächtnis gerufen werden soll (vgl. Kap. 1.4.1): *Complacency* wird als ein Merkmal der Mensch-Maschine-Interaktion verstanden, das sich auf der Verhaltensebene in einer unzureichenden Überwachung automatisierter Systeme durch den Menschen widerspiegelt, mit einem übersteigerten Vertrauen in die Funktions- und Leistungsfähigkeit der Automation zusammenhängt und als Folge zu einem Übersehen kritischer Systemzustände führen kann.

Ein Problem der bisherigen *complacency*-Studien besteht darin, dass auf eine Erfassung der subjektiven Komponente des Konstrukts im Sinne eines übersteigerten Vertrauens in die Funktions- und Leistungsfähigkeit der Automation meist verzichtet wurde. Überwachungsdefizite – operationalisiert über das Übersehen von Automationsfehlern – könnten jedoch auch unabhängig von einem übersteigerten Vertrauen auftreten und wären damit definitionsgemäß nicht mehr als *complacency*-Effekt zu werten. Eine Ausnahme bildet hier die Studie von Bailey und Scerbo (2007), in der der oft postulierte Zusammenhang von Überwachungsverhalten und Vertrauen in Automation auch empirisch belegt wurde.

Eine zweite grundsätzliche, methodische Schwäche der bisherigen Arbeiten muss auch in der verhaltensbasierten Operationalisierung von *complacency* gesehen werden. In fast allen vorliegenden Untersuchungen wurde das Übersehen von Automationsfehlern als empirischer Indikator für *complacency* genutzt. Wie Moray (Moray, 2003; Moray & Inagaki, 2000) aber zeigt, ist das Übersehen von Automationsfehlern zwar eine mögliche Folge von *complacency*, es kann darin aber keine hinreichende Bedingung für eine unzureichende Überwachung gesehen werden. So ist es durchaus möglich, dass Automationsfehler trotz sorgfältiger Überwachung eines automatisierten Systems übersehen werden. Moray (2003) verdeutlicht dies anhand eines Gedankenexperiments: Angenommen, die Aufgabe eines Operators ist es, mehrere Subsysteme zu überwachen, wie es ja auch der für *complacency* charakteristischen Mehrfachaufgabensituation entspricht. Geht man nun davon aus, dass in Subsystem A potenziell kritische Veränderungen – entsprechend der Dynamik der zu überwachenden Parameter – alle 5 s möglich sind, während in Subsystem B dies alle

7 s der Fall ist, so würde der Operateur idealerweise zu Sekunde „5“, „10“, „15“ etc. Subsystem A und zu Sekunde „7“, „14“, „21“ etc. Subsystem B kontrollieren. Dies ist zunächst problemlos möglich, bis bei Sekunde „35“ beide Subsysteme kontrolliert werden müssten, der Operateur sich aber für eines der beiden Systeme entscheiden muss und dadurch trotz optimaler Überwachungsstrategie einen ggf. im anderen System vorliegenden Fehler übersehen würde.

Zweifelsohne ist dies eine etwas überspitzte Darstellung; die grundlegende Problematik wird aber deutlich. Die Tatsache, dass Fehler einer Automation übersehen werden, muss nicht zwangsläufig in einer unzureichenden Überwachung des Systems, und damit *complacency*, begründet sein.

Eine in diesem Sinne überzeugendere Operationalisierung von *complacency* würde voraussetzen, dass das tatsächliche Überwachungsverhalten einer Person über geeignete Methoden (z.B. Logfiles, Blickbewegungsmessung etc.) erfasst und dann mit einem normativen Modell verglichen wird, das beschreibt, wie ein angemessenes Überwachungsverhalten der gegebenen Automation aussehen sollte (Moray & Inagaki, 2000; Moray, 2003). Bagheri und Jamieson (2004a; b) griffen diesen Ansatz teilweise auf und kontrastierten das üblicherweise verwandte Fehlerdetektionsmaß mit Blickbewegungsdaten (*Mean Time Between Fixations*). Dabei zeigte sich, dass erst die Blickbewegungsdaten eine detaillierte Analyse des Überwachungsverhaltens ermöglichten. So zeigte sich in der einen Studie (Bagheri & Jamieson, 2004b) nur in den Blickbewegungsdaten, nicht aber mit Blick auf das herkömmliche Fehlerdetektionsmaß ein Effekt der Bedingungsmanipulation. In der anderen Studie (Bagheri & Jamieson, 2004a) traten erst durch die Analyse der Blickbewegungsdaten wesentliche Unterschiede der Überwachungsleistung im zeitlichen Verlauf zutage. Allerdings wurde in diesen Arbeiten kein „normatives Modell“ einer optimalen Überwachungsstrategie definiert – wenngleich dies der Schwierigkeit geschuldet sein mag, bei der gewählten kontinuierlichen Überwachungsaufgabe überhaupt optimale Abruffrequenzen bezüglich der zu überwachenden Parameter begründen und festlegen zu können.

Aus den bisherigen Ausführungen lassen sich folgende Anforderungen für eine angemessene Operationalisierung ableiten: Erstens ist *complacency* direkt über das Überwachungsverhalten und nicht über mögliche Verhaltenskonsequenzen wie etwa das Übersehen von Automationsfehlern zu operationalisieren. Dabei muss es zweitens Ziel sein, eine Experimentalumgebung zu wählen, die die Formulierung eines normativen Modells optimalen Überwachungsverhaltens erlaubt, um so in Abgleich zum tatsächlich gezeigten Überwachungsverhalten über das Vorliegen von *complacency* entscheiden zu können. Drittens ist neben objektiven Verhaltensmaßen auch das subjektive Vertrauen in Automation sowie das Selbstvertrauen in die manuelle Ausführung der entsprechenden Funktion zu erfassen. In der vorliegenden Arbeit soll dies umgesetzt werden.

Nachdem in den vorausgehenden Kapiteln der Versuch unternommen wurde, den bisherigen Stand theoretischer Überlegungen wie auch des empirischen Forschungsstandes zum Phänomen *complacency* umfassend darzustellen, soll im Folgenden das zweite, für die vorliegende Arbeit zentrale Phänomen, *automation bias*, Gegenstand der Betrachtung sein.

1.5 AUTOMATION BIAS

1.5.1 Definition des Begriffs *automation bias*

Automation bias stellt wie *complacency* ein Phänomen dar, das in Zusammenhang mit einem übersteigerten Vertrauen in automatisierte Systeme gebracht werden kann. Während *complacency*-Effekte insbesondere im Kontext von Aufgaben der kontinuierlichen Überwachung von automatisierten Systemen beschrieben werden, findet der Begriff des *automation bias* seinen Ursprung im Umgang mit Assistenzsystemen. Solche Systeme haben mittlerweile Einzug in viele Arbeitsdomänen gehalten, um den Menschen bei komplexen Diagnose-, Planungs- oder Entscheidungsaufgaben zu unterstützen (z.B. bei der Erkennung und Behebung von Störungen in komplexen technischen Systemen, der Konfliktlösung in Flugleitzentralen oder der Überwachung von Patienten in der Intensivmedizin). Doch auch in den privaten Bereich finden Assistenzsysteme Einzug, wie beispielsweise im Fahrzeug in Form von Navigationssystemen. Berichte in der Tagespresse von Autofahrern, die ihrem Navigationssystem „blind“ vertrauen und ihre Fahrt unbeabsichtigt im Fluss, im Acker oder auch im falschen Land beenden, liefern einen anekdotischen Beleg dafür, dass auch hier ein übersteigertes Vertrauen negative Verhaltenskonsequenzen nach sich ziehen kann. Diese und ähnliche Vorkommnisse in der Interaktion mit automatisierten Assistenzsystemen werden von Mosier und Skitka (1996) unter dem Begriff *automation bias* zusammengefasst. Allgemein verstehen sie darunter „the tendency to use automated cues as a heuristic replacement for vigilant information seeking and processing“ (S. 205). Als mögliche Folgen werden zwei Fehlertypen beschrieben: *Omission* und *commission* Fehler. Ein *omission* Fehler liegt dann vor, wenn ein Operateur bzw. Nutzer eines Assistenzsystems einen kritischen Systemzustand, der vom Assistenzsystem nicht angezeigt wird, übersieht. Ein *commission* Fehler besteht dagegen im Befolgen fehlerhafter Direktiven eines Assistenzsystems. Als mögliche Ursache hierfür nennen Skitka, Mosier und Burdick (1999) „not seeking out confirmatory or disconfirmatory information, or discounting other sources of information in the presence of computer-generated cues“ (p. 993).

Die im Exkurs in Kapitel 1.3.3 beschriebene Havarie der „Royal Majesty“ liefert für beide Fehlertypen Beispiele: Begreift man die Abdrift vom geplanten Kurs als kritischen Systemzustand, der durch die Automation (hier: GPS) nicht angezeigt wurde, so mag im Nichterkennen

desselben für eine Dauer von mehr als 34 Stunden ein wiederholter (die Positionsdaten hätten nach Vorschrift stündlich abgeglichen werden müssen) *omission* Fehler gesehen werden. Die Ausrichtung der Navigation an der fehlerhaften Positionsinformation kann als Befolgen der automatisch generierten Hinweise und damit als *commission* Fehler gewertet werden, wobei in diesem Beispiel offenbar beide von Skitka et al. (1999) genannten Ursachen beigetragen haben: Zum einen unterblieb eine Kontrolle der Positionsinformationen anhand anderer verfügbarer Quellen (LORAN-C, Sichtprüfung der Tonnen). Zum anderen wurden widersprechende Informationen (Lichter werden wahrgenommen, die nicht da sein dürften; wichtige Markierungstonnen werden weder vom Radar angezeigt noch direkt gesehen) zwar registriert, aber offenbar nicht berücksichtigt oder als Falsifikation der automatisch generierten Informationen gewertet.

Mosier, Skitka, Heers und Burdick (1998) sehen *automation bias* in der allgemeinen Tendenz begründet, den kognitiven Aufwand für eine Aufgabe soweit wie möglich zu reduzieren. Die Nutzung von Heuristiken ist dabei ein möglicher Weg, um komplexe Entscheidungen zu vereinfachen (Tversky & Kahneman, 1984). Wird der Mensch nun durch ein Entscheidungsassistenzsystem unterstützt, das in der Regel angemessene Empfehlungen generiert, mag das Verlassen auf diese als zeit- und kostensparende Heuristik erscheinen. Dabei wird – je nach erforderlichem Aufwand – entweder ganz auf den Abruf anderer verfügbarer Informationen, die eine Überprüfung der automatisch generierten Empfehlung ermöglichen würden, verzichtet, oder aber andere Information verzerrt verarbeitet. Offenbar können hier verschiedene Verzerrungsmechanismen zum Tragen kommen. Mosier et al. (1998) unterscheiden diesbezüglich drei Formen: *assimilation*, *confirmational* und *discounting bias*. Erstgenannter greift dann, wenn Umgebungsinformationen durch Ambiguität gekennzeichnet sind, aber als automationskonsistent interpretiert werden, obwohl auch andere Interpretationen zulässig gewesen wären. Der *confirmational bias* besteht darin, dass nur Informationen wahrgenommen und verarbeitet werden, die der automatisch generierten Empfehlung entsprechen, während widersprechende Informationen ausgeblendet werden. Bisweilen werden widersprechende Informationen aber auch abgerufen und registriert, dann jedoch in ihrer Bedeutung abgewertet und bei der letztendlichen Entscheidungsfindung nicht berücksichtigt (*discounting bias*). Auch gibt es Hinweise darauf, dass Probanden sich bisweilen an Informationen erinnern, die die Automation vermeintlich bestätigen, obwohl diese nicht vorhanden waren (*false memory*; Mosier et al., 1998).

Folgt man Mosier et al. (1998), so ist davon auszugehen, dass die Einführung von Assistenzsystemen immer die Gefahr mit sich bringt, dass Operateure deren Empfehlungen nicht als einen zusätzlichen Hinweis zur Unterstützung der eigenen Entscheidungsfindung nutzen, sondern die eigene aktive Entscheidungsfindung durch die Automation „ersetzt“ wird.

Inwieweit diese Annahme auch empirisch belastbar ist und welche Faktoren das Risiko von *omission* und *commission* Fehlern reduzieren können, ist Gegenstand des nachfolgenden Kapitels.

1.5.2 Empirische Studien des Phänomens *automation bias*

Mosier, Palmer und Degani (1992) lieferten einen der ersten empirischen Nachweise von *automation bias* bzw. genauer des *commission* Fehlers. In einer Flugsimulationsumgebung (Advanced Concept Flight Simulator, ACFS) legte eine automatisch vom System abgearbeitete Checkliste es nahe, dass das linke Triebwerk Feuer gefangen habe und daher abzuschalten sei. Allerdings wiesen andere verfügbare Triebwerk-Parameter darauf hin, dass nicht das linke, sondern das rechte Triebwerk defekt war. Dennoch entschieden sich 75 % der Crews (Berufspiloten) fälschlicherweise dafür, das linke Triebwerk abzuschalten – im Vergleich zu nur 25 % in einer Vergleichsbedingung, der eine rein papierbasierte, nicht-automatisierte Checkliste zur Verfügung gestellt wurde. Zudem zeigte die Analyse von Tonbandaufzeichnungen, dass in der automatisierten Bedingung im Entscheidungsvorfeld weniger diskutiert wurde und offenbar weniger Informationen einbezogen wurden, um eine Entscheidung zu fällen. Hier führte also das bloße Vorhandensein einer automatisierten Assistenz zu einer deutlichen Reduktion der Informationssuche und in der Folge zu *commission* Fehlern, die ohne Assistenz nicht vorkamen. Einschränkend ist jedoch darauf hinzuweisen, dass in dieser Studie die Stichprobe aus nur $N = 12$ Zweipersonen-Crews bestand und somit keine statistische Absicherung des Effekts möglich war.

Auch in anderen Studien wurde untersucht, welchen Einfluss die Einführung von Assistenzsystemen auf potenzielle *omission* und *commission* Fehler hat. Als Versuchsumgebung diente in Studien mit studentischen Stichproben oftmals eine an Flugaufgaben orientierte, aber stark vereinfachte Simulation (Workload/PerformANcE Simulation software, W/PANES; z.B. beschrieben in Skitka et al., 1999). In dieser Simulation bearbeiten die Probanden parallel drei Aufgaben: Eine Tracking-Aufgabe, eine Wegpunktaufgabe, die das wiederkehrende Quittieren von Wegpunkten, die auf einer Karte ebenso wie die eigene Position angezeigt werden, vorsieht, und eine Überwachungsaufgabe, bei der Abweichungen von vier Parameteranzeigen in den jeweils kritischen Bereich per Tastendruck korrigiert werden müssen. Sowohl die Wegpunkt- als auch die Überwachungsaufgabe können dabei durch ein Assistenzsystem unterstützt werden. Bei auftretenden Ereignissen (Abweichung der Parameter, Passieren eines Wegpunktes) gibt das System zum einen an, welches Ereignis vorliegt (z.B. „Temp is failing“) und welche Taste dementsprechend zu drücken ist (z.B. „Press button 3“).

Mit dieser Versuchsumgebung verglichen Skitka et al. (1999) die Leistung von Probanden mit und ohne Assistenzsystem zur Unterstützung der Wegpunkt- und Überwachungsaufgabe.

Sechs von hundert Ereignissen (Abweichungen der Parameter bzw. Passieren von Wegpunkten) wurden in der automatisierten Bedingung vom Assistenzsystem nicht angezeigt – und in der Folge davon im Mittel 41 % von den Probanden übersehen (*omission* Fehler). Hingegen wurden in einer nicht-automatisierten Kontrollbedingung nahezu alle dieser kritischen Systemzustände erkannt. Für die Analyse wurden genau dieselben sechs Ereignisse herangezogen, bei denen in der automatisierten Bedingung kein automatischer Hinweis erfolgte. Ein Vergleich der beiden Gruppen über jene insgesamt 88 Ereignisse, bei denen das Assistenzsystem korrekte Hinweise lieferte, zeigte hingegen einen Vorteil der automatisierten Bedingung (im Mittel 83 vs. 72 korrekte Reaktionen). In Bezug auf die Trackingaufgabe ergaben sich keine Unterschiede zwischen den experimentellen Gruppen.

In weiteren sechs Fällen zeigte das Assistenzsystem falsche Hinweise an. Diese hätten anhand bloßer Inspektion des entsprechenden Bildschirmareals einfach und eindeutig als fehlerhaft identifiziert werden können, zumal die Probanden darüber informiert worden waren, dass das Assistenzsystem fehlerhaft arbeiten könne, während die übrigen Systemanzeigen zu 100 % zuverlässig seien. Dennoch folgten in 65 % der Fälle die Probanden der Automation und über 99 % der Probanden begingen zumindest in einem der sechs Fälle einen *commission* Fehler.

Zusammenfassend lässt sich festhalten, dass obwohl die Probanden in der automatisierten Bedingung darauf hingewiesen wurden, dass das Assistenzsystem hoch, aber nicht perfekt zuverlässig arbeitet, *omission* und *commission* Fehler mit Raten von 41 bzw. 65 % auftraten. Auch in einer anderen Studie mit derselben Versuchsumgebung wurden vergleichbare Fehlerraten von 31 (*omission*) und 54 % (*commission*) gefunden (Skitka, Mosier, Burdick & Rosenblatt, 2000). Bei realitätsnäheren Szenarien steigt der Anteil sogar auf 48 bzw. 55 % für *omission* Fehler und 93 bzw. 100 % für *commission* Fehler (Mosier, Skitka, Dunbar & McDonnel, 2001; bzw. Mosier et al., 1998). Diese Zahlen verdeutlichen die Relevanz des Problems: Die Einführung von Assistenzsystemen scheint nicht nur die Gefahr von automationsbezogenen Fehlern in sich zu bergen, sondern diese sogar sehr wahrscheinlich zu machen.

Doch welche Aspekte können die drohende Gefahr von *omission* und *commission* Fehlern abwenden oder zumindest mindern? Untersucht wurden in diesem Zusammenhang die Aspekte der Verantwortungsübernahme, verschiedener Trainingsalternativen und der Aufgabenbearbeitung im Team statt durch eine Einzelperson. Diese Studien und ihre Ergebnisse werden nachfolgend dargestellt.

1.5.2.1 Automation bias und Verantwortungsübernahme

Skitka, Mosier und Burdick (2000) untersuchten, inwiefern *commission* und *omission* Fehler durch experimentell variierte Grade von sozialer Verantwortungsübernahme beeinflusst werden können. Als Versuchsumgebung diente wiederum die bereits im vorangegangenen Kapitel (1.5.2.) beschriebene W/PANES-Simulation. Im Unterschied zu allen übrigen Studien wurde hier jedoch auch versucht, das Verifikationsverhalten der Probanden zu operationalisieren. Hierzu diente ein gesonderter Bildschirm, auf dem eine Zellenmatrix dargestellt war. Immer dann, wenn das Assistenzsystem ein Ereignis (passierter Wegpunkt bzw. Parameterabweichung) mit korrespondierender Reaktionsempfehlung (z.B. „Press button 3“) meldete, konnte der Proband die Reaktionsempfehlung mithilfe der Matrix verifizieren. Nach Klick(s) auf eine bis maximal drei Zellen der Matrix erhielt er die 100-prozentig zuverlässige Information, welche Taste zu drücken ist. Die Probanden wurden in fünf Gruppen aufgeteilt, wobei eine der Gruppen als Kontrollbedingung ohne Verantwortungsübernahme diente. Die übrigen vier Gruppen erhielten alle die Instruktion, dass ihre Leistungsdaten aufgezeichnet würden und sie ihre Leistung mit Blick auf eine der folgenden Hauptaufgaben maximieren sollten: a) Gesamtleistung, b) Genauigkeit oder c) Reaktionsgeschwindigkeit bezüglich der Überwachungs- und Wegpunktaufgabe oder d) Tracking-Leistung. Zudem erhielten alle Probanden der Experimentalbedingungen die Information, dass sie sich im Anschluss an das Experiment im Rahmen eines Interviews für die von ihnen erreichte Leistung erklären und rechtfertigen müssten. Die Ergebnisse zeigen, dass Probanden, die für die Gesamtleistung oder für die Genauigkeit ihrer Entscheidungen verantwortlich gemacht worden waren, insgesamt weniger *commission* und *omission* Fehler begingen. *Commission* Fehler wurden über das Befolgen einer falschen Direktive des Assistenzsystems, *omission* Fehler über das Nichtentdecken von Wegpunkt- oder Parameterereignissen, die das Assistenzsystem nicht anzeigte, operationalisiert. Zudem verifizierten die Probanden dieser beiden Gruppen die Hinweise des Assistenzsystems auch in stärkerem Umfang als jene Gruppen, die für die Reaktionsschnelligkeit, die Tracking-Leistung oder für keine Aufgabe verantwortlich gemacht worden waren. Interessanterweise gingen mit der beschriebenen Minderung von *automation bias*-Effekten keine negativen Kosteneffekte bezüglich der Reaktionsgeschwindigkeit und der Tracking-Aufgabe einher. Vor diesem Hintergrund scheint ein möglicher Ansatzpunkt zur Reduktion von *automation bias* darin zu liegen, dem Operateur zu verdeutlichen, dass er für die Gesamtleistung und – sofern übertragbar – für die Treffsicherheit seiner Entscheidungen – verantwortlich ist. Allerdings dürften dem von Skitka, Mosier und Burdick (2000) aufgebauten Rechtfertigungsdruck in der Praxis auch gewichtige Argumente entgegenstehen, wie etwa die Fokussierung auf den Menschen als Fehlerursache.

Hinzu kommt, dass in einer ähnlichen Studie, in der jedoch Piloten als Probanden dienten und eine vergleichsweise realitätsnähere Flugsimulationsaufgabe (Advanced Concepts Flight Simulator, ACFS) eingesetzt wurde, die Induktion von Verantwortungsübernahme für den Umgang mit automatisierten Cockpit-Systemen keine signifikanten Effekte zeigte (Mosier et al., 1998). Allerdings beschrieben sich jene Probanden, die weniger *omission* Fehler begingen, in einer Nachbefragung als stärker verantwortlich für ihre Leistung. In ähnlicher Weise ergab sich auch in der Studie von Skitka et al. (1999) ein positiver Zusammenhang von *omission* und *commission* Fehlerrate und dem Ausmaß, in dem die Probanden berichteten, die Verantwortung dem Assistenzsystem überlassen zu haben. Dies könnte darauf deuten, dass die Instruktion von Mosier et al. (1998) ungeeignet war, um Verantwortlichkeit zu induzieren – ob dem tatsächlich so ist, muss jedoch offen bleiben, da keine Ergebnisse bezüglich eines möglichen Manipulationschecks anhand der subjektiven Daten berichtet werden.

Zusammengefasst ist davon auszugehen, dass ein ausgeprägtes Verantwortungsgefühl für die jeweilige Tätigkeit sich positiv auf die Leistung auswirkt und mit einer geringeren *omission* und *commission* Fehlerrate verbunden ist. Ob eine solche Verantwortungsübernahme jedoch per Instruktion bzw. Training erreicht werden kann, bleibt fraglich. Möglicherweise handelt es sich dabei auch um ein interindividuell unterschiedlich stark ausgeprägtes Personenmerkmal, das sich nur sehr begrenzt induzieren lässt.

1.5.2.2 Automation bias und mögliche Trainingsinterventionen

Skitka, Mosier, Burdick und Rosenblatt (2000) untersuchten, inwiefern durch bestimmte Trainingsmaßnahmen *automation bias*-assoziierten Problemen vorgebeugt werden kann. Als Versuchsumgebung diente wiederum W/PANES (vgl. Kap. 1.5.2.) und als Probanden Studierende. Drei Trainings- bzw. Systemvarianten wurden verglichen. Eine erste Gruppe erhielt ein Training, in dem das Problem des *automation bias* explizit thematisiert wurde. Eine zweite Gruppe wurde instruiert, dass die Direktiven des Assistenzsystems überprüft werden *müssen* (anhand der verfügbaren Systemindikatoren; eine Zellenmatrix zur Verifikation wurde in dieser Studie nicht eingesetzt), während die dritte Gruppe die Information erhielt, dass sie die Hinweise überprüfen *könnten*. Jeweils die Hälfte der Probanden erhielt zusammen mit einem Hinweis des Assistenzsystems die Aufforderung, diesen zu verifizieren, auf dem Bildschirm eingeblendet. Alle Probanden wurden darüber informiert, dass das Assistenzsystem fehlerhaft arbeiten kann, während die übrigen Systemindikatoren zu 100 % zuverlässig sind.

Die Ergebnisse zeichnen folgendes Bild: Erstens hatte die direkte bildschirmgestützte Aufforderung zur Verifikation keinen Einfluss auf das Verhalten der Probanden. Zweitens wirkte

sich die Trainingsvariation nicht auf den *omission* Fehler, sondern ausschließlich auf den *commission* Fehler aus. Dabei erwies sich das erste Training, indem die Probanden explizit auf die Problematik des *automation bias* hingewiesen wurden, den anderen Varianten überlegen und führte im Vergleich zu den anderen Trainingsalternativen zu einer geringeren Anzahl von *commission* Fehlern.

Mosier et al. (2001) überprüften die Ergebnisse der Studie von Skitka, Mosier, Burdick und Rosenblatt (2000) an einer Stichprobe von Piloten und unter Verwendung der realitätsnäheren Flugsimulationsaufgabe ACFS. Auch hier hatte die direkte Bildschirmaufforderung zur Verifikation keinen Einfluss. Allerdings verschwanden hier auch die Unterschiede zwischen den verschiedenen Trainingsbedingungen: Ein explizites *automation bias* Training vermochte es hier nicht, die *commission* Fehlerrate zu reduzieren.

Offenbar handelt es sich bei *omission* und *commission* Fehlern um relativ robuste Phänomene. Weder durch die Vorabinstruktion, dass das Assistenzsystem überprüft werden *muss*, noch durch die wiederholte Aufforderung hierzu während des Versuchs wurde die *commission* Fehlerrate beeinflusst. Ob dies daran lag, dass die Probanden diese Aufforderung schlicht ignorierten oder aber den Versuch einer Überprüfung machten, die abgerufenen Informationen dann aber verzerrt verarbeiteten, kann nicht beurteilt werden, da das Verifikations- bzw. Überprüfungsverhalten in keiner der beiden Studien erfasst wurde. Mit etwas besseren Erfolgsaussichten verbunden scheint dagegen die Aufklärung über *automation bias* im Training. Allerdings ist auch hier die Befundlage heterogen; nur die Studierenden, nicht aber die Piloten profitierten davon.

1.5.2.3 Automation bias: Teams versus Einzelpersonen

In den beiden im vorausgehenden Abschnitt geschilderten Studien (Skitka, Mosier, Burdick & Rosenblatt, 2000; Mosier et al., 2001) wurde als weiterer experimenteller Faktor auch das Bearbeitungssetting variiert: Ein Teil der Probanden arbeitete individuell, ein anderer im Zweierteam. Zugrunde lag die sozialpsychologisch begründbare Vermutung, dass durch das Hinzukommen einer zweiten Person Leistungsveränderungen eintreten. Diese können jedoch in beide Richtungen gehen. Leistungssteigerungen sind etwa denkbar durch *social facilitation* (Zajonc, 1965). Leistungsminderungen hingegen könnten sich im Vergleich zur individuellen Bearbeitung durch Verantwortungsdiffusion oder *social loafing* (Williams & Karau, 1991) ergeben. In beiden Studien zeigten Teams und Individuen jedoch keinerlei Unterschiede mit Blick auf *commission* oder *omission* Fehler. Die Unterstützung durch eine zweite Person bei der Überwachung des Systems führte weder zu einer Steigerung noch zu einer Reduktion von *omission* und *commission* Fehlern. Dagegen konnten in einer Studie von Domeinski, Wagner, Schöbel und Manzey (2007) durch die Information, dass eine zweite Person dieselbe Aufgabe räumlich getrennt parallel bearbeitet, eine Leis-

tungsminderung festgestellt werden. Diese bestand darin, dass die vom Assistenzsystem vorgeschlagenen Direktiven in dieser Redundanzbedingung im Vergleich zur individuellen Bearbeitung in geringerem Maße überprüft wurden. Wenn die Probanden die Information erhielten, dass die zweite Person ihre Leistung bezüglich der betreffenden Aufgabe als gering einschätzte, verringerte sich der im Sinne des *social loafing* interpretierbare Effekt, er verschwand jedoch nicht. Dieser Befund passt gut zu den geschilderten Ergebnissen zur Verantwortungsübernahme (vgl. Kap. 1.5.2.1). Nahe liegt die Interpretation, dass sich die Probanden von Domeinski et al. (2007) durch das Wissen, dass zusätzlich eine zweite Person ihrer Tätigkeit nachging, weniger verantwortlich dafür fühlten. Allerdings wurden in der Studie keine *commission* und *omission* Fehler untersucht, sodass diesbezüglich keine Aussage getroffen werden kann.

Auch mit Blick auf den Einfluss von Team- vs. individueller Arbeit auf *automation bias* ist die Befundlage somit uneindeutig. Vorsichtig interpretiert kann jedoch angenommen werden, dass Teams zumindest nicht besser bei der Überprüfung von Assistenzsystemen und im Vermeiden von *commission* und *omission* Fehlern im Vergleich zu Individuen sind.

Neben den Hauptbefunden zu den Themen Verantwortung, Training und Teams sind den bisherigen *automation bias* Studien noch einige weitere interessante Ergebnisse zu entnehmen, wenngleich diese eher explorativer Natur sind. Mosier et al. (1998) fanden einen positiven Zusammenhang zwischen *omission* Fehlerrate und Flugstunden bzw. den Flugerfahrungsjahren der Probanden. Möglicherweise zeigt sich hier eine Parallele der bereits in Kapitel 1.4.3 beschriebenen positiven Rückkopplungsschleife: Mit der Erfahrung, dass die Automation im Cockpit fehlerfrei funktioniert, wird eine Überprüfung derselben offenbar zunehmend als überflüssig erachtet und unterlassen. Wird die Automation jedoch als fehlerhaft arbeitend wahrgenommen, ändert sich das Bild: In der Studie von Skitka et al. (1999) ergab eine Nachbefragung der Probanden einen negativen Zusammenhang zwischen der *omission* sowie *commission* Fehleranzahl und dem geschätzten Prozentsatz falscher Direktiven des Assistenzsystems. Hierzu passt auch der auch in dieser Studie gefundene positive Zusammenhang mit der eingeschätzten Genauigkeit des Assistenzsystems. Offenbar werden *omission* und *commission* Fehler eher begangen, wenn das Assistenzsystem subjektiv als wenig fehleranfällig und genau arbeitend empfunden wird. Eine nahe liegende Begründung hierfür wäre, dass mit jedem wahrgenommenen Automationsfehler das Vertrauen in das System sinkt und damit die Bereitschaft bzw. subjektive Notwendigkeit, das System zu überprüfen, größer wird. Andersherum dürfte auch der Verzicht auf Überprüfungen und damit einhergehende *omission* und *commission* Fehler bei einem als hochzuverlässig empfundenen System zumindest in Teilen über ein übersteigertes Vertrauen erklärbar sein.

Ein weiterer bemerkenswerter Befund zeigte sich in der Studie von Mosier et al. (1998) bei einer Nachbefragung der Probanden bezüglich des *commission* Fehlers, den alle Probanden begingen: Hier gaben 67 % der Probanden an, sich an zusätzliche Hinweise zu erinnern, die die automatische Meldung eines Feuers im linken Triebwerk bestätigten – obwohl de facto keine solchen Hinweise vorhanden waren. Zudem war die Anzahl der erinnerten zusätzlichen Hinweise positiv korreliert mit der Geschwindigkeit, mit der das Triebwerk (fälschlicher Weise) abgeschaltet wurde. Dieses Ergebnis weist darauf hin, dass bei der Entstehung von *automation bias* und speziell des *commission* Fehlers tatsächlich auch Gedächtnisfehler oder Wahrnehmungsverzerrungen beteiligt sein könnten (vgl. Kap. 1.5.1).

1.5.3 Methodische Probleme der bisherigen Studien zu *automation bias*

Insgesamt ist festzuhalten, dass bislang vergleichsweise wenige experimentelle Studien zum Phänomen des *automation bias* existieren und damit auch wenige Erkenntnisse bezüglich beeinflussender Faktoren vorliegen. Zudem wurde in allen Untersuchungen die Luftfahrt als Anwendungsbereich gewählt, was im Übrigen auch ausnahmslos für die Untersuchungen zu *complacency* gilt. Hinzu kommt, dass bei beiden Phänomenen jeweils meist dieselben Versuchsumgebungen eingesetzt wurden; bei *complacency* v. a. die MAT-Batterie, für *automation bias* W/PANES oder ACFS. Inwiefern vergleichbare Auftretenshäufigkeiten auch mit grundlegend anders gestalteten Versuchsumgebungen gefunden werden können, und ob die Befunde auch auf andere Anwendungsdomänen übertragbar sind, ist offen. Um die oft zitierte Relevanz der Phänomene belegen zu können, wäre es jedoch essentiell, zunächst deren Unabhängigkeit von Versuchsumgebung und Anwendungsdomäne zu zeigen.

Neben dieser eher allgemeinen Schwäche der bisherigen Forschung sind jedoch auch Defizite auf Detailebene der einzelnen Studien festzustellen: Ähnlich der Kritik an *complacency*-Studien (vgl. Kap. 1.4.4) ist auch mit Blick auf das Phänomen *automation bias* ein Hauptkritikpunkt die Operationalisierung, und zwar insbesondere des *commission* Fehlers. Definitionsgemäß zeigt sich dieser darin, dass der Proband einer falschen Direktive der Automation folgt und zwar entweder, weil widersprechende Informationen anderer Quellen abgewertet und nicht berücksichtigt werden oder, weil die Direktive gar nicht erst anhand anderer verfügbarer Informationen überprüft wird. In der Definition wird also explizit darauf verwiesen, unter welchen Voraussetzungen das Befolgen falscher Direktiven einer Automation als *commission* Fehler zu werten ist. Auf empirischer Ebene wurden diese Voraussetzungen jedoch kaum adressiert und sich stattdessen auf die Erfassung fälschlich befolgter Automationshinweise beschränkt. Die einzige *automation bias* Studie, in der zusätzlich zu möglichen *commission* Fehlern auch das Überprüfungsverhalten der Pro-

banden erfasst wurde, ist die Studie von Skitka, Mosier und Burdick (2000; vgl. Kap. 1.5.2.1). Allerdings wählten die Autoren ein sehr vereinfachtes und letztlich künstliches Verifikationsparadigma. Die Probanden mussten auf maximal drei Zellen einer Matrix klicken, bis letztendlich eine mit der Empfehlung des Assistenzsystems identische oder eben abweichende Anweisung erschien. Warum aber die Matrix zuverlässigere Empfehlungen als das Assistenzsystem liefern sollte, dürfte für die Probanden kaum verständlich gewesen sein, zumal die letztlich gelieferte Information strukturell identisch gestaltet war (press button X) und auch nicht ersichtlich war, ob für diese Entscheidung andere, unabhängige Informationsquellen herangezogen worden waren. Zudem berichten Skitka, Mosier und Burdick (2000) keine Datenanalyse, in der das Verifikationsverhalten und mögliche *commission* Fehler zueinander in Bezug gesetzt worden wären. Gerade das wäre aber wichtig, um die Frage, warum *commission* Fehler auftreten, zu klären und ggf. zwischen den beiden verschiedenen Ausprägungen desselben mit Blick auf die Voraussetzungen unterscheiden zu können.

1.6 INTEGRATION DER KONZEPTE AUTOMATION BIAS UND COMPLACENCY

In den vorausgehenden Kapiteln wurden die beiden Konzepte *automation bias* und *complacency* getrennt voneinander vorgestellt und mögliche beeinflussende Faktoren auf Basis der bisherigen empirischen Forschung berichtet. In diesem Kapitel soll versucht werden, die Phänomene in ein gemeinsames Rahmenmodell zu integrieren. Obwohl deren Untersuchung in zwei getrennten Forschungslinien (überwachende Kontrolle versus Umgang mit Assistenzsystemen) erfolgte und sich in den Studien kaum explizite Querbezüge auf das jeweils andere Phänomen finden, weisen sie doch in verschiedener Hinsicht Gemeinsamkeiten auf. So kann konzeptuell als Basis beider Phänomene ein übersteigertes Vertrauen in Automation gesehen werden. Dies ist zwar nur für *complacency* auch empirisch belegt worden (Bailey & Scerbo, 2007), doch auch für *automation bias* gibt es Hinweise darauf, dass das Vertrauen in die Funktions- und Leistungsfähigkeit der Automation das Auftreten von *commission* und *omission* Fehlern beeinflusst (Skitka et al., 1999, vgl. Kap. 1.5.2.3).

Betrachtet man die Definitionen der Konzepte, zeigt sich eine weitere Parallele: *Complacency* manifestiert sich in der mangelnden Überwachung eines automatisierten Systems. *Commission* Fehler, als Befolgen der fehlerhaften Direktiven eines Assistenzsystems, resultieren, weil widersprechende Informationen abgewertet werden oder aber aufgrund einer mangelnden Überprüfung der Automation. Wenn auch nicht explizit erwähnt, so ist hier doch ein direkter Bezug der beiden Phänomene erkennbar: *Complacency* im Sinne einer mangelnden Überwachung bzw., je nach Au-

tomationstyp, mangelnden Überprüfung eines automatisierten Systems kann als eine von zwei möglichen Ursachen für *commission* Fehler verstanden werden.

Auch mit Blick auf den *omission* Fehler ist ein ähnlicher Zusammenhang feststellbar: Die Tatsache, dass Fehler einer Automation übersehen werden, kann verschiedene Ursachen haben, wie etwa Vigilanzprobleme oder die Fokussierung auf eine andere konkurrierende Aufgabe, die momentan wichtiger scheint. Eine mögliche Ursache ist aber auch in einer mangelnden Überwachung des Systems zu sehen. Somit stellt *complacency* ebenfalls eine von mehreren möglichen Ursachen für *omission* Fehler dar.

Fraglich ist jedoch, ob eine solche Übertragung der Phänomene mit Blick auf die unterschiedlichen Anwendungsbereiche der überwachenden Kontrolle bzw. der Nutzung von Assistenzsystemen gerechtfertigt ist. Beim *omission* Fehler legen es die bisherigen Untersuchungen nahe, dass eine Übertragung ohne weiteres möglich ist: Sowohl *omission* Fehler als auch *complacency* wurden über die Detektionsrate von Automationsfehlern bei einer sehr ähnlichen Parameterüberwachungsaufgabe erfasst¹. Einzig der Automationsgrad variierte hier: In den *complacency*-Studien (MAT-Batterie) bestand die Automation in einer automatischen Zurücksetzung der Parameter bei Parameterabweichungen, wobei jedes Eingreifen über eine Warnlampe angezeigt wurde (z.B. Parasuraman et al., 1993, vgl. Kap. 1.4.2.1). In den *automation bias*-Studien (W/PANES) wies die Automation nur auf die Abweichung hin; die Zurücksetzung musste von den Probanden aber selbst vorgenommen werden (z.B. Skitka et al., 1999, vgl. Kap. 1.5.2.). Automationsfehler äußerten sich in beiden Fällen darin, dass das System eine Abweichung nicht korrigierte bzw. keine Korrekturmaßnahmen empfahl. Um solche Automationsfehler kompensieren zu können und damit *omission* Fehlern vorzubeugen, hätten die Probanden in beiden Fällen die Parameteranzeigen kontinuierlich überwachen müssen. Insofern handelt es sich beim *omission* Fehler letztlich um ein Überwachungsproblem, das bezüglich der zugrunde liegenden Automationsform unspezifisch ist.

Doch auch die umgekehrte Übertragung, von *complacency* auf den Kontext der Nutzung von Assistenzsystemen, wie sie für den oben skizzierten Zusammenhang zum *commission* Fehler vorgenommen wird, erfordert nur eine geringfügige Erweiterung des Konzepts. Vorauszusetzen ist dabei lediglich, dass *complacency* sich nicht nur im Verlassen auf das rechtzeitige Eingreifen bzw. die rechtzeitige Alarmierung durch eine Automation, sondern auch im Verlassen auf die Diagno-

¹ Wie schon in Kapitel 1.4.4 erläutert, wird diese Operationalisierung für *complacency* als nicht angemessen erachtet. Denn damit werden nur *omission* Fehler erfasst, diese können aber auch andere Ursachen haben und auch *complacency* muss sich nicht zwingend darin manifestieren. Auch wenn ein kausaler Zusammenhang angenommen wird, rechtfertigt dies nicht die Gleichsetzung der Phänomene auf operationaler Ebene.

se- und Entscheidungskompetenz einer Automation zeigen kann. Während *complacency* sich im ersten Fall in einer mangelnden Überwachung zeigt, besteht im zweiten Fall die Verhaltensmanifestation in einer mangelnden Überprüfung der vom System gelieferten Diagnose bzw. Entscheidung.

Der geschilderte Zusammenhang zwischen Anwendungsbereich, *complacency*, *commission* sowie *omission* Fehlern ist in Abbildung 5 dargestellt.

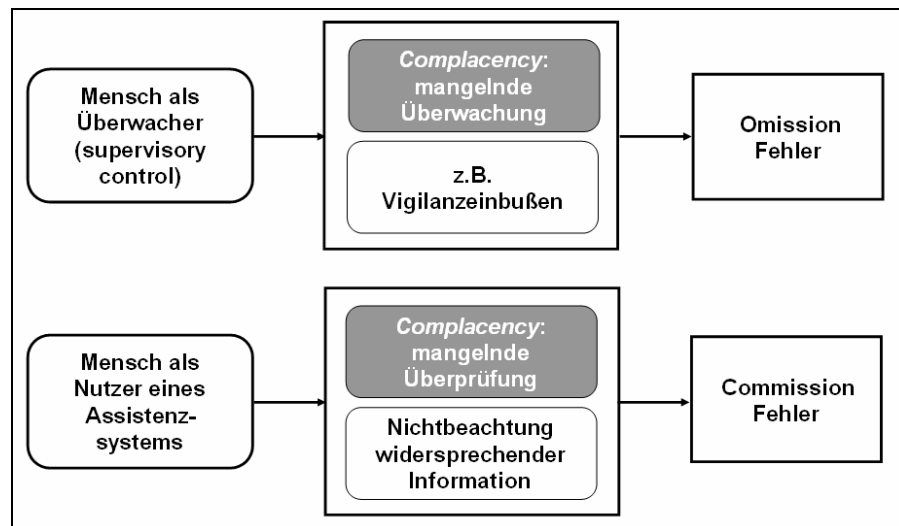


Abb. 5: Complacency im Kontext von Überwachung und Assistenzsystemnutzung

Begreift man ein übersteigertes Vertrauen in Automation als Basis beider Phänomene, liegt es nahe, dass jene Faktoren, die vertrauenssteigernd wirken, das Auftreten sowohl von *complacency* als auch von *automation bias* begünstigen. Folgt man zudem dem Vorschlag dieser Arbeit, *complacency* als mögliche Ursache für *commission* und *omission* Fehler zu sehen, so ist davon auszugehen, dass all jene Aspekte, die *complacency* fördern, auch das Risiko für die beiden Fehlertypen erhöhen.

Ein wesentlicher Unterschied der beiden Phänomene besteht jedoch darin, dass sowohl *commission* als auch *omission* Fehler im Gegensatz zu *complacency* das Vorliegen von Automationsfehlern (im Sinne von falschen Direktiven bzw. von Auslassungen) voraussetzen. Daraus ergeben sich unterschiedliche Implikationen mit Blick auf die Persistenz von *complacency* und *automation bias*: *Complacency* wird, solange die Automation fehlerfrei arbeitet, für den Operateur keine negativen Verhaltenskonsequenzen nach sich ziehen und somit im Sinne einer positiven Rückkopplung den Operateur darin bestärken, dass eine umfangreiche und genaue Überprüfung und Überwachung der Automation verzichtbar ist (vgl. Manzey & Bahner, 2005, vgl. Kap. 1.4.3). Zudem ist anzunehmen, dass Operateure die Reliabilität der Automation mit steigender Dauer eines fehlerfreien Betriebs zunehmend höher einschätzen, was eine vergleichsweise oberflächliche Überwachung und Überprüfung zusätzlich bestärkt. Demgegenüber werden *omission* oder *commission* Fehler in der Regel negative Konsequenzen nach sich ziehen und auch die vorausgegangenen Automati-

onsfehler nicht auf Dauer unbemerkt bleiben. Wie sich diese Erfahrung auf das künftige Verhalten des Operators auswirkt, wird dabei von den Merkmalen des Automationsfehlers abhängen. Wie in Kapitel 1.3.2 geschildert, ist davon auszugehen, dass Fehler sich vor allem dann, wenn sie nicht nachvollziehbar und vorhersehbar, schlecht kompensierbar sowie mit hohen Kosten verbunden sind, in einer Vertrauensminderung äußern. Sofern ein Fehler das Vertrauen in die Automation zu mindern vermag, ist auch eine Steigerung der Überwachungs- und Überprüfungintensität und eine Reduktion potenzieller *commission* und *omission* Fehler zu erwarten.

In Abbildung 6 wird versucht, den skizzierten Zusammenhang der Konzepte *complacency* und *automation bias* zu veranschaulichen.

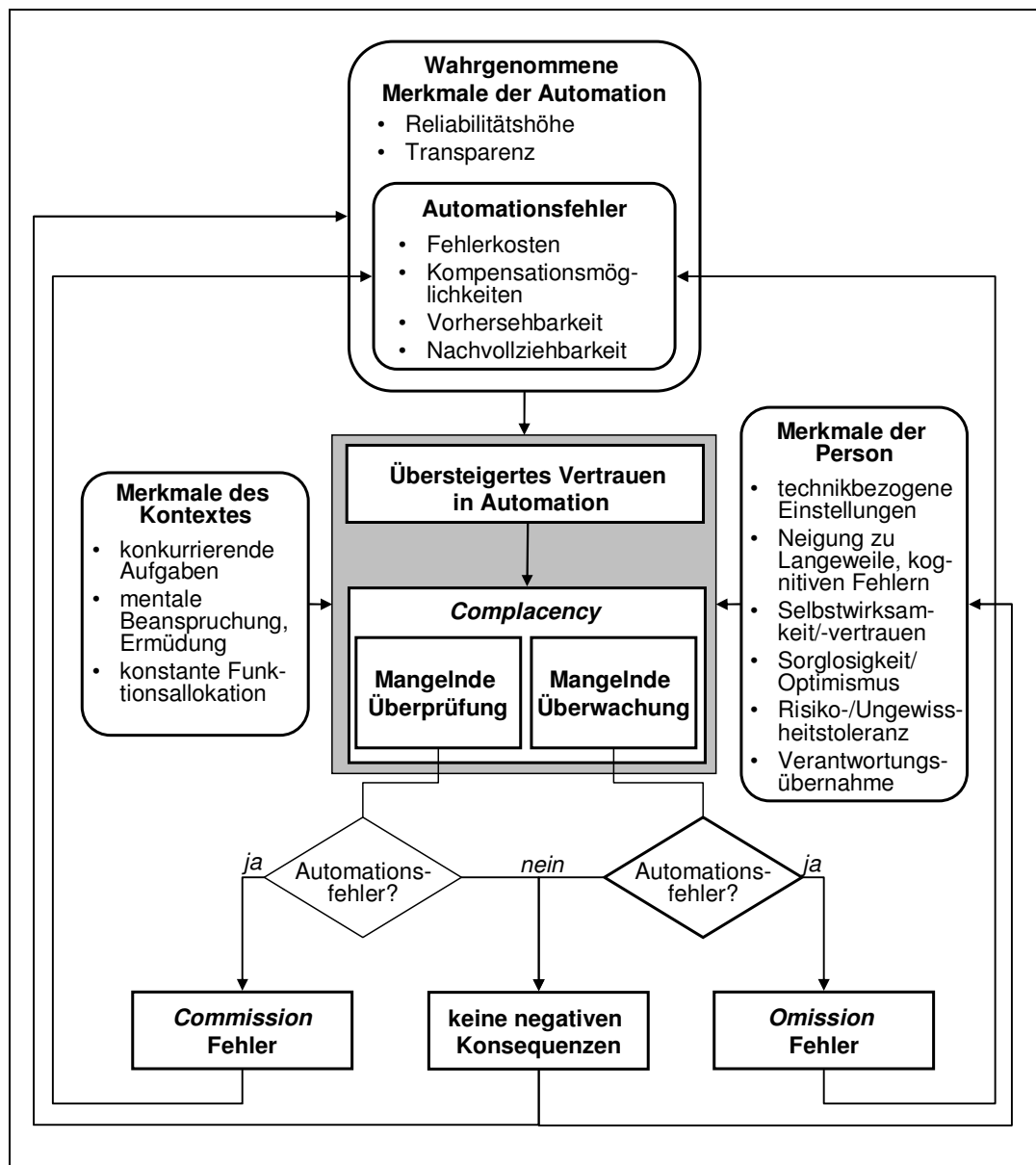


Abb. 6: Integration der Konzepte *complacency* und *automation bias*

Darin wird in Anlehnung an das Rahmenmodell von Manzey und Bahner (2005, vgl. Kap. 1.4.3) davon ausgegangen, dass ein übersteigertes Vertrauen in Automation sowie *complacency* maßgeblich von den Merkmalen der Person, des Kontextes und den wahrgenommenen Merkmalen der Automation determiniert werden. Diese wurden im Modell gemäß dem gegenwärtigen Forschungsstand angepasst. So wurde zwar die Reliabilitätshöhe, aufgrund der heterogenen Befundlage aber nicht die Konstanz der Reliabilität als beeinflussendes Automationsmerkmal aufgenommen. Auch fanden bisherige Befunde von *automation bias*-Studien Berücksichtigung. Allerdings hat sich hier bislang nur der Einfluss von subjektiver „Verantwortungsübernahme“ (vgl. Kap. 1.5.2.1) als empirisch belastbar erwiesen; bezüglich der anderen untersuchten Einflussfaktoren stellt sich die Befundlage zu widersprüchlich dar, um eine Aufnahme in das Modell zu rechtfertigen. Zudem wurde auf einige Details des Rahmenmodells von Manzey und Bahner (2005) zu Gunsten der Klarheit der Darstellung verzichtet, etwa auf die Differenzierung zwischen *complacency* als Verhaltenstendenz und *complacency* als manifestes Verhalten.

Ein weiterer Unterschied zum Rahmenmodell von Manzey und Bahner (2005) ist bezüglich der Automationsmerkmale festzustellen: Aufgrund der besonderen Bedeutung von Automationsfehlern für die Entstehung eines übersteigerten Vertrauens und *commission* sowie *omission* Fehlern, wurden deren Merkmale in einer eigenen Unterkategorie allgemein wahrgenommener Automationsmerkmale zusammengefasst. Treten Automationsfehler vermehrt auf und sind schlecht kompensierbar, vorhersehbar und nachvollziehbar sowie mit hohen Kosten verbunden, ist die Wahrscheinlichkeit, dass es zu einem übersteigerten Vertrauen kommt, vergleichsweise gering. Im Umkehrschluss ist davon auszugehen, dass eine fehlerfrei arbeitende Automation die Entstehung eines übersteigerten Vertrauens und in der Folge von *complacency* im Sinne einer mangelnden Überprüfung oder Überwachung begünstigt. Kommt es dann unerwartet zu Automationsfehlern, können *commission* Fehler die Folge einer mangelnden Überprüfung sein sowie *omission* Fehler aus einer mangelnden Überwachung resultieren. Diese Fehlererfahrung wirkt dann wiederum zurück, vermittelt über die wahrgenommenen Merkmale der Automation, auf das dem System entgegengebrachte Vertrauen, wobei eine Minderung desselben zu erwarten ist. Arbeitet das System dagegen weiterhin fehlerfrei, schließt sich der Kreis im Sinne einer positiven Rückkopplung: Je länger das System zuverlässig arbeitet, desto höher dürfte zum einen die Reliabilität desselben wahrgenommen werden und zum anderen dürfte dadurch auch ein Beitrag zur Aufrechterhaltung bzw. Ausbildung bestimmter Personmerkmale geleistet werden (z.B. positive technikbezogene Einstellungen, Sorglosigkeit/Optimismus). Diese beiden Aspekte fördern wiederum ein übersteigertes Vertrauen in Automation und führen somit zur Aufrechterhaltung einer mangelnden Überwachung bzw. Überprüfung des Systems.

1.7 ZIEL- UND FRAGESTELLUNG DER ARBEIT

Zentrale Ziele der vorliegenden Arbeit sind es, zu einem besseren Verständnis von Problemen im Umgang mit Assistenzsystemen beizutragen und mögliche Ansatzpunkte zu deren Vermeidung zu untersuchen. Im Fokus stehen dabei die Phänomene *complacency* und *automation bias*, die bisher immer nur getrennt voneinander betrachtet wurden. Dass diese beiden Phänomene jedoch, obwohl sie aus unterschiedlichen Kontexten stammen, einen engen konzeptuellen Bezug aufweisen, ist im vorhergehenden Kapitel (1.6) dargelegt worden. Das dort vorgestellte integrative Modell (Abb. 6) der Konzepte soll als heuristischer Rahmen für die vorliegende Arbeit zugrunde gelegt werden.

Wie in den vorangegangenen Kapiteln beschrieben (vgl. insbesondere Kap. 1.4.4 sowie 1.5.3), ist die bisherige empirische Basis zu den beiden Phänomenen begrenzt und zudem durch einige methodische Schwächen gekennzeichnet, die in erster Linie die Operationalisierung betreffen: Bezogen auf *complacency* sind dies insbesondere, dass auf eine Erfassung des Vertrauens in Automation verzichtet wurde und dass *complacency* nicht direkt sondern meist über mögliche Folgen operationalisiert wurde, zumindest aber kein normatives Modell „optimalen“ Informations-suchverhaltens herangezogen wurde. Mit Blick auf das Phänomen des *automation bias* liegt eine Begrenzung bisheriger Studien darin, dass *commission* Fehler bislang nie mit dem vorausgehenden Überprüfungsverhalten in Bezug gesetzt wurden. Zudem wurden beide Phänomene bisher nur im Kontext der Luftfahrt als Anwendungsbereich und fast immer in denselben, vergleichsweise künstlichen, Versuchsumgebungen untersucht. Obwohl sich in den letzten 15 Jahren eine Vielzahl empirischer Studien den Phänomenen gewidmet hat, fehlt bis heute ein letztendlich überzeugendes experimentelles Paradigma.

Daraus ergibt sich als Grundvoraussetzung der hier vorgestellten Arbeit, einen experimentellen Rahmen zu entwickeln, der die aufgeführten Probleme überwindet. Insbesondere gilt es, Operationalisierungen zu finden, die auch auf empirischer Ebene eine klare Trennung von *complacency* und möglichen Folgen realisieren. Nur so kann letztendlich der Kritik von Dekker und Hollnagel (2004, vgl. Kap. 1.4.1) begegnet und *complacency* des Status eines vorwissenschaftlichen „Alltagskonzepts“ enthoben und in ein wissenschaftlich brauchbares Konzept transferiert werden. Dies stellt die Grundlage dar, um die Rolle von *complacency* für mögliche Folgeeffekte im Sinne von *automation bias* und potenziellen Fertigkeitsverlusten untersuchen zu können.

Im Folgenden werden die Ziele und Fragestellungen der Arbeit sowie in groben Zügen der jeweilige Untersuchungsansatz im Detail vorgestellt.

(1) Übertragung auf einen neuen und gemeinsamen Anwendungsbereich

Complacency und *automation bias* sollen auf den Anwendungsbereich der Prozesssteuerung übertragen und in einer im Vergleich zu bisherigen Studien realitätsnäheren Versuchsumgebung untersucht werden. In dieser Versuchsumgebung werden bestimmte Parameter automatisch geregelt, wobei verschiedene Fehlfunktionen auftreten können. Bei der Detektion, Diagnose und Behebung dieser Fehler werden die Probanden von einem Entscheidungsassistenzsystem unterstützt, das somit sowohl eine Alarm- als auch eine Diagnosefunktion übernimmt. Mit Blick auf diese beiden Funktionen des Assistenzsystems erfolgt die Erfassung von *complacency* und *automation bias*.

Beantwortet werden soll so die Frage, ob die Phänomene tatsächlich, wie oft behauptet aber bislang empirisch nicht nachgewiesen, auch außerhalb des Luftfahrtkontextes und eng umschriebener Versuchskonstellationen Relevanz besitzen. Zum anderen soll geprüft werden, ob sich *complacency*, wie im vorangehenden Kapitel (1.6) beschrieben, auch auf den Kontext von Assistenzsystemen übertragen lässt.

(2) Operationalisierung von *complacency* unabhängig von möglichen Verhaltenskonsequenzen

Unter Berücksichtigung der Kritik an bisherigen Studien soll *complacency* direkt über das Informationssuchverhalten unabhängig von möglichen Folgen, wie etwa dem Übersehen von Automationsfehlern operationalisiert werden. Wie dem Rahmenmodell der Konzepte *complacency* und *automation bias* (Kap. 1.6, Abb. 6) zu entnehmen ist, kann sich *complacency* in einer mangelnden Überwachung und einer mangelnden Überprüfung äußern. Eine mangelnde Überwachung ist etwa dann gegeben, wenn ein Operateur sich auf die Alarmfunktion eines Assistenzsystems völlig verlässt. Zeigen würde sich letzteres z.B. darin, dass der Operateur – sofern kein Alarm vorliegt – keine Informationen über den tatsächlichen Systemzustand abrufen. Eine mangelnde Überprüfung bezieht sich dagegen auf Situationen, in denen ein eben solches Assistenzsystem einen kritischen Zustand indiziert und eine Fehlerdiagnose mit entsprechenden Behebungsschritten vorschlägt; folgt der Operateur den Direktiven, ohne sie vorher anhand verfügbarer Informationen angemessen zu verifizieren, wäre dies als mangelnde Überprüfung zu werten. In der vorliegenden Arbeit sollen beide Formen von *complacency* einer Untersuchung zugänglich gemacht werden. Mit Blick auf den Aspekt der Überprüfung soll für jede Fehlerdiagnose vorab definiert werden, welche Informationen idealerweise abgerufen werden müssten. Damit wird es möglich, das Verhalten der Probanden mit einem Optimum (normatives Modell sensu Moray, 2003) zu vergleichen und *complacency* entsprechend erst dann zu attestieren, wenn eine mangelnde Überprüfung vorliegt.

Zudem soll die Abweichung vom Optimum quantifiziert werden, um differenzierte Aussagen über den Ausprägungsgrad von *complacency* treffen zu können.

Mit Blick auf den Überwachungsaspekt soll das Informationssuchverhalten in vom Assistenzsystem als fehlerfrei indizierten Phasen analysiert werden. Allerdings ist es hier mit den verfügbaren Mitteln nicht möglich, ein normatives Modell optimalen Informationssuchverhaltens festzuschreiben. Geschuldet ist dies der generellen Schwierigkeit, bei kontinuierlichen Überwachungsaufgaben optimale Abruffrequenzen zu definieren. Entsprechend kann hier nicht absolut über eine mangelnde Überwachung geurteilt werden, sondern das Informationssuchverhalten nur im Vergleich zwischen Probandengruppen als geringer oder stärker ausgeprägt beurteilt werden. Insofern handelt es sich dabei nicht um eine direkte Operationalisierung von *complacency* im Sinne einer mangelnden Überwachung, wenngleich der Umfang des Informationssuchverhaltens damit in engem Zusammenhang stehen dürfte. Dem Rechnung tragend, soll im weiteren Verlauf für diese Variable nicht der Begriff „*complacency*“ verwendet werden, sondern allgemein vom „Informationssuchverhalten in fehlerfreien Phasen“ die Rede sein.

Darüber hinaus soll das Vertrauen der Probanden in die Automation erfasst werden, um so Aussagen darüber treffen zu können, inwiefern eine mangelnde Überprüfung bzw. vergleichsweise geringe Überwachung tatsächlich mit einem erhöhten Vertrauen in Automation in Verbindung steht. Mit Blick auf die Operationalisierung des Vertrauens in Automation stellt sich dabei die in Kapitel 1.3.1 beschriebene Schwierigkeit, dass bislang kein geteiltes Begriffsverständnis auszumachen ist und auch die Frage der zugrunde liegenden Vertrauensdimensionen nicht abschließend geklärt ist. Dies wäre jedoch die Voraussetzung, um ein entsprechendes Erhebungsinstrument theoretisch begründbar für die vorliegende Studie zu entwickeln. Vor diesem Hintergrund scheint es am angemessensten, Probanden direkt nach ihrem Vertrauen in die Automation (bzw. das Assistenzsystem) zu fragen. Mithilfe einer solchen einfachen Abfrage konnten bereits in einer der ersten Studien zum Vertrauen in Automation Unterschiede zwischen den Probanden abgebildet werden (Lee & Moray, 1992; verwendet wurde dort das Item: „Overall, how much do you trust the system?“). Zusätzlich soll explorativ auch das Selbstvertrauen in die manuelle Fehlerdiagnose und -behebung erhoben werden, da anzunehmen ist, dass das Verhalten der Operateure gegenüber der Automation auch von diesem beeinflusst wird (vgl. Kap. 1.3.2).

Davon ausgehend, dass sich die beschriebenen Operationalisierungsansätze in der praktischen Umsetzung als tauglich erweisen, wird damit Folgendes erreicht: Erstens kann das Verhalten der Probanden gegenüber zwei Aspekten des Assistenzsystems – seiner Alarm- sowie Diagnosefunktion – getrennt betrachtet und analysiert werden. Zweitens wird *complacency* über die Überprüfung der automatisch generierten Diagnosen erstmals unabhängig von möglichen Folgen

operationalisiert und im Vergleich zu einer „optimalen“ Überprüfung quantifiziert. Sofern eine „mangelnde Überprüfung“ beobachtbar ist, wäre dies ein erster direkter empirischer Nachweis von *complacency*. Drittens wird über die Analyse des Vertrauens in Automation auch der zwar theoretisch implizierte, bislang aber kaum empirisch belegte Bezug zur subjektiven Komponente von *complacency* zugänglich gemacht.

(3) Analyse möglicher Maßnahmen zur Vorbeugung von *complacency*

Wie in den vorangehenden Kapiteln beschrieben, gibt es bislang wenig gesicherte Kenntnisse über wirksame Maßnahmen, mit denen *complacency* oder auch *automation bias* vorgebeugt werden kann. Allerdings geht aus dem in Kapitel 1.6 beschriebenen Integrationsversuch der Konzepte hervor, dass der wahrgenommenen Reliabilität bzw. auftretenden Automationsfehlern für deren Entstehung und Aufrechterhaltung – und letztlich auch im Falle von *omission* und *commission* Fehlern als Manifestationsvoraussetzung – eine zentrale Rolle zukommt. Dabei ist davon auszugehen, dass die Erfahrung von Automationsfehlern zu einer Minderung des Vertrauens in die Automation führt und somit auch *complacency* und *automation bias* entgegenwirken kann. Folglich könnte die gezielte Einstreuung von Automationsfehlern während des Trainings eine hilfreiche Gegenmaßnahme darstellen. Während anzunehmen ist, dass über die Erfahrung von Automationsfehlern im Training *complacency*-Effekte reduziert werden können, ist zu vermuten, dass die reine Information darüber, dass Fehler auftreten können, dies nicht zu leisten vermag. So konnten dadurch etwa *automation bias*-Effekte nicht verhindert werden (vgl. Kap. 1.5.2.2). Basierend auf diesem Rational sollen in der vorliegenden Arbeit zwei experimentelle Gruppen miteinander verglichen werden: Eine Gruppe, die während des Trainings Automationsfehler erfährt, und eine Gruppe, die nur über die Eventualität von Automationsfehlern informiert wird.

Aus dem Rahmenmodell (Kap. 1.6, Abb. 6) sind einige weitere Aspekte ableitbar, die bei der experimentellen Umsetzung Berücksichtigung finden sollen: Bezüglich der Gestaltung der Automationsfehler ist darauf zu achten, dass diese für die Probanden nicht vorhersehbar auftreten, mit Kosten verbunden, schlecht kompensierbar und wenig nachvollziehbar („einfache Fehler“) sind, da insbesondere dann vertrauensmindernde Effekte zu erwarten sind. In einem ersten Experiment soll dies über Fehldiagnosen eines Assistenzsystems realisiert werden, während in einem zweiten Experiment der Einfluss von Systemausfällen untersucht werden soll. Bei der Gestaltung der Versuchssituation soll eine hinreichende mentale Beanspruchung durch konkurrierende Aufgaben sichergestellt, eine konstante Funktionsallokation gewählt und die Automation im Allgemeinen transparent gestaltet werden.

Sollten trotz der genannten *complacency*-begünstigenden Rahmenbedingungen durch die Erfahrung von Automationsfehlern *complacency* und *automation bias* reduziert werden können, bestün-

de in einer solchen Trainingsintervention ein für die Praxis wertvoller Ansatzpunkt, um diesen Automationsproblemen vorzubeugen.

(4) Analyse möglicher Folgen von *complacency*

4a) *Commission* Fehler

Es wird angenommen, dass *commission* Fehler eine mögliche Folge von *complacency* gegenüber der Diagnosefunktion darstellen. In der vorliegenden Arbeit soll dies empirisch geprüft werden. Hierzu soll gegen Ende der beiden Experimente vom Assistenzsystem eine Fehldiagnose generiert werden. Sofern die Probanden dieser fälschlicherweise folgen, ist dies als *commission* Fehler zu werten. Im Unterschied zu bisherigen Studien soll zudem bei auftretenden *commission* Fehlern eine detaillierte Analyse des vorausgehenden Überprüfungsverhaltens erfolgen.

Sofern *commission* Fehler beobachtbar sind, könnte damit die Frage beantwortet werden, welche Informationssuchmuster hinter einem *commission* Fehler stehen und inwiefern *complacency* ursächlich für einen solchen Fehler sein kann. Zudem soll geprüft werden, inwiefern *commission* Fehler durch die Erfahrung von Automationsfehlern im Training reduziert werden können.

4b) *Omission* Fehler

Auch mit Blick auf *omission* Fehler wird angenommen, dass diese Folge von *complacency* gegenüber der Alarmfunktion bzw. einer verminderten Überwachung des Systems sein können. Auch das soll in der vorliegenden Arbeit empirisch geprüft werden. Um *omission* Fehler erfassen zu können, soll gegen Ende des Experiments ein Systemausfall simuliert werden. Dieser soll sich so äußern, dass das System auftretende Fehler der zugrunde liegenden Prozesssteuerung nicht meldet (*automation miss* bzw. Auslassungsfehler), sodass die Probanden vorliegende Fehler ohne Automationsunterstützung entdecken, diagnostizieren und beheben müssen. Sofern ein solcher Fehler nicht entdeckt wird, ist dies als *omission* Fehler zu werten. Geprüft werden soll, inwiefern *omission* Fehler mit dem vorausgehenden Überwachungsverhalten verbunden sind. Dieses soll über das Informationssuchverhalten in fehlerfreien Phasen operationalisiert werden, wobei Aussagen über die Ausprägung des Überwachungsverhaltens nur relativ im Vergleich der Probanden untereinander getroffen werden können.

Sofern *omission* Fehler auftreten, würde diese Analyse Aufschluss über deren Beeinflussung durch den Umfang der Automationsüberwachung geben. Zudem soll die Frage beantwortet werden, ob *omission* Fehler durch die Erfahrung von Automationsfehlern im Training reduziert werden können. *Omission* Fehler werden ausschließlich im zweiten Experiment betrachtet.

4c) Fertigkeitsverlust

In Erweiterung der Kernfragestellung sollen als eine weitere mögliche Folge von *complacency* potenzielle Fertigkeitsverluste empirisch zugänglich gemacht werden. Dies soll dadurch geschehen, dass die gewohnte Unterstützung durch das Assistenzsystem nach einer Weile, für die Probanden unerwartet, ausfällt und eine rein manuelle Fehlerdiagnose und -behebung erforderlich wird. Ein möglicher Fertigkeitsverlust läge dann bei einer Leistungsver schlechterung im Vergleich zum – ebenfalls zunächst manuellen – Training vor.

Vermutet wird, dass *complacency* zu einem Verlust der manuellen Fehlerdiagnose-Kompetenz beitragen kann. Dies liegt nahe, da durch den Verzicht auf eine Überprüfung von automatisch generierten Fehlerdiagnosen auch die eigene Fehlerdiagnosekompetenz nicht mehr trainiert wird. Wenn die Fehlererfahrung im Training zu einer Reduktion von *complacency* führt, müsste damit zudem auch eine Verminderung potenzieller Fertigkeitsverluste verbunden sein. Dies zu prüfen ist ebenfalls Ziel der vorliegenden Arbeit.

Zur Beantwortung der beschriebenen Fragestellungen wurden zwei experimentelle Studien durchgeführt. In beiden wurde jeweils im Training die Erfahrung von Automationsfehlern versus die Information über mögliche Automationsfehler variiert. In der ersten Studie bestanden Automationsfehler darin, dass ein Assistenzsystem falsche Fehlerdiagnosen stellte. In der zweiten Studie wurde dagegen der Einfluss von Ausfällen des Assistenzsystems untersucht.

Bevor jedoch auf die Studien im Einzelnen eingegangen wird, soll im nachfolgenden Kapitel die Versuchsumgebung „AutoCAMS“ vorgestellt werden, die in beiden Experimenten Einsatz fand.

2 EXPERIMENTALUMGEBUNG AUTOCAMS

2.1 ALLGEMEINE BESCHREIBUNG DES SYSTEMS

Als Experimentalumgebung beider im Rahmen der vorliegenden Arbeit durchgeführten Studien diente eine Modifikation der Prozesssteuerungs-Mikrowelt AutoCAMS (Lorenz, Di Nocera, Röttger & Parasuraman, 2002). Das System basiert auf dem von Hockey, Wastell und Sauer (1998; Sauer, Hockey & Wastell, 2000) entwickelten „Cabin Air Management System“ (CAMS), ist aber im Unterschied zu diesem mit einem Assistenzsystem zur Unterstützung von Fehlerdiagnose und -management ausgestattet.

AutoCAMS simuliert das Lebenserhaltungssystem einer Raumstation und besteht aus den fünf Subsystemen Sauerstoff, Druck, Kohlenstoffdioxid, Temperatur und Luftfeuchtigkeit. Um atmosphärische Bedingungen auf der Raumstation sicherzustellen, müssen diese innerhalb ihres jeweiligen Sollbereichs gehalten werden, wobei dies in der Regel autonom durch das System erfolgt. Allerdings ist es möglich, dass aufgrund von Fehlfunktionen Sollwertabweichungen der Parameter auftreten. Die Aufgabe des Operators besteht in der überwachenden Kontrolle der Parameter und, sofern eine Fehlfunktion im System vorliegt, in deren Diagnose und Behebung. Fehlerdetektion, -diagnose und -management werden jedoch durch ein Assistenzsystem (*Automated Fault Identification and Recovery Aid*, AFIRA) unterstützt. Sobald eine Fehlfunktion im System erkennbar ist, wird der Operator darauf hingewiesen, indem ein fehlerunspezifischer Masteralarm erscheint und zeitgleich eine Fehlerdiagnose und eine Sequenz von Handlungsschritten, die zur Behebung der Fehlerfunktion erforderlich sind, von AFIRA vorgeschlagen werden. Die Realisierung der Fehlerbehebung erfolgt dann jedoch manuell durch den Operator. AFIRA besteht also aus zwei funktionellen Komponenten, einem fehlerunspezifischen Alarmsystem sowie einem fehlerspezifischen Entscheidungsassistenzsystem. Letzteres ist dabei mit Sheridan und Verplank (1978) als „level 4“ Automation einzuordnen, bei der das System eine Alternative für die Entscheidungen vorschlägt, deren letztendliche Auswahl und Umsetzung jedoch dem Menschen obliegt.

Neben der Überwachung des Systems und der Behebung von Fehlern sind vom Operator zwei zusätzliche Aufgaben zu erledigen: Erstens wird durchschnittlich einmal pro Minute ein Verbindungssymbol eingeblendet, das möglichst schnell per Mausklick quittiert werden muss, um die Verbindung mit der Raumstation zu bestätigen (Reaktionszeitaufgabe). Zweitens ist zu jeder vollen Minute der CO₂-Wert abzulesen, in ein dafür vorgesehenes Feld einzutragen und zu spei-

chern (prospektive Gedächtnisaufgabe). Diese beiden Sekundäraufgaben dienen der Realisierung einer Mehrfachaufgabenbedingung, da empirisch belegt ist, dass es nur unter dieser Voraussetzung zu *complacency*-Effekten kommt (Parasuraman et al., 1993).

2.2 BENUTZEROBERFLÄCHE UND STEUERUNG VON AUTO-CAMS

Alle Bedien- und Anzeigeelemente, die der Steuerung und Überwachung von AutoCAMS dienen, werden auf einem Bildschirm dargeboten (siehe Abb. 7). Bedieneingaben erfolgen maus- bzw. tastaturgesteuert. Genutzt wurden in der Untersuchung Pentium 4 Rechner (2.80 GHz; 265 MB Arbeitsspeicher) mit 17 Zoll Monitoren. Die einzelnen Bildschirmbereiche werden nachfolgend oben links beginnend im Uhrzeigersinn im Detail beschrieben.

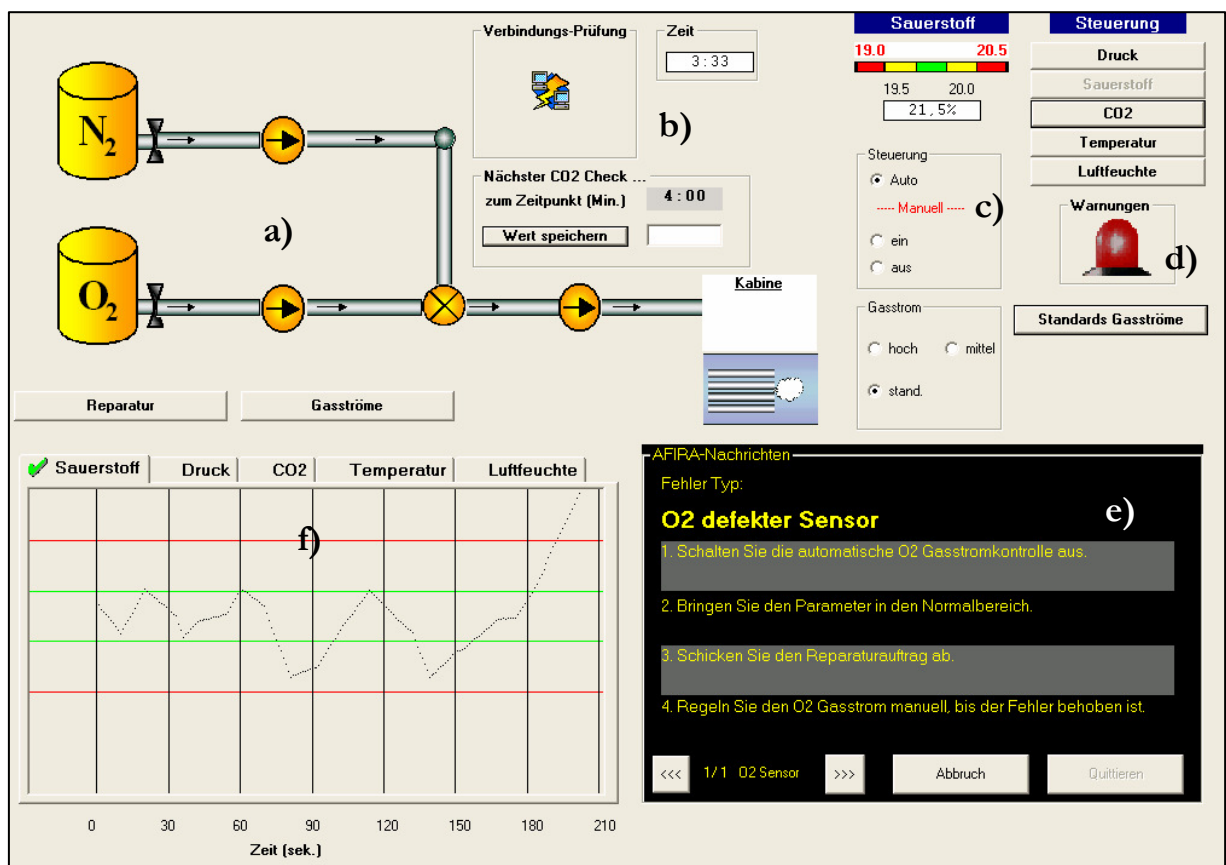


Abb. 7: Benutzeroberfläche von AutoCAMS: a) schematische Darstellung des Sauerstoff- und Stickstoffsystems, b) Bildschirmbereich für Sekundäraufgaben, c) manueller Steuerungsbereich, d) Masteralarm, e) Assistenzsystem AFIRA, f) Verlaufsanzeigenbereich

2.2.1 Schematische Darstellung des Sauerstoff- und Stickstoffsystems

Oben links auf Abbildung 7 (a) ist eine schematische Darstellung der Sauerstoff- und Stickstofftanks und deren Gasleitungen zur Kabine zu sehen. Ausgehend vom jeweiligen Tank strömt

jedes der beiden Gase zunächst durch ein zufuhrregelndes Ventil, wobei nach dem Ventil das erste Gasstrommessgerät schematisch visualisiert ist (Pfeilsymbol). Danach folgt das Mixerventil und zwischen diesem und der Kabine ein weiteres Gasstrommessgerät. Durch Klick auf den unterhalb der schematischen Darstellung befindlichen Button „Gasströme“ werden die Tankfüllstände sowie die durch das jeweilige Gasstrommessgerät erfasste Stromstärke angezeigt (siehe Abb. 8). Diese Informationen werden nach ihrem Abruf für 15 s angezeigt, können danach aber jederzeit wieder aufgerufen werden. Rechts unten in Abbildung 8 ist die Kabine visualisiert, in die Sauerstoff und Stickstoff entsprechend der dargestellten Gasstromstärken eingeleitet werden. Das weiße Feld unterhalb der Bezeichnung „Kabine“ dient der Rückmeldung zur Temperaturregelung: Hier wird gegebenenfalls „Heizung an“ bzw. „Kühler an“ eingeblendet. Im blauen Feld unterhalb davon indizieren weiße Wolken, dass das Ablassventil gegenwärtig geöffnet ist um den Kabinendruck zu reduzieren.

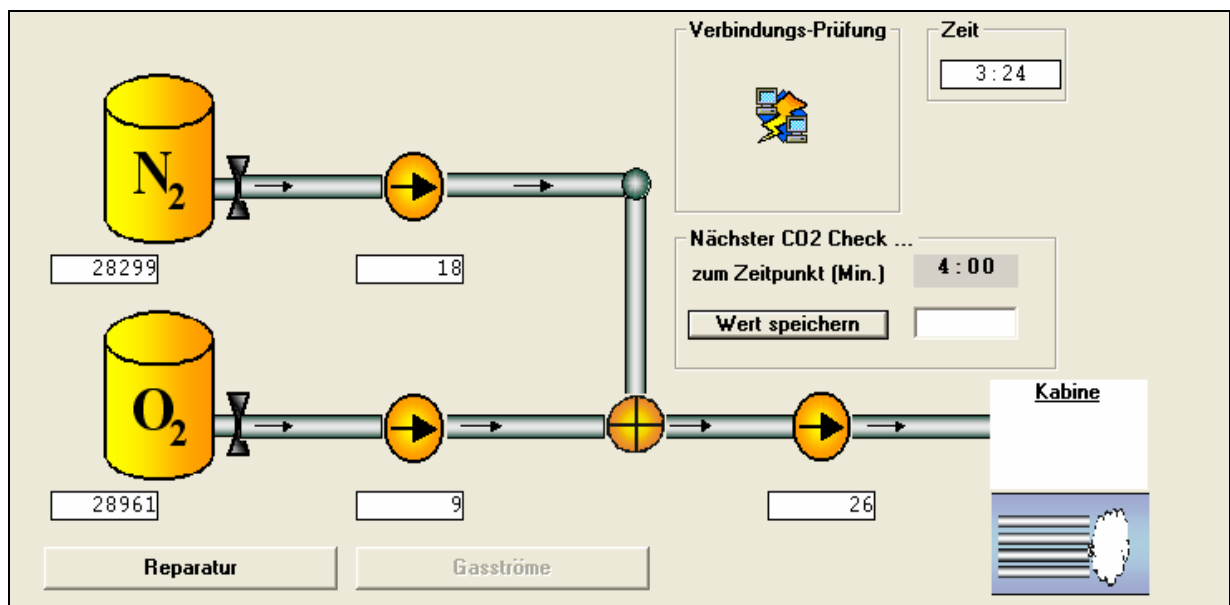


Abb. 8: Schematische Darstellung des Sauerstoff- und Stickstoffsystems

2.2.2 CO₂-Eingabefeld und Verbindungssymbol

Oberhalb der Kabinendarstellung findet sich das Eingabefeld für den CO₂-Wert, der nach der Eingabe per Mausklick auf den Button „Wert speichern“ jeweils zur vollen Minute gespeichert werden soll. Abzulesen ist dieser Wert über die Steuerungsmenüs (vgl. Kap. 2.2.3). Um den richtigen Zeitpunkt für die Eintragung zu ermitteln, steht den Probanden die Zeitanzeige rechts neben dem Verbindungssymbol (siehe Abb. 8) zur Verfügung. Diese zeigt durchgehend die seit Start des Programms vergangene Zeit sekundengenau an. Da die Probanden keinen direkten Hinweis auf den Ablauf einer vollen Minute erhalten und somit selbst für die Überwachung der Zeit zuständig sind, handelt es sich hierbei um eine prospektive Gedächtnisaufgabe.

Oberhalb des Eingabefelds für die CO₂-Werte wird in unregelmäßigen Abständen und durchschnittlich einmal pro Minute das Verbindungssymbol, wie auf Abbildung 8 ersichtlich, eingeblendet. Es verschwindet, sobald der Operateur es per Mausklick auf das Symbol quittiert (Reaktionszeitaufgabe).

2.2.3 Manuelle Steuerungsmenüs

Oben rechts auf der Bildschirmoberfläche (siehe Abb. 7) befindet sich der Bereich zur manuellen Steuerung der Parameter. Per Mausklick auf einen der Buttons („Druck“, „Sauerstoff“, „CO₂“, „Temperatur“, „Luftfeuchte“) erscheint links daneben das jeweilige Steuerungsmenü wie in Abbildung 9 für den Parameter Sauerstoff dargestellt.

Jedes dieser Menüs informiert im oberen Bereich über die zulässigen Soll- und Normalbereiche (in Abb. 9 Sollbereich: 19.0 – 20.5, Normalbereich: 19.5 – 20.0) und den aktuellen Messwert des jeweiligen Parameters in der Kabine (hier 19.7 %). Unterhalb davon werden die jeweils verfügbaren Steuerungsoptionen angezeigt. Im Normalfall werden alle Parameter automatisch geregelt. Der Operateur hat jedoch immer die Möglichkeit, von automatischer auf manuelle Steuerung umzuschalten, was beim Vorliegen einer Fehlfunktion zwingend erforderlich ist. Bezüglich der Parameter Stickstoff und Sauerstoff besteht die Möglichkeit, den Gasstrom ein- oder auszuschalten (Radiobuttons „ein“ bzw.

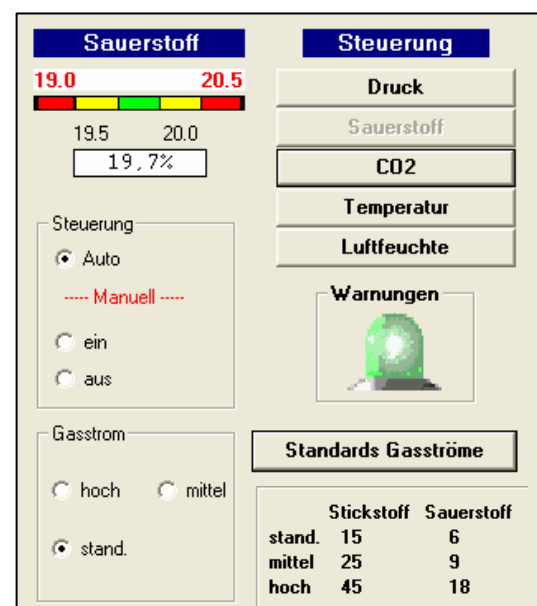


Abb. 9: Steuerungsmenü und Standards Gasströme

„aus“) sowie die Gasstromstärke zu variieren („hoch“, „stand.“ für „standard“ und „mittel“). Die systemseitigen Einstellungen für die verschiedenen Gasstromstärken können über den Button „Standards Gasströme“ aufgerufen werden. Nach Klick auf diesen Button werden die Werte für beide Parameter unterhalb des Buttons für 15 s eingeblendet. Änderungen der Gasstromstärke sind sowohl im manuellen wie auch im automatischen Betrieb möglich.

2.2.4 Masteralarm

Zwischen dem Bereich zur manuellen Steuerung und dem Button zum Aufruf der Standardeinstellungen für die Gasströme, befindet sich die fehlerunspezifische Masteralarm-Anzeige

(siehe Abb. 9). Ist diese in grüner Farbe dargestellt, bedeutet dies, dass gegenwärtig kein Fehler im System vorliegt. Sobald jedoch eine Fehlfunktion im System – ggf. auch durch den Operateur – potenziell erkennbar ist, wechselt die Farbe des Masteralarms von grün auf rot. Die Alarmanzeige bleibt solange rot, bis alle Fehler behoben sind und alle Parameterwerte wieder in den Sollbereich zurückgeführt wurden.

2.2.5 Assistenzsystem AFIRA

Der Bereich unten rechts der Bildschirmoberfläche ist dem Assistenzsystem AFIRA vorbehalten. In fehlerfreien Phasen ist nur eine schwarze Fläche dargestellt. Sobald eine Fehlfunktion vom System erkannt wird, erscheint hier, zeitgleich mit dem Masteralarm, eine Fehlermeldung, wie in Abbildung 7 beispielhaft für einen defekten Sauerstoffsensor dargestellt. Angegeben wird neben dem Fehlertyp (Diagnosevorschlag) auch die Handlungssequenz zur Behebung des Fehlers. Diese beinhaltet neben dem Hinweis, dass der Parameter bis zur Fehlerbehebung manuell zu steuern ist, immer auch die Aufforderung, einen entsprechenden Reparaturauftrag abzuschicken.

Um dies zu realisieren, kann per Mausklick auf den Button „Reparatur“ (links oberhalb der Verlaufsanzeige, vgl. Abb. 7) ein *pop-up*-Menü aufgerufen werden, das in Abbildung 10 dargestellt ist. Es enthält alle Fehlfunktionen, die im System auftreten können, wobei der Operateur die aktuell vorliegende Fehlfunktion per Radio-Button auswählen und deren Reparatur über den Button „Reparieren“ initiieren kann. Nachdem ein Reparaturauftrag initiiert wurde, ist es für die Dauer von 60 s nicht möglich, einen anderen Reparaturauftrag zu aktivieren. Dies wurde so realisiert, um falsche Fehlerdiagnosen mit

Abb. 10: Menü zur Auswahl und Aktivierung eines Reparaturauftrags

Kosten zu verbinden: Zum einen ist für die Dauer der Fehlfunktion die manuelle Steuerung des betreffenden Parameters erforderlich, was eine erhöhte Beanspruchung für den Operateur bedeutet. Zum anderen führt die manuelle Steuerung in der Regel zu einem erhöhten Ressourcen-

verbrauch, der je nach Fehlertyp (z.B. Leck) noch zusätzlich erhöht sein kann. Nach Ablauf von 60 s ist die Fehlerreparatur abgeschlossen und AFIRA zeigt je nach dem, ob der gewählte Reparaturauftrag dem tatsächlich vorliegenden Fehler entsprach, eine der beiden in Abbildung 11 dargestellten Rückmeldungen an.

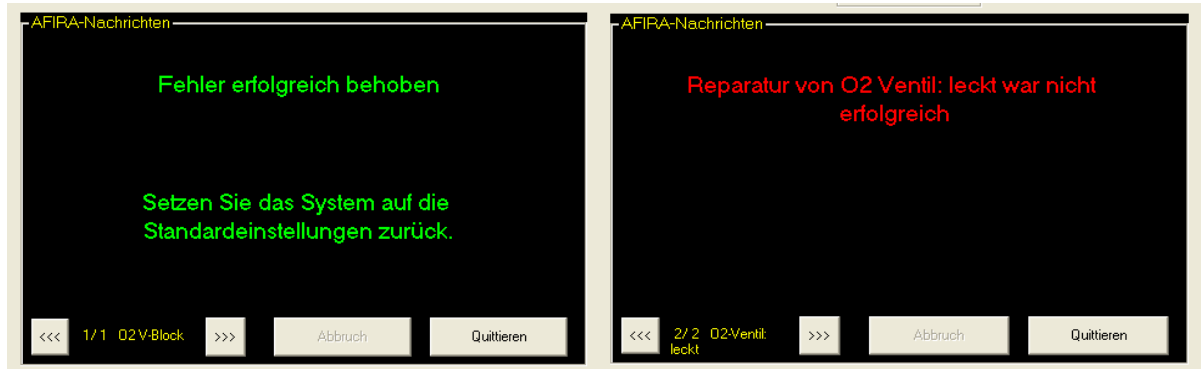


Abb. 11: AFIRA-Rückmeldungen nach einer erfolgreichen bzw. nicht erfolgreichen Fehlerreparatur

Im Falle einer nicht erfolgreichen Reparatur (rechts in Abb. 11) wird dabei als Fehlertyp immer der mittels Reparaturauftrag gewählte ausgegeben. Es ist möglich, dass zu einem Zeitpunkt mehrere Fehlfunktionen im System vorliegen, etwa weil ein früherer Fehler noch nicht erfolgreich repariert wurde, mittlerweile aber bereits ein weiterer Fehler aufgetreten ist. Um dem Rechnung zu tragen, ist AFIRA ähnlich einem Registerkartenprinzip aufgebaut. Unten links auf der Maske ist jeweils ersichtlich, welche Meldung in Abhängigkeit vom Fehlertyp gerade angezeigt wird und wie viele Meldungen gegenwärtig vorliegen (in Abb. 11 unten rechts bedeutet „2/2“, dass gerade die zweite von insgesamt zwei Fehlermeldungen dargestellt wird). Sofern mehr als eine Meldung vorliegt, kann der Operateur zwischen diesen über die Pfeiltasten hin- und herwechseln. Bei Eingang einer neuen Meldung (z.B. Reparaturrückmeldung oder Feststellung einer neuen Fehlfunktion) wird jedoch vom System automatisch die aktuelle Meldungskarte angezeigt. Bei jeder einzelnen Fehlermeldung hat der Operateur die Möglichkeit, vor erfolgter Reparatur die Unterstützung für diesen Fehler durch das Assistenzsystem abubrechen (Mausklick auf den Button „Abbruch“). Dies führt dazu, dass die Fehlermeldung nicht mehr angezeigt wird, allerdings erhält der Operateur in jedem Fall eine Rückmeldung über den Erfolg eines ggf. initiierten Reparaturauftrags. Mit Blick auf Reparaturrückmeldungen hat der Operateur die Möglichkeit, diese zu quittieren (Mausklick auf den Button „Quittieren“), was dann ebenso dazu führt, dass die Meldung nicht mehr angezeigt wird. War jedoch der abgeschickte Reparaturauftrag falsch, sodass die Fehlfunktion immer noch vorliegt, meldet AFIRA die Fehlfunktion erneut. Das bedeutet, dass im Fall einer von AFIRA generierten Fehldiagnose, der vom Operateur zunächst bezüglich des gewählten Reparaturauftrags gefolgt wird, zunächst eine negative Reparaturrückmeldung („nicht erfolgreich“) erfolgt, danach aber auch der tatsächlich vorliegende Fehler von AFIRA angezeigt wird.

Damit wird sichergestellt, dass Fehldiagnosen von AFIRA vom Operateur, wenn auch im Nachhinein, als solche erkannt werden. Festzuhalten ist, dass der Operateur den von AFIRA vorgeschlagenen Diagnosen und den dazu gehörenden Handlungsschritten zwar folgen kann, aber keinesfalls folgen muss. So könnte er etwa durch Inspektion der verfügbaren Systeminformationen (Gasströme, Verlaufsanzeigen) zu einer abweichenden Fehlerdiagnose gelangen und entsprechend andere Fehlerbehebungsschritte wählen.

2.2.6 Verlaufsanzeigen

Unten links in Abbildung 7 ist die Verlaufsanzeige der verschiedenen Parameter dargestellt, in diesem Fall für Sauerstoff. Angezeigt werden jeweils der Verlauf der in der Kabine gemessenen Parameterwerte für die letzten 240 s sowie der parameterspezifische Sollbereich (gekennzeichnet durch die roten Linien) und Normalbereich (gekennzeichnet durch die grünen Linien). Die Verlaufsanzeige eines Parameters wird durch Mausklick auf die entsprechende Registerkarte aufgerufen und bleibt dann für 15 s sichtbar. Nach Ablauf dieses Zeitraums wird das Verlaufs Fenster grau maskiert und muss ggf. erneut per Mausklick aufgerufen werden. Ein Wechsel zwischen den Parametern ist – ebenfalls per Mausklick – jederzeit möglich.

2.3 FEHLERTYPEN, -DIAGNOSE UND -MANAGEMENT

In der hier vorgestellten Studie traten insgesamt acht verschiedene Fehlfunktionen auf. Diese acht Fehlfunktionen lassen sich auf vier Fehlertypen reduzieren, die jeweils entweder im Stickstoff- oder aber im Sauerstoffsystem auftreten können:

Blockierungen des Sauerstoff- oder Stickstoffventils: Folge ist, dass weniger Gas als nötig in die Kabine strömt, wobei auch weniger Gas als normal den Tank verlässt.

Leck des Sauerstoff- oder Stickstoffventils: Auch hier wird weniger Gas als benötigt in die Kabine eingeleitet, die Abgabemenge aus dem Tank ist aber nicht vermindert.

Offen steckendes Sauerstoff- oder Stickstoffventil: In diesem Fall strömt permanent Gas in die Kabine, was im Falle eines offen steckenden Sauerstoffventils zu einem kontinuierlichen Anstieg der Konzentration führt. Bei einem offen steckenden Stickstoffventil wird durch die automatische Öffnung des Ablassventils die Konzentration in der Kabine zwar im oberen Bereich, jedoch noch innerhalb des Normalbereichs, gehalten.

Sauerstoff- oder Stickstoffsensordefekt: Bei einem Sensordefekt hängt das Symptommuster davon ab, ob der Defekt auftritt, während die jeweilige Gaskonzentration in der Kabine gerade ansteigt oder abnimmt. Im ansteigenden Fall führt der Sensordefekt zu einer immer weiter zunehmenden Gaskonzentration in der Kabine, da der Sensor nicht das Erreichen des oberen Normwerts rückmeldet. Sofern der Sauerstoffsensor betroffen ist, muss über eine etwas aufwändigere Prüfprozedur das Vorliegen eines offen steckenden Ventils ausgeschlossen werden, da dieses sich zunächst identisch äußert. Die Abgrenzung von offen steckendem Stickstoffventil und Stickstoffsensordefekt (bei steigenden Werten) ist dagegen einfach vorzunehmen: Bei einem Sensordefekt öffnet sich das Ablassventil nicht und die Konzentration steigt über den Normalbereich hinaus an. Bei einem offen steckenden Ventil öffnet sich jedoch das Ablassventil und die Konzentration bleibt innerhalb des Normalbereichs. Wenn ein Sensordefekt auftritt, während die Gaskonzentration in der Kabine gerade abnimmt, ist die Diagnose vergleichsweise einfach festzustellen: Obwohl die Konzentration in der Kabine immer weiter sinkt, verlässt kein Gas den Sauerstoff- bzw. Stickstofftank.

In Anhang A ist im Detail dargestellt, welche Prüfschritte erforderlich sind, um den jeweiligen Fehler sicher diagnostizieren zu können, und mit welchen Handlungsschritten dessen Behebung verbunden ist. Jede Fehlfunktion erfordert den Abruf mehrerer Informationen, um einen gegebenen Diagnosevorschlag verifizieren zu können (vgl. hierzu Kap. 3.3.5.3). Sofern kein Diagnosevorschlag gegeben ist, ist die Fehlerdiagnose selbstverständlich aufwändiger, da die Einschränkung des Entscheidungssuchraumes erst durch den Probanden vorgenommen werden muss. Somit ist auch bei vollständiger Überprüfung der AFIRA-Diagnosen ein Vorteil gegenüber der rein manuellen Fehlerdiagnose gegeben. Mit Blick auf die Parameter Temperatur, Luftfeuchte und Kohlenstoffdioxid traten in der hier beschriebenen Untersuchung keine Fehlfunktionen auf.

3 EXPERIMENT I: ZUM EINFLUSS VON FEHLDIAGNOSEN

3.1 VORBEMERKUNG

Die im Folgenden berichtete Studie wurde an der Technischen Universität Berlin durchgeführt. Die Datenerhebung erfolgte dabei im Rahmen einer von der Autorin dieser Arbeit betreuten Diplomarbeit (Hüper, 2005), wobei ein Teil der erhobenen Daten bereits für die Diplomarbeit herangezogen wurde (vgl. auch Hüper, Bahner & Manzey, 2005). Die hier berichteten Fragestellungen umfassen in Teilen jene der Diplomarbeit, gehen jedoch deutlich über diese hinaus.

3.2 HYPOTHESEN EXPERIMENT I

Hypothese 1: Probanden, die im Training Fehldiagnosen eines Assistenzsystems erfahren, zeigen in der späteren Arbeit mit dem System eine **geringere Ausprägung von *complacency* gegenüber der Diagnosefunktion** als Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert worden sind.

Diese Erwartung liegt nahe, da davon auszugehen ist, dass die Erfahrung von Fehlern eines automatisierten Systems sich in einer unmittelbaren Minderung des Vertrauens in Automation niederschlägt und in der Folge auch *complacency* gegenüber der als fehlerhaft erlebten Teilfunktion zu reduzieren vermag.

Hypothese 2: Probanden, die im Training Fehldiagnosen eines Assistenzsystems erfahren, werden in der späteren Arbeit mit dem System in vom Assistenzsystem als fehlerfrei angezeigten Phasen **mehr Informationen über den Systemzustand abrufen** als Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert worden sind.

Da Muir und Moray (1996; vgl. Kap. 1.3.2) zeigten, dass Fehler einer bestimmten automatisierten Funktion sich auch mit Blick auf andere Automationsfunktionen vertrauensmindernd auswirken, wird davon ausgegangen, dass Fehldiagnosen auch in der hier berichteten Studie generell das Vertrauen in das Assistenzsystem mindern. Wenn eine solche Generalisierung des Vertrauens in die Diagnosefunktion des Assistenzsystems auf dessen Alarmfunktion stattfindet, müsste sich dies auch auf Verhaltensebene in einer vermehrten Überwachung des Systems in vermeintlich fehlerfreien Phasen widerspiegeln.

Hypothese 3: Probanden, die im Training Fehldiagnosen eines Assistenzsystems erfahren haben, werden in der späteren Arbeit mit dem System **weniger *commission* Fehler** begehen als

Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert worden sind.

Diese Annahme folgt aus den Überlegungen, dass Fehldiagnosen im Training zu einer Minderung von *complacency* gegenüber der Diagnosefunktion des Systems führen (vgl. Hypothese 1) und dass *commission* Fehler eine mögliche Folge von *complacency* darstellen. Folglich ist zu erwarten, dass die Fehlererfahrung im Training auch *commission* Fehler zu vermindern vermag.

Hypothese 4: Probanden, die einen ***commission* Fehler** begehen, zeigen **höhere *complacency*-Ausprägungen gegenüber der Diagnosefunktion** im Vergleich zu Probanden, die keinen *commission* Fehler begehen.

Diese Hypothese adressiert direkt den Zusammenhang von *automation bias* und *complacency* und reflektiert die Annahme, dass *complacency* im Sinne einer mangelnden Überprüfung bzw. Überwachung eine mögliche Ursache für *commission* als auch *omission* Fehler darstellt. In der hier berichteten Studie werden jedoch nur *commission* Fehler untersucht, sodass auch nur diesbezüglich eine Hypothese aufgestellt werden kann.

Hypothese 5: Probanden, die im Training Fehldiagnosen eines Assistenzsystems erfahren haben, zeigen **geringere Fertigkeitsverluste** bei der manuellen Fehlerdiagnose, die aufgrund eines Ausfalls des Assistenzsystems plötzlich erforderlich wird, als Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert worden sind.

Wenn die Erfahrung von Fehldiagnosen im Training dazu führt, dass Probanden Diagnosen des Systems im weiteren Verlauf des Experiments in stärkerem Maße überprüfen (vgl. Hypothese 1), so würde dies auch eine vermehrte Übung der erforderlichen Diagnoseschritte bedeuten. Wird dann bei einem Systemausfall eine rein manuelle Fehlerdiagnose erforderlich, so ist zu erwarten, dass Probanden, die Fehler erfahren haben, jenen, die nur über deren Eventualität informiert worden sind, überlegen sind. Einen Beleg dafür, dass geringere *complacency*-Ausprägungen mit besseren manuellen Diagnoseleistungen bei Ausfällen des Assistenzsystems einhergehen können, liefern Befunde von Lorenz et al (2002).

Hypothese 6: Probanden, die im Training Fehldiagnosen eines Assistenzsystems erfahren, berichten ein **geringeres Vertrauen** in die Automation als Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert worden sind.

Diese letzte Hypothese leitet sich aus der Annahme ab, dass die direkte Erfahrung von Automationsfehler unmittelbar zu einer Minderung des Vertrauens in Automation führt, während

die abstrakte Information, dass Fehler auftreten können, das subjektive Vertrauen weniger stark beeinflusst.

3.3 METHODE EXPERIMENT I

3.3.1 Stichprobe

$N = 24$ Studentinnen und Studenten der Ingenieurwissenschaften der Technischen Universität Berlin nahmen an der Studie teil. Die Probanden wurden über Bekanntgabe in Vorlesungen und E-Mail-Verteilern geworben. Die zwölf weiblichen und zwölf männlichen Probanden waren zwischen 18 und 26 Jahren alt ($M = 23.08$, $SD = 3.04$) und verfügten nicht über Vorerfahrungen mit der genutzten Versuchsumgebung. Als Aufwandsentschädigung erhielten sie nach der Teilnahme 40 Euro.

3.3.2 Versuchsplan

Der hier beschriebenen experimentellen Studie lag ein einfaktorieller (Fehlererfahrung vs. Fehlerinformation) Versuchsplan zugrunde, wobei die beiden Faktorstufen in verschiedenen Probandengruppen realisiert wurden (*between Design*).

Um sicherzustellen, dass die beiden experimentellen Gruppen sich nicht von vornherein in zentralen abhängigen Variablen unterschieden, wurde die Gruppenzuteilung basierend auf der Leistung am ersten Untersuchungstag, der dem manuellen Fehlerdiagnose- und Fehlermanagementtraining der Probanden gewidmet war, vorgenommen. Herangezogen wurden hierzu die folgenden Variablen:

- Manuelle Fehlerdiagnosezeit. Für die Berechnung dienten die Daten aus dem CAMS-Training, wobei über alle zehn Fehlfunktionen hinweg die mittlere Dauer zwischen Auftreten einer Fehlfunktion und dem Absenden des richtigen Reparaturauftrags ermittelt wurde.
- Vertrauen in CAMS. Grundlage war hier ein entsprechendes Item des Fragebogens zum allgemeinen Vertrauen in Automation (Item 4a, vergleiche nachfolgendes Kap. 3.3.3.1).
- Selbstvertrauen bezüglich der Bedienung von CAMS. Auch hier diente ein entsprechendes Item des Fragebogens als Grundlage (Item 5a, vergleiche nachfolgendes Kap. 3.3.3.1).

Die Mittelwerte der beiden experimentellen Gruppen nach der Einteilung sind Anhang B zu entnehmen.

Die Bedingungsmanipulation wurde am zweiten Untersuchungstag während des AutoCAMS-Trainings (siehe Kap. 3.3.4.2) realisiert. Alle Probanden erhielten zu Beginn des Trainings die Information, dass Fehler des Assistenzsystems vorkommen können und daher die Diagnose und das vorgeschlagene Fehlermanagement überprüft werden sollten. Im Rahmen der Trainingseinheit generierte AFIRA jedoch nur bei der Hälfte der Probanden ($n = 12$) bei zwei von insgesamt zehn Fehlfunktionen falsche Diagnosen (**Erfahrungsgruppe**). Die Fehldiagnosen traten dabei für die Probanden nicht vorhersehbar auf und waren mit Kosten verbunden, da durch sie eine aufwändige manuelle Fehlerdiagnose und -behebung erforderlich wurde. Zudem handelte es sich bei den Fehldiagnosen um aus Probandensicht nicht nachvollziehbare, vermeidbare Fehler: Der tatsächlich vorliegende Fehler und der von AFIRA gemeldete Fehler wiesen keine Überschneidungen in ihren Symptommustern auf. Die Voraussetzungen für einen vertrauensmindernden Effekt der Automationsfehler wurden somit sichergestellt.

Demgegenüber waren die AFIRA-Diagnosen für die andere Hälfte der Probanden ($n = 12$) immer korrekt. Damit diente dieser Gruppe die Information im Rahmen der Instruktion als einziger Hinweis auf eine nicht 100-prozentige Zuverlässigkeit des Systems (**Informationsgruppe**).

Die abhängigen Variablen wurden in einer auf das AutoCAMS-Training folgenden Testphase erfasst und werden – der besseren Nachvollziehbarkeit wegen – im Anschluss an die Beschreibung der Durchführung in einem gesonderten Kapitel (3.3.5) beschrieben.

3.3.3 Material

3.3.3.1 Fragebogen zum allgemeinen Vertrauen in Automation

Da sowohl *complacency* als auch *automation bias* mit einem übersteigerten Vertrauen verbunden sind, wurde in der hier beschriebenen Studie das Vertrauen in Automation sowie das Selbstvertrauen in die manuelle Ausführung der Funktion erhoben. Um zu vermeiden, den zentralen Untersuchungsgegenstand für die Probanden, etwa durch eine isolierte Erhebung der beiden Variablen (Vertrauen in Automation und Selbstvertrauen), transparent werden zu lassen, wurden die entsprechenden Items in einen am NASA-Task-Load-Index (Hart & Staveland, 1988) orientierten Fragebogen eingebettet. Zwei Varianten des Fragebogens wurden eingesetzt, wobei sich die eine auf CAMS (ohne AFIRA) und die andere auf AutoCAMS (mit AFIRA) bezog. Je nachdem mit Blick auf welches System die Befragung erfolgte, wurden einige Items in abgeänderter Form bzw. nicht erfragt. Beide Fragebogenversionen wurden rechnergestützt dargeboten. Items, Antwortformate und Zielvariablen des eingesetzten Fragebogens sind Tabelle 1 zu entnehmen, wobei aus der ersten Spalte die Zuordnung der Items zu den beiden Fragebogenvarianten hervorgeht. I-

temnummern mit dem Zusatz „a“ fanden im CAMS-Fragebogen, mit dem Zusatz „b“ im Auto-CAMS-Fragebogen Verwendung.

Tab. 1: Exp. I: CAMS- und AutoCAMS-Fragebogen

Nr.	Item	Zielvariable	Antwort-format
Bitte geben Sie an, wie die einzelnen Größen bei der Arbeit mit CAMS [für Variante a; für Variante b: „AutoCAMS“] ausgeprägt waren:			
1 a, b	Geistige Anforderungen	Füllitem	
2 a, b	Körperliche Anforderungen	Füllitem	
3 a, b	Zeitliche Anforderungen	Füllitem	
4 a	Vertrauen in CAMS	Vertrauen in CAMS	gering/ hoch (5 stufig)
4 b	Vertrauen in AFIRA	Vertrauen in AFIRA	
5 a	Selbstvertrauen bezüglich der Bedienung von CAMS	Selbstvertrauen (Auto-CAMS)	
5 b	Selbstvertrauen bezüglich der Bedienung von AutoCAMS	Selbstvertrauen (CAMS)	
6 a, b	Eigene Leistung	Füllitem	gut/ schlecht (5 stufig)
7 a, b	Eigene Anstrengung	Füllitem	gering/ hoch
8 a, b	Frustration	Füllitem	(5 stufig)

3.3.3.2 Handout zur Fehlerdiagnose und -behebung

In bestimmten Untersuchungsphasen (vgl. nachfolgendes Kapitel 3.3.4) stand den Probanden ein Handout zur Verfügung, das die erforderlichen Fehlerdiagnose und -behebungsschritte übersichtlich zusammenfasste. Ein Ausdruck desselben findet sich in Anhang C.

3.3.4 Durchführung

Die Durchführung der Untersuchung erfolgte im Frühjahr 2005 an der Technischen Universität Berlin. Jeweils vier Probanden nahmen gleichzeitig teil und kamen an zwei verschiedenen Tagen für jeweils ca. vier Stunden zur Versuchsdurchführung.

3.3.4.1 Erster Untersuchungstag

Der erste Tag der Untersuchung diente dem Erwerb der manuellen Fehlerdiagnose- und -managementfertigkeiten. Ein Ausdruck der unterstützend eingesetzten MS Powerpoint-Präsentation findet sich im Anhang D. Zunächst erläuterte die Versuchsleiterin allgemein die Funktionsweise von CAMS und der fünf Subsysteme, die einzelnen Elemente der Benutzeroberfläche sowie die Aufgaben der Probanden während der Arbeit mit CAMS. Letztere wurden ohne instruktive Priorisierung als Halten der Parameter im Sollbereich, Fehlermanagement, Ressourcen sparen, Bestätigung der Verbindung mit der Raumstation und Erfassung des Kohlenstoffdioxidwerts in regelmäßigen Abständen beschrieben.

Anschließend wurden die verschiedenen Fehlfunktionen von CAMS vorgestellt und jeweils erläutert, durch welches Symptommuster sich diese äußern, und wie sie behoben werden können. Parallel bearbeiteten die Probanden am Rechner Übungsaufgaben, um sich zum einen mit dem System und mit ihren Aufgaben vertraut zu machen, zum anderen, um die manuelle Fehlerdiagnose und -behebung einzuüben. Jede Fehlfunktion wurde dabei, nachdem sie von der Versuchsleiterin vorgestellt worden war, einmal geübt. Im Anschluss daran folgte eine Übung, in der die Probanden drei, zwar vorher geübte aber bezüglich Auftretenszeitpunkt und Typ nicht antizipierbare, Fehlfunktionen diagnostizieren und beheben mussten. Dabei stand den Probanden unterstützend das Handout zur Fehlerdiagnose und -behebung (siehe Anhang C) zur Verfügung.

Im Anschluss an diese einführende Übungsphase folgte ein einstündiges CAMS-Training, in dem die Probanden vier mal 15 Minuten mit dem System arbeiteten. Nur während des ersten 15-Minuten-Blocks stand den Probanden das Handout zur Fehlerdiagnose und -behebung zur Verfügung. Über die vier Übungsblöcke verteilt, traten zehn, weder bezüglich des Fehlertyps noch des Auftretenszeitpunkts vorhersehbare, Fehlfunktionen auf. Tabelle 2 zeigt den Ablauf des Trainings mit Blick auf die einzelnen Blöcke und die darin auftretenden Fehlfunktionen.

Tab. 2: Exp. I: Ablauf des CAMS-Trainings mit Blick auf auftretende CAMS-Fehlfunktionen

Block (jeweils 15 Min)	Nr.	Fehlfunktion	Auftretenszeitpunkt (Min. nach Blockbeginn)
Block 1	1	N ₂ Blockierung des Ventils	2
	2	O ₂ Sensordefekt	9
Block 2	3	N ₂ offen steckendes Ventil	1
	4	O ₂ Blockierung des Ventils	6
	5	O ₂ Leck	11

Block (jeweils 15 Min)	Nr.	Fehlfunktion	Auftretenszeitpunkt (Min. nach Blockbeginn)
Block 3	6	O ₂ Sensordefekt	2
	7	O ₂ Blockierung des Ventils	9
Block 4	8	N ₂ Sensordefekt	1
	9	O ₂ offen steckendes Ventil	6
	10	N ₂ Leck	11

Wie ersichtlich, trat jede Fehlfunktion mindestens einmal auf. Mit Blick auf den Auftretenszeitpunkt der Fehler wurde darauf geachtet, dass jeweils mindestens fünf Minuten zwischen zwei Fehlern lagen, um so die Wahrscheinlichkeit, jeden Fehler vor Auftreten des nächsten Fehlers zu beheben, zu maximieren. Entsprechende Erfahrungswerte lagen durch Vorversuche vor.

Im Anschluss an das CAMS-Training und als Abschluss des ersten Untersuchungstages füllten die Probanden den in Kapitel 3.3.3.1 beschriebenen CAMS-Fragebogen aus.

3.3.4.2 Zweiter Untersuchungstag

Zwischen erstem und zweitem Untersuchungstag lagen für die Probanden im Mittel 9.29 ($SD = 3.77$) Tage. Bei der Terminplanung wurde darauf geachtet, dass jeweils vier Probanden derselben Bedingung teilnahmen, um so in der Informationsgruppe „stellvertretenden“ Fehlererfahrungen durch Berichte von Probanden der Erfahrungsgruppe vorzubeugen. Zu Beginn des zweiten Tages wurden kurz die Funktionsweise von CAMS und die Aufgaben der Probanden wiederholt. Es folgte eine präsentationsgestützte Einführung in AutoCAMS (vgl. Anhang E). Hinsichtlich der Reliabilität des Assistenzsystems erhielten alle Probanden die folgende Information:

„Da die Möglichkeit besteht, dass Fehler des Assistenzsystems vorkommen, sollten die Diagnose und das vorgeschlagene Fehlermanagement überprüft werden.“

Nach dieser Einführung folgte ein einstündiges, in vier 15-minütige Blöcke aufgeteiltes Training mit Assistenzsystemunterstützung, wobei es für alle Probanden galt, insgesamt zehn Fehlfunktionen zu diagnostizieren und zu beheben. Bei der Hälfte der Probanden generierte AFIRA bei zwei dieser zehn Fehlfunktionen falsche Diagnosen (**Erfahrungsgruppe**), während für die andere Hälfte der Probanden stets korrekte Diagnosen von AFIRA vorgeschlagen wurden (**Informationsgruppe**).

Alle Probanden füllten nach dem zweiten der vier Blöcke den AutoCAMS-Fragebogen aus, wobei die Fehldiagnosen in der Erfahrungs-Gruppe erst danach, am Ende des dritten bzw. zu Beginn des vierten Blocks auftraten. In der nachfolgenden Tabelle 3 ist der Ablauf der Trainingsphase mit Blick auf die auftretenden Fehlfunktionen festgehalten, die in unregelmäßigen, für die Probanden nicht antizipierbaren Zeitabständen auftraten. Zwischen zwei Fehlfunktionen lagen mindestens vier Minuten, sodass den Probanden ausreichend Zeit zur vollständigen Behebung eines Fehlers vor Beginn des nächsten Fehlers zur Verfügung stand.

Tab. 3: Exp. I: Ablauf des AutoCAMS-Trainings mit Blick auf auftretende CAMS-Fehlfunktionen

Block (jeweils 15 Min)	Nr.	Fehlfunktion (tatsächlich vorliegend)	Auftretenszeit- punkt (Min. nach Blockbeginn)	Ggf. falsche AFIRA- Diagnosen (nur Erfahrungsgruppe)
Block 1	1	O ₂ Blockierung des Ventils	1	
	2	N ₂ Sensordefekt	7	
Block 2	3	N ₂ Blockierung des Ventils	1	
	4	O ₂ Leck	5	
	5	O ₂ Sensordefekt	10	
Block 3	6	O ₂ Sensordefekt	2	
	7	N ₂ Leck	6	
	8	N ₂ Sensordefekt	10	AFIRA: O ₂ Sensordefekt
Block 4	9	N ₂ Sensordefekt	4	
	10	N ₂ Leck	10	AFIRA: N ₂ offen steckendes Ventil

Nach dem AutoCAMS-Training und damit nach der Bedingungsmanipulation folgte zunächst eine weitere Erhebung des AutoCAMS-Fragebogens. Daran schloss sich die 90-minütige Testphase an, unterteilt in drei Blöcke mit jeweils 15-minütiger Dauer und einem vierten Block mit 30-minütiger Dauer. Zwischen den vier Blöcken erhielten die Probanden die Gelegenheiten 5-minütiger Pausen. Tabelle 4 zeigt den Ablauf der Testphase mit Blick auf die CAMS-Fehlfunktionen und deren Auftretenszeitpunkt im jeweiligen Block.

Tab. 4: Exp. I: Ablauf der Testphase mit Blick auf auftretende CAMS-Fehlfunktionen

Block	Nr.	Fehlfunktion (tatsächlich vorliegend)	Auftretenszeit- punkt (Min. nach Blockbeginn)	AFIRA-Fehlfunktionen (alle Probanden)
Block 1	1	O ₂ Sensordefekt	3	
	2	O ₂ offen steckendes Ventil	11	
Block 2	3	O ₂ Blockierung des Ventils	1	
	4	O ₂ Sensordefekt	6	
	5	N ₂ Blockierung des Ventils	9	
Block 3	6	O ₂ Leck	2	
	7	N ₂ Sensordefekt	9	
	8	N ₂ offen steckendes Ventil	13	
Block 4	9	N ₂ Leck	2	
	10	O ₂ offen steckendes Ventil	7	AFIRA: O ₂ Blockierung des Ventils
	11	N ₂ Blockierung des Ventils	17	AFIRA-Ausfall (keine Diagnose)
	12	O ₂ Sensordefekt	23	AFIRA-Ausfall (keine Diagnose)

Während dieser Testphase generierte AFIRA für die ersten neun CAMS-Fehlfunktionen jeweils die richtige Diagnose, bei der zehnten Fehlfunktion wurde von AFIRA jedoch eine falsche Diagnose ausgegeben: Statt des tatsächlich vorliegenden offen steckenden Sauerstoffventils, diagnostizierte das Assistenzsystem eine Blockierung des Sauerstoffventils. Sowohl die tatsächlich vorliegende Fehlfunktion als auch der fälschlich von AFIRA vorgeschlagene Fehler waren vorher wiederholt aufgetreten (2 bzw. 3 Mal), sodass davon auszugehen ist, dass alle Probanden hinreichend geübt waren, um die AFIRA-Diagnose falsifizieren zu können. Zudem war bei der vorliegenden Kombination von tatsächlicher Fehlfunktion und Diagnosevorschlag eine Falsifikation vergleichsweise einfach möglich. Aufgrund des gegensätzlichen Symptommusters (Anstieg versus Abnahme der Sauerstoffkonzentration) genügte bereits der Abruf einer Informationsquelle (Verlaufsanzeige oder Steuerungsmenü O₂), um die vorgeschlagene Diagnose als fehlerhaft zu identifizieren.

Für die nachfolgenden Fehler 10 und 12 wurde ein Ausfall des Assistenzsystems simuliert. Das Vorliegen einer Fehlfunktion wurde zwar jeweils durch den Masteralarm („rot“) indiziert,

AFIRA zeigte jedoch weder eine Fehlerdiagnose noch die entsprechende Handlungssequenz zur Behebung der beiden Fehlfunktionen an.

Nach der Testphase füllten die Probanden ein letztes Mal den AutoCAMS-Fragebogen aus.

3.3.5 Abhängige Variablen

Soweit nicht anders vermerkt, wurden die abhängigen Variablen aus den über Logfiles protokollierten Mausklicks und Tastatureingaben der Probanden unter Berücksichtigung des jeweils vorliegenden Systemzustands abgeleitet.

3.3.5.1 Fehlerdiagnosezeit

Die Fehlerdiagnosezeit für die ersten neun Fehlfunktionen von CAMS, bei denen AFIRA korrekte Diagnosen lieferte, diente als allgemeiner Leistungsparameter für die Fehlerdiagnose. Operationalisiert wurde diese über die Zeitdauer (s) zwischen dem Auftreten einer CAMS-Fehlfunktion und dem Abschicken des richtigen Reparaturauftrags. Potenziell falsche Reparaturaufträge, die vor dem richtigen Auftrag abgeschickt wurden, blieben folglich außer Acht. Der Auftretenszeitpunkt einer Fehlfunktion wurde definiert als Zeitpunkt, ab dem der Fehler für den Probanden prinzipiell über das jeweils spezifische Symptommuster diagnostizierbar war. Zu diesem Zeitpunkt wurde der Fehler auch durch AFIRA gemeldet und der Masteralarm wechselte auf „rot“, sofern dieser nicht ohnehin schon aufgrund einer noch nicht behobenen vorausgehenden Fehlfunktion „rot“ anzeigte.

3.3.5.2 Informationssuchverhalten allgemein in Fehlerdiagnosephasen

Als allgemeiner, quantitativer Indikator des Informationssuchverhaltens in der Fehlerdiagnosephase wurde berechnet, wie oft Informationen zum Systemstatus überhaupt abgerufen wurden. Dies geschah wiederum bezogen auf die ersten neun Fehlfunktionen mit korrekten AFIRA-Diagnosen. Als Fehlerdiagnosephase wurde die Zeit zwischen dem Auftreten einer Fehlfunktion und dem Abschicken des ersten Reparaturauftrages definiert. Ausgezählt wurde für diese Phasen, wie oft die Probanden Informationsquellen per Mausklick öffneten (Informationsabrufe gesamt).

Zusätzlich wurde betrachtet, wie hoch der Anteil relevanter Informationsabrufe an der Gesamtzahl von Abrufen war. Als relevant wurden dabei diejenigen definiert, die für eine Überprüfung der vorgeschlagenen Diagnose nützlich waren, während in die Gesamtzahl auch „irrelevante“ Informationsquellen, die keinen Bezug zur Diagnose hatten, einbezogen wurden.

In Anhang F findet sich eine Übersicht über relevante und irrelevante Informationsquellen in Abhängigkeit der verschiedenen Fehlerdiagnosen.

3.3.5.3 Complacency

Das Überprüfungsverhalten der Probanden gegenüber den von AFIRA generierten Diagnosen diente der Operationalisierung von *complacency*. Der Forderung von Moray (2000; vgl. Kap. 1.4.4) Rechnung tragend, *complacency* erst dann zu attestieren, wenn ein im Vergleich zu einem „normativen Modell“ unzureichendes Informationssuchverhalten vorliegt, wurde zunächst festgelegt, welche Informationsquellen notwendiger Weise abzurufen sind, um von AFIRA vorgeschlagene Fehlerdiagnosen zu überprüfen. Die folgende Übersicht (Tab. 5) zeigt, welche Informationsquellen abgerufen werden müssen, um die von AFIRA vorgeschlagenen Diagnosen zu verifizieren bzw. zu falsifizieren (vgl. hierzu auch Anhang A).

Tab. 5: Zur Diagnoseüberprüfung notwendige Prüfelemente

Diagnosevorschlag	Zur Überprüfung notwendige Prüfelemente (Informationsquellen bzw. Systemeingriffe)	Anzahl
O ₂ Leck im Ventil	Gasströme, O ₂ Verlaufsanzeige ODER O ₂ Steuerungsmenü	2
O ₂ Blockierung des Ventils	Gasströme, O ₂ Verlaufsanzeige ODER O ₂ Steuerungsmenü, Standards Gasströme	3
O ₂ Offen steckendes Ventil	O ₂ Verlaufsanzeige, O ₂ Steuerungsmenü, Steuerung aus, Steuerung automatisch	4
O ₂ Sensordefekt (steigend)	O ₂ Verlaufsanzeige, O ₂ Steuerungsmenü, Steuerung aus, Steuerung automatisch	4
O ₂ Sensordefekt (sinkend)	O ₂ Verlaufsanzeige ODER O ₂ Steuerungsmenü, Gasströme	2
N ₂ Leck im Ventil	Gasströme, N ₂ Verlaufsanzeige ODER N ₂ Steuerungsmenü	2
N ₂ Blockierung	Gasströme, N ₂ Verlaufsanzeige ODER N ₂ Steuerungsmenü, Standards Gasströme	3
N ₂ Offen steckendes Ventil	N ₂ Verlaufsanzeige ODER N ₂ Steuerungsmenü, O ₂ Verlaufsanzeige ODER O ₂ Steuerungsmenü	2
N ₂ Sensordefekt (steigend)	Gasströme, N ₂ Verlaufsanzeige ODER N ₂ Steuerungsmenü	2
N ₂ Sensordefekt (sinkend)	N ₂ Verlaufsanzeige ODER N ₂ Steuerungsmenü	1

Aus vorangehender Tabelle geht hervor, dass bei den meisten Fehlfunktionen die Verlaufsanzeige des jeweiligen Parameters „oder“ das Steuerungsmenü als notwendige Informationsquelle klassifiziert wurden. Dies liegt darin begründet, dass die Probanden zwar daraufhin trainiert wurden, die Verlaufsanzeige abzurufen, um sich über Anstieg bzw. Abnahme der Parameter zu informieren, diese Information kann jedoch auch dem Steuerungsmenü entnommen werden, da dort der jeweilige Parameterwert in digitaler Form ausgewiesen wird. Daher wurde in diesen Fällen der Abruf nur eines der beiden Parameter als notwendig definiert. Hierin besteht auch der Unterschied zu den für das allgemeine Informationssuchverhalten als relevant definierten Informationsquellen (vgl. vorhergehendes Kapitel 3.3.5.2): Beide Informationsquellen sind zwar relevant, aber nur eine der beiden Quellen ist notwendig für die Diagnoseprüfung.

Mit Blick auf den N₂ und O₂ Sensordefekt hängen die Symptommuster und damit auch die notwendigen Prüfelemente davon ab, ob die Fehlfunktion bei gerade steigendem oder sinkendem Parameterverlauf auftritt. Dieser Verlauf ist durch den vorgegebenen und für alle Probanden konstant gehaltenen Auftretenszeitpunkt einer Fehlfunktion nicht eindeutig bestimmt, sondern durch mögliche vorhergehende manuelle Steuereingriffe der Probanden beeinflussbar. Daher wurde zunächst anhand der Logfiles für jeden Probanden und jeden Sensordefekt festgestellt, ob die Fehlfunktion bei steigendem oder sinkendem Parameterverlauf auftrat und die notwendigen Prüfelemente in Abhängigkeit davon definiert.

Diese fehlerspezifische Klassifikation von notwendigen Informationsquellen diene als Grundlage, um für die ersten neun, stets korrekten AFIRA-Diagnosen feststellen zu können, ob die Probanden einer optimalen Überprüfung der Automation nachkamen oder aber im Sinne eines *complacency*-Effekts unzureichend überprüften. Hierzu wurde analysiert, wie viele der eigentlich notwendigen Informationsquellen von den Probanden im Zuge einer Diagnosenprüfung, d.h. zwischen dem Auftreten einer Fehlfunktion und dem Absenden des ersten Reparaturauftrags (unabhängig davon ob dieser richtig oder falsch war), herangezogen wurden.

Im Detail diene die Anzahl herangezogener notwendiger Informationsquellen geteilt durch die Gesamtanzahl notwendiger Informationsquellen als Maß für die Diagnoseüberprüfung.

Da sich die einzelnen Fehlfunktionen mit Blick auf die Anzahl erforderlicher Informationsquellen erheblich unterscheiden, wurden bezüglich der ersten neun Fehlfunktionen mit korrekten AFIRA-Diagnosen jeweils drei aufeinander folgende Fehlfunktionen (Fehler 1-3, 4-6, 7-9) zusammengefasst. Die Summe der für die drei Fehler herangezogenen notwendigen Informationsquellen wurde dann geteilt durch die Summe der für die drei Fehler insgesamt notwendigen Informationsquellen. Der sich daraus ergebende Wert variiert zwischen 0 und 1, wobei 0 als keine Überprüfung und damit maximale Ausprägung von *complacency* und 1 als vollständige Überprüfung

fung der Diagnosen bzw. Abwesenheit von *complacency* zu interpretieren sind. Diese Operationalisierung ermöglichte somit auch eine Quantifizierung verschiedener *complacency*-Ausprägungen.

3.3.5.4 Informationssuchverhalten in fehlerfreien Phasen

Auch in den fehlerfreien Phasen, d.h. zwischen der vollständigen Behebung einer Fehlfunktion und dem Auftreten der nächsten Fehlfunktion, wurde das Informationssuchverhalten der Probanden analysiert. Die Dauer dieser fehlerfreien Phasen ist von den individuell eingesetzten Fehlerdiagnose- und -behebungsstrategien abhängig und variierte folglich stark interindividuell und von Fehlfunktion zu Fehlfunktion. Um dennoch eine vergleichbare Datenbasis für alle Fehlfunktionen und alle Probanden zu schaffen, wurden jeweils nur die letzten 120 s vor Auftreten des nächsten Fehlers in Betracht gezogen. Sofern eine fehlerfreie Phase kürzer als 120 s ausfiel, wurde diese nicht in die Datenauswertung einbezogen und als fehlender Wert interpretiert. Bezogen auf diese 120 s wurde analysiert, wie oft von den Probanden Informationen abgerufen worden waren. Im Fokus stand zum einen wieder die Gesamtmenge abgerufener Informationen, operationalisiert über die Anzahl der Mausklicks auf Informationsquellen. Die so gewonnenen Häufigkeitswerte wurden jeweils durch zwei geteilt, so dass sie die durchschnittliche Anzahl von Informationsabrufen pro Minute widerspiegeln.

Zum anderen interessierte auch, inwiefern die Informationsabrufe gezielt der Überwachung der Alarmfunktion von AFIRA dienten. Hierzu wurde der Anteil relevanter Informationsabrufe an der Gesamtzahl der Informationsabrufe errechnet. Relevante Informationen repräsentieren dabei jene, die potenziell als Hinweis auf das Vorliegen einer Fehlfunktion interpretiert werden können, wie etwa die O₂ oder N₂ Verlaufsanzeige. Als irrelevant wurden dagegen jene gewertet, die keinen Bezug zu den trainierten Fehlertypen hatten, wie etwa die Verlaufsanzeige von Temperatur und Luftfeuchte. In Anhang G findet sich ein Überblick über die einbezogenen relevanten und irrelevanten Parameter. Auch bei diesen beiden Variablen erfolgte die Analyse bezüglich der ersten neun Fehlfunktionen, bei denen AFIRA korrekte Diagnosen lieferte.

3.3.5.5 Commission Fehler

Commission Fehler konnten beim zehnten Fehler, bei dem AFIRA eine falsche Diagnose generierte, auftreten. Wenn ein Proband dieser Diagnose dennoch folgte, indem er als ersten Reparaturauftrag nach Auftreten der Fehlfunktion einen der falschen Diagnose entsprechenden Reparaturauftrag abschickte, wurde dies als *commission* Fehler gewertet. Weitere, im späteren Verlauf abgeschickte und ggf. der tatsächlich vorliegenden Fehlfunktion entsprechende Reparaturaufträge wurden nicht berücksichtigt.

3.3.5.6 „Return-to-manual“-Leistung

Nach der Fehldiagnose beim zehnten Fehler wurde für zwei Fehlfunktionen ein Systemausfall von AFIRA simuliert, sodass eine rein manuelle Fehlerdiagnose und -behebung erforderlich wurde. Hier dienten zwei Maße als Indikator für die „return-to-manual“-Leistung bzw. möglicherweise auftretende Fertigkeitsverluste der Probanden. Zum einen wurde die Anzahl korrekt abgeschickter Reparaturaufträge analysiert (0, 1 oder 2), wobei jedoch nur die jeweils ersten nach Auftreten der Fehlfunktion abgeschickten Aufträge Berücksichtigung fanden. Zum anderen wurde die durchschnittliche Fehlerdiagnosezeit, d.h. die Zeit zwischen Auftreten einer Fehlfunktion und Absenden des richtigen Reparaturauftrags (unabhängig davon, ob vorher bereits falsche Aufträge abgeschickt worden waren) erfasst. Um Aussagen über mögliche Fertigkeitsverluste treffen zu können, wurde die durchschnittliche Fehlerdiagnosezeit für Fehler 11 und 12 verglichen mit den durchschnittlichen Fehlerdiagnosezeiten im CAMS-Training für dieselben Fehlertypen (Fehler Nummer 1, 2 und 6).

3.3.5.7 Leistung in der prospektiven Gedächtnis- und Reaktionszeitaufgabe

Die Leistung in der prospektiven Gedächtnisaufgabe wurde über die Abweichung (Dauer in s) der Speicherung des Kohlenstoffdioxidwertes vom geforderten Eintragezeitpunkt (volle Minute) operationalisiert. Mittlere Abweichungsmaße wurden dabei getrennt für fehlerfreie Phasen (Masteralarm „grün“) und Fehlerphasen (Masteralarm „rot“) berechnet, wobei auch hier die ersten neun Fehlfunktionen von CAMS (bzw. die fehlerfreien Phasen jeweils davor) die Datengrundlage bildeten. Berücksichtigt wurden dabei nur jene Eintragungen, die eindeutig einer der beiden Phasen zuzuordnen waren, d.h. innerhalb einer Phase fällig und auch erledigt wurden.

Als Maß für die Reaktionszeitaufgabe wurde die Zeitdauer (s) zwischen dem Erscheinen des Verbindungssymbols und dessen Quittierung durch den Probanden per Mausklick definiert. Auch hier wurden getrennt für fehlerfreie Phasen und Fehlerphasen bezüglich der ersten neun Fehlfunktionen mittlere Reaktionszeiten berechnet, wobei wiederum nur eindeutig zuordenbare Sequenzen in die Auswertung einbezogen wurden.

3.3.5.8 Vertrauen in Automation und Selbstvertrauen

Mithilfe des (Auto)CAMS-Fragebogens wurde das Vertrauen in Automation (Item 4a bzw. 4b) sowie das Selbstvertrauen der Probanden (Item 5a bzw. 5b) auf fünfstufigen Ratingskalen (gering – hoch) erfasst. Zur Auswertung wurden die Selbsteinschätzungen in eine numerische Skala von 1 (gering) bis 5 (hoch) übertragen. Insgesamt erfolgten vier Erhebungen der Variablen:

- Nach dem CAMS-Training (CAMS-Fragebogen)
- Vor der Bedingungsmanipulation (nach dem zweiten Block des AutoCAMS-Trainings; AutoCAMS-Fragebogen)
- Nach der Bedingungsmanipulation (im Anschluss an das AutoCAMS-Training; AutoCAMS-Fragebogen)
- Nach der Testphase (AutoCAMS-Fragebogen)

3.3.6 Statistische Auswertung

Die statistische Auswertung der Daten erfolgte mit der Statistiksoftware SPSS (für Windows, Version 15.0). Für alle Tests wurde das Alphaniveau auf .05 gesetzt. Die Überprüfung der Hypothesen erfolgte zum Großteil mithilfe von Varianzanalysen für Messwiederholungsdesigns (ANOVA). Die Voraussetzung der Sphärizität wurde jeweils mit dem Test nach Mauchly überprüft. Sofern die Voraussetzung nicht erfüllt war, erfolgte eine Korrektur nach Greenhouse-Geisser. Sofern Häufigkeitsverteilungen für eine Fragestellung von Interesse waren, fand der exakte Fisher-Yates-Test Einsatz, da die Voraussetzungen des χ^2 -Tests (maximal 20 % der Felder der erwarteten Häufigkeiten kleiner als 5) nicht erfüllt waren.

3.4 ERGEBNISSE EXPERIMENT I

3.4.1 Fehlerdiagnosezeit

Die Fehlerdiagnosezeiten für die ersten neun Fehlfunktionen, bei denen AFIRA stets korrekte Diagnosen lieferte, sind in Abbildung 12 dargestellt. Um einerseits mögliche *time-on-task*-Effekte zu untersuchen, andererseits um intraindividuelle Varianz zu reduzieren, wurden jeweils drei aufeinander folgende Fehlfunktionen zusammengefasst (Fehler 1-3, 4-6 und 7-9) und die entsprechenden Fehlerdiagnosezeiten blockweise gemittelt. Wie in der Abbildung ersichtlich, benötigten die Probanden der Erfahrungsgruppe durchweg mehr Zeit ($M = 40.88$, $SD = 12.57$ s) für die Fehlerdiagnose als Probanden der Informationsgruppe ($M = 29.86$, $SD = 12.68$ s). Die Daten wurden mithilfe einer 2(Gruppe) x 3(Fehlerblock) Varianzanalyse (ANOVA) mit Messwiederholung (Faktor: „Fehlerblock“) ausgewertet. Die Analyse ergab einzig einen signifikanten Haupteffekt für den Faktor „Gruppe“, $F(1, 22) = 4.57$, $p = .04$. Weder der Messwiederholungseffekt noch die Interaktion der Faktoren „Gruppe“ und „Fehlerblock“ wurde signifikant, $F(2, 44) = 1.41$, $p = .26$ bzw. $F < 1$.

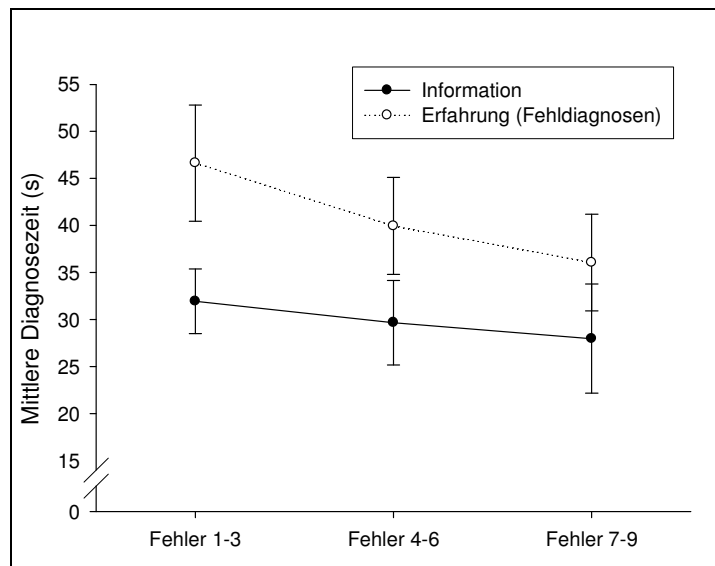


Abb. 12: Exp. 1: Fehlerdiagnosezeit (experimentelle Gruppen)

3.4.2 Überprüfung der Diagnosefunktion

3.4.2.1 Informationssuchverhalten allgemein in Fehlerdiagnosephasen

Auch für die Analyse des Informationssuchverhaltens wurden die ersten neun Fehler in drei Blöcke eingeteilt und die erhobenen Maße für jeden Block gemittelt. Im Mittel wurden vor Abschicken des ersten Reparaturauftrags insgesamt 7.28 ($SD = 2.79$) Informationen abgerufen, wobei der durchschnittliche Anteil relevanter Informationsabrufe .64 ($SD = .12$) betrug. Letzterer nahm über die Zeit hinweg ab, wie in Abbildung 13 ersichtlich wird. Die statistische Auswertung erfolgte wie auch bei der Fehlerdiagnosezeit über 2(Gruppe) x 3(Fehlerblock) Varianzanalysen mit Messwiederholung (ANOVAs). Mit Blick auf die Gesamtanzahl an Informationsabrufen in den Fehlerdiagnosephasen ergab die Auswertung keine signifikanten Effekte, Haupteffekt „Gruppe“ $F(1, 22) = 1.42, p = .25$, Haupteffekt „Fehlerblock“ $F(2, 44) = 3.03, p = .06$, Interaktionseffekt $F(2, 44) = 1.21, p = .31$. Nur bezüglich des Anteils relevanter Informationsabrufe zeigte sich der beschriebene Haupteffekt „Fehlerblock“, $F(2, 44) = 13.38, p < .001$. Gruppenunterschiede wurden jedoch auch bei diesem Maß nicht beobachtet, Haupteffekt „Gruppe“ $F < 1$, wie auch kein Interaktionseffekt zu Tage trat, $F(2, 44) = 1.28, p = .29$. Vor dem Hintergrund, dass keine Abnahme der Informationsabrufe insgesamt beobachtet wurde, kann geschlossen werden, dass über die Zeit hinweg weniger relevante, dafür aber mehr irrelevante Informationen abgerufen wurden.

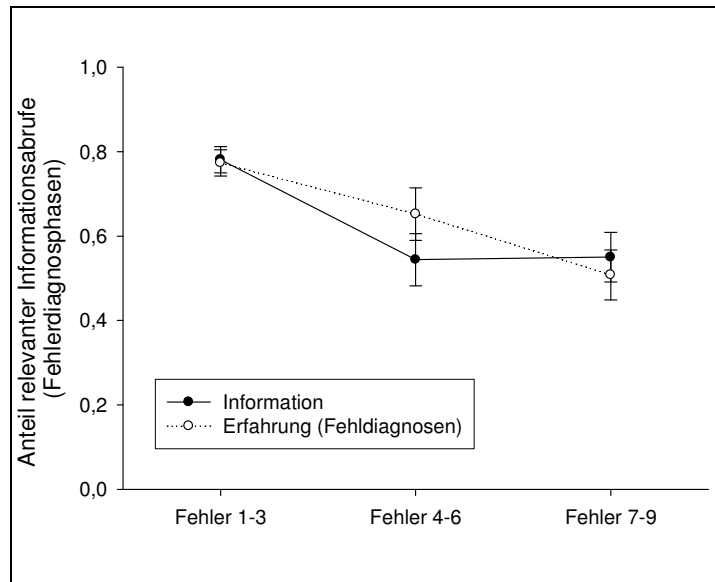


Abb. 13: Exp. I: Anteil relevanter Informationsabrufe in Fehlerdiagnosephasen (experimentelle Gruppen)

3.4.2.2 Complacency

Der Umfang, in dem AFIRA-Diagnosen von den Probanden überprüft wurden, ist in Abbildung 14 dargestellt.

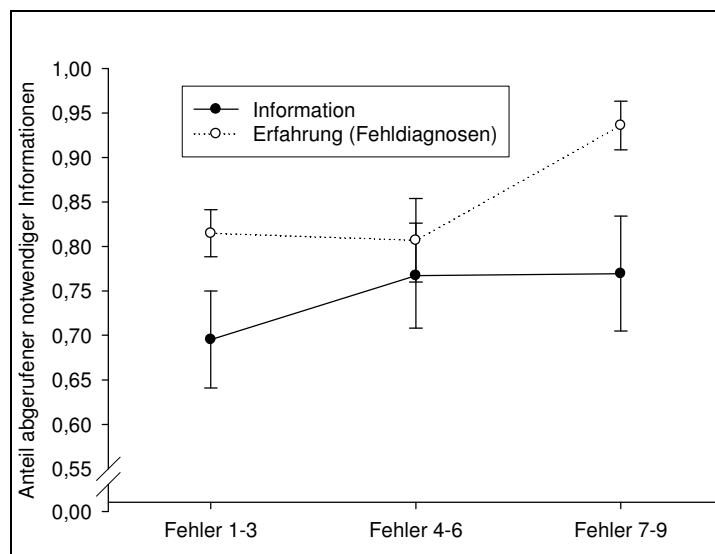


Abb. 14: Exp. I: *Complacency* (experimentelle Gruppen)

Wie zu sehen ist, überprüfte keine der beiden Gruppen die vorgeschlagenen Diagnosen vollständig, bevor sie den Diagnosen folgten und entsprechende Reparaturaufträge abschickten. Bei den Gruppen ist somit *complacency* in gewissem Umfang zu attestieren. Allerdings weisen die Ergebnisse eine deutlich geringere *complacency*-Ausprägung in der Erfahrungsgruppe im Vergleich zur Informationsgruppe aus, wobei Probanden der Erfahrungsgruppe im Mittel 84 % ($M = 0.84$, $SD = 0.07$) der zur Diagnoseprüfung notwendigen Parameter und Probanden der Informations-

gruppe durchschnittlich 73 % ($M = 0.73$, $SD = 0.15$) abriefen. Eine 2(Gruppe) x 3(Fehlerblock) ANOVA mit Messwiederholung ergab einzig einen signifikanten Haupteffekt für den Faktor „Gruppe“, $F(1, 22) = 5.00$, $p = .04$. Für den Faktor „Fehlerblock“ zeigte sich kein Haupteffekt, $F(2, 44) = 2.75$, $p = .08$, auch wurde kein Interaktionseffekt beobachtet, $F(2, 44) = 1.14$, $p = .33$.

3.4.3 Überwachung der Alarmfunktion: Informationssuchverhalten in den fehlerfreien Phasen

In den fehlerfreien Phasen riefen die Probanden durchschnittlich 7.14 ($SD = 6.29$) Informationsquellen pro Minute ab. Der Anteil relevanter Informationsabrufe betrug dabei .55 ($SD = .27$). Die Auswertung erfolgte mithilfe von 2(Gruppe) x 3(Fehlerblock) ANOVAs mit Messwiederholung (Faktor: „Fehlerblock“) und ergab mit Blick auf die Gesamtanzahl von Informationsabrufen keine Unterschiede zwischen der Erfahrungs- und der Informationsgruppe, Haupteffekt „Gruppe“ $F < 1$. Auch zeigte sich kein Messwiederholungs- oder Interaktionseffekt, $F < 1$ bzw. $F(2, 44) = 1.11$, $p = .34$. Wie schon beim Informationssuchverhalten in den Fehlerdiagnosephasen zeigte sich jedoch auch hier eine Abnahme des Anteils relevanter Informationsabrufe über die Zeit (siehe Abb. 15).

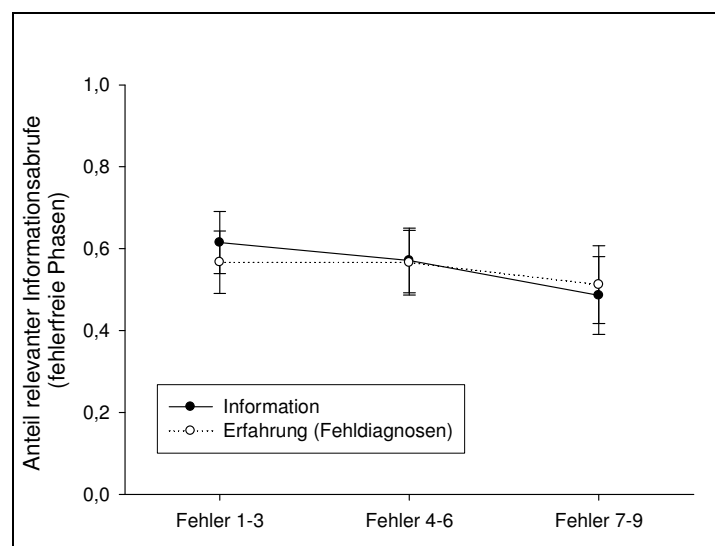


Abb. 15: Exp. I: Anteil relevanter Informationsabrufe in fehlerfreien Phasen (experimentelle Gruppen)

Statistisch schlug sich dies in einem signifikanten Haupteffekt „Fehlerblock“, $F(2, 44) = 3.64$, $p = .03$ nieder, für alle übrigen Effekte $F < 1$.

3.4.4 Commission Fehler

Beim zehnten Fehler generierte AFIRA eine Fehldiagnose. Insgesamt 5 der 24 Probanden folgten dieser, indem sie den vorgeschlagenen, aber falschen, Reparaturvorschlag abschickten, und begingen somit einen *commission* Fehler. Alle Probanden, auch jene, die der Fehldiagnose folgten, hatten vor Abschicken des Reparaturauftrages zumindest die Verlaufsanzeige des Sauerstoffsubsystems eingesehen und damit der vorgeschlagenen Diagnose eindeutig widersprechende Informationen (O_2 -Abnahme statt -Anstieg) abgerufen. Allerdings stammten die Probanden zu gleichen Teilen aus den beiden experimentellen Gruppen (Erfahrungsgruppe: $n = 2$; Informationsgruppe: $n = 3$). Folglich hatte die Bedingungsmanipulation keinen Einfluss auf das Auftreten von *commission* Fehlern, $p = .50$, Fisher-Yates-Test (einseitig).

Vor dem Hintergrund dieses unerwarteten Befundes interessierte umso mehr, was die Probanden, die der Fehldiagnose folgten, gegenüber den Übrigen auszeichnete. Daher wurde in einem zweiten Schritt untersucht, ob Unterschiede zwischen jenen 5 (*commission* Fehler ja) und den 19 Probanden, die die Fehldiagnose als falsch erkannten (*commission* Fehler nein), mit Blick auf das Informationssuchverhalten bei den ersten neun Fehlern existierten. Die Auswertung erfolgte wiederum über $2(\text{commission Fehler ja/nein}) \times 3(\text{Fehlerblock})$ ANOVAs mit Messwiederholung für den zweiten Faktor. Die Ergebnisse zeigen, dass Probanden, die der Fehldiagnose folgten, sich bereits bei den vorherigen neun Fehlern durch signifikant höhere *complacency*-Ausprägungen auszeichneten (siehe Abb. 16).

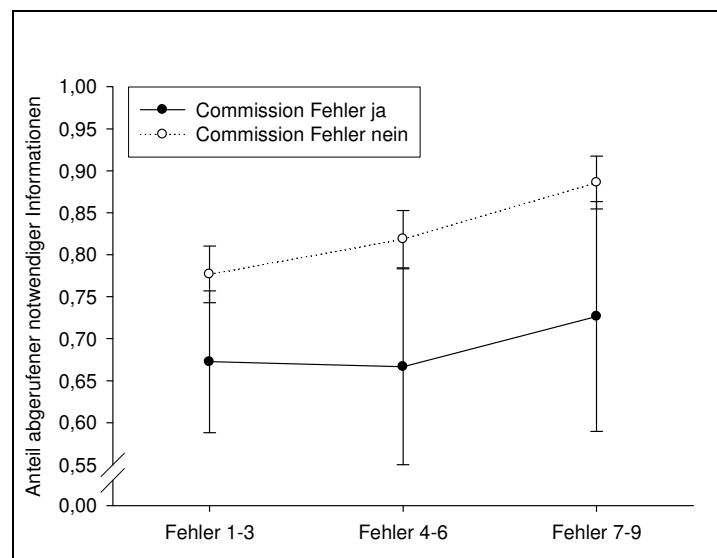


Abb. 16: Exp. I: *Complacency* (*commission* Fehler ja/nein)

Probanden, die die Fehldiagnose erkannten, investierten deutlich mehr Aufwand in die Überprüfung der Diagnosen. Sie riefen einen höheren Anteil der für die Diagnoseprüfung notwendigen Informationsquellen ab ($M = 0.81$, $SD = 0.10$) als jene fünf Probanden, die der Fehldiag-

nose folgten ($M = 0.68$, $SD = 0.19$). Statistisch manifestierte sich dies in einem Haupteffekt für den Faktor „Gruppe“, $F(1, 22) = 5.45$, $p = .03$, während weder ein Haupteffekt „Fehlerblock“ noch ein Interaktionseffekt beobachtet wurde, $F(2, 44) = 1.29$, $p = .29$ bzw. $F < 1$.

Dieser Effekt spiegelte sich auch in höheren Fehlerdiagnosezeiten der Probanden, die einen *commission* Fehler vermeiden konnten, wider (siehe Abb. 17). Diese nahmen sich fast doppelt so viel Zeit für die Fehlerdiagnose ($M = 39.28$, $SD = 11.90$) wie die Probanden, die einen solchen Fehler begingen ($M = 20.51$, $SD = 8.50$), Haupteffekt „*commission* Fehler ja/nein“ $F(1, 22) = 10.80$, $p = .003$, für alle übrigen Effekte $F < 1$.

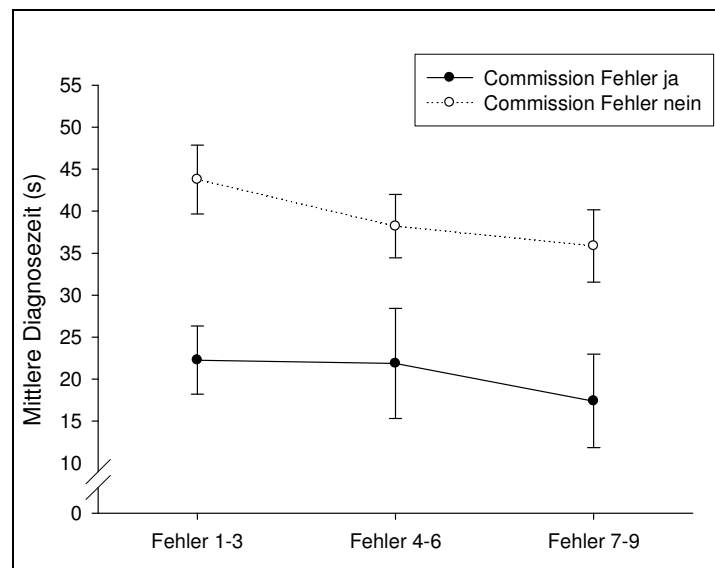


Abb. 17: Exp. I: Fehlerdiagnosezeit (*commission* Fehler ja/nein)

Zudem zeigte sich, dass Probanden die der falschen Diagnose folgten, insgesamt weniger Informationen in den Fehlerdiagnosephasen abrufen (siehe Abb. 18): Die Gesamtanzahl der Informationsabrufe lag hier durchschnittlich bei 4.13 ($SD = 2.36$), während sie für Probanden, die die Fehldiagnose entdeckten, bei 8.11 ($SD = 2.29$) lag, Haupteffekt „*commission* Fehler ja/nein“ $F(1, 22) = 11.85$, $p = .002$. Auch hier wurden keine weiteren signifikanten Effekte beobachtet, Haupteffekt „Fehlerblock“ $F(2, 44) = 1.51$, $p = .23$, Interaktionseffekt $F < 1$.

Eine Analyse des Anteils relevanter Informationsabrufe in der Fehlerdiagnosephase zeigte, wie schon bei der experimentellen Gruppeneinteilung, eine Abnahme dieses Anteils über die Zeit, Haupteffekt „Fehlerblock“ $F(2, 44) = 9.38$, $p < .001$. Weitere Effekte wurden jedoch nicht beobachtet, Haupteffekt „*commission* Fehler ja/nein“ $F < 1$, Interaktionseffekt $F(2, 44) = 1.99$, $p = .15$.

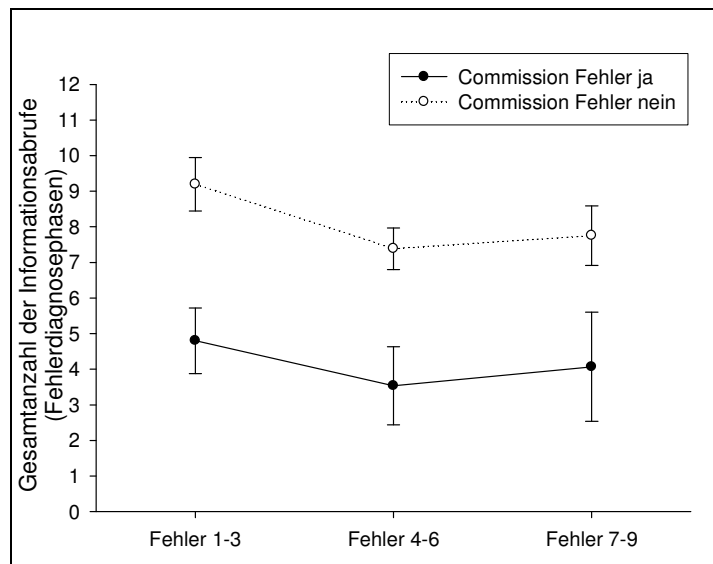


Abb. 18: Exp. I: Informationsabrufe in Fehlerdiagnosephasen
(*commission* Fehler ja/nein)

Mit Blick auf das Informationssuchverhalten in fehlerfreien Phasen zeigten sich ebenfalls keine Unterschiede zwischen den Gruppen. Allerdings zeigten Probanden, die einen *commission* Fehler begingen, gegen Ende des Versuchs eine Abnahme der insgesamt getätigten Informationsabrufe, während jene, die die Fehlerdiagnose erkannten, eine Zunahme beim letzten Fehlerblock beobachten ließen (siehe Abb. 19). Statistisch signifikant wurde entsprechend der Interaktionseffekt, $F(2, 44) = 6.83, p = .003$. Haupteffekte wurden für die Faktoren „Fehlerblock“ und „*commission* Fehler ja/nein“ nicht beobachtet, beide $F < 1$.

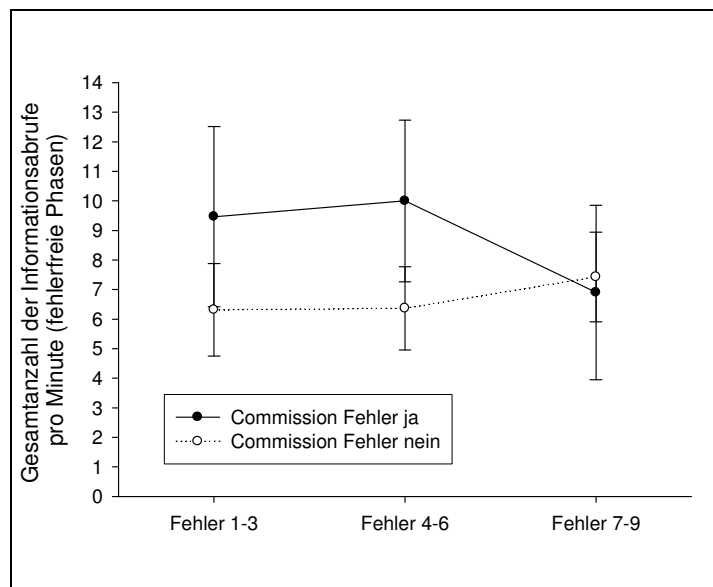


Abb. 19: Exp. I: Informationsabrufe in fehlerfreien Phasen
(*commission* Fehler ja/nein)

In Bezug auf den Anteil relevanter Informationsabrufe zeigten sich keine signifikanten Effekte, Haupteffekt „Fehlerblock“ $F(2, 44) = 2.98, p = .06$, alle übrigen $F < 1$.

3.4.5 „Return-to-manual“-Leistung

Die letzten beiden Fehler (11 und 12) der Testphase mussten manuell diagnostiziert und behoben werden. Nur 3 der 24 Probanden stellten zunächst bei einer der beiden Fehlfunktionen eine falsche Diagnose (Absenden eines falschen Reparaturauftrages), wobei ein Proband im Training die Erfahrung von Fehldiagnosen gemacht hatte, während die übrigen beiden Probanden der Informationsgruppe entstammten. Die experimentelle Manipulation hatte damit keinen Einfluss auf das Auftreten manueller Fehldiagnosen, $p = .50$, Fischer-Yates-Test (einseitig). Auch zeigte sich kein Zusammenhang mit dem Begehen eines *commission* Fehlers, $p = .52$, Fischer-Yates-Test (einseitig).

Als zweite Variable diente die durchschnittliche Fehlerdiagnosezeit, d.h. die Zeit zwischen Auftreten der Fehlfunktion und Absenden des richtigen Reparaturauftrags (unabhängig davon, ob vorher bereits falsche Aufträge abgeschickt worden waren). Diese wurde für die beiden Fehlfunktionen 11 und 12 gemittelt und mit der mittleren Fehlerdiagnosezeit der korrespondierenden drei Fehlfunktionen im CAMS-Training verglichen. Die Ergebnisse sind in Abbildung 20 dargestellt. Wie ersichtlich zeigte sich für beide experimentellen Gruppen kein Fertigskeitsverlust, sondern vielmehr ein deutlicher Trainingseffekt. Während des CAMS-Trainings benötigten die Probanden durchschnittlich 132.66 s ($SD = 71.40$) zur manuellen Fehlerdiagnose, während am Ende der Testphase der richtige Reparaturauftrag bereits nach $M = 63.92$ s ($SD = 28.13$) abgeschickt wurde. Eine 2(Gruppe) x 2(Training vs. Systemausfall) ANOVA mit Messwiederholung für den zweiten Faktor diente der Analyse der Daten. Signifikant wurde der Haupteffekt für den Faktor „Training vs. Systemausfall“, $F(1, 22) = 22.01$, $p < .001$. Weder ein Haupteffekt für den Faktor „Gruppe“, $F < 1$, noch ein Interaktionseffekt, $F(1, 22) = 2.01$, $p = .17$, wurden beobachtet.

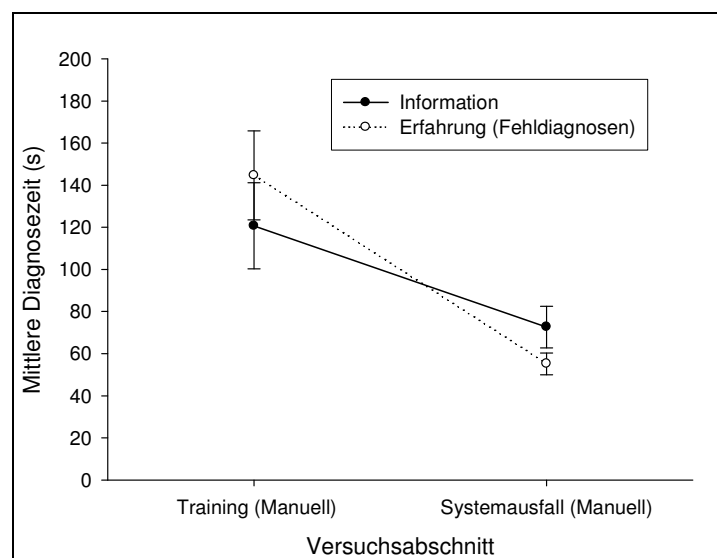


Abb. 20: Exp. I: Fehlerdiagnosezeit beim Systemausfall im Vergleich zum Training (experimentelle Gruppen)

In ähnlicher Weise zeigten sich auch keine Unterschiede in der Fehlerdiagnosezeit zwischen Probanden mit *commission* Fehler und solchen ohne. Eine 2(*commission* Fehler ja/nein) x 2(Training vs. Systemausfall) ANOVA mit Messwiederholung ergab weder einen Unterschied zwischen den Gruppen, Haupteffekt „*commission* Fehler ja/nein“, $F(1, 22) = 1.29, p = .27$, noch einen Interaktionseffekt, $F < 1$. Signifikant wurde jedoch wieder der Haupteffekt für den Faktor „Training vs. Systemausfall“, $F(1, 22) = 17.08, p < .001$.

3.4.6 Leistung in der prospektiven Gedächtnis- und Reaktionszeitaufgabe

Für die Analyse der Daten in den beiden Zusatzaufgaben wurde erneut der Zeitraum der ersten neun Fehlfunktionen, bei denen AFIRA korrekte Diagnosen lieferte, zugrunde gelegt. Für jede fehlerfreie Phase vor einer der neun Fehlfunktionen und jede Fehlerphase wurde die mittlere Abweichung vom geforderten Eintragungszeitpunkt (prospektive Gedächtnisaufgabe) bzw. die mittlere Reaktionszeit (Reaktionszeitaufgabe) berechnet. Anschließend wurden die Messwerte wieder über jeweils drei aufeinander folgende Fehlfunktionen (Fehler 1-3, 4-6 und 7-9) gemittelt. Basierend auf diesen Daten erfolgte die Auswertung mit Hilfe von 2(Gruppe) x 3(Fehlerblock) x 2(Fehlerfrei vs. Fehler) ANOVAs mit den letzten beiden Faktoren als *within*-Faktoren.

Die Darstellung der Ergebnisse erfolgt zuerst für die prospektive Gedächtnisaufgabe, dann für die Reaktionszeitaufgabe und dabei jeweils getrennt für die beiden Gruppeneinteilungen (experimentell sowie *commission* Fehler ja/nein).

Bei der prospektiven Gedächtnisaufgabe erfolgten die Eintragungen des CO₂-Wertes mit einer mittleren Abweichung von 2.63 s ($SD = 1.44$) vom geforderten Zeitpunkt (volle Minute). Die Leistung in dieser Aufgabe wurde dabei weder durch die experimentelle Gruppeneinteilung noch durch das Vorliegen einer Fehlfunktion beeinflusst, Haupteffekte „Fehlerblock“ $F(2, 44) = 1.03, p = .37$, „Fehlerfrei vs. Fehler“ $F(1, 22) = 3.37, p = .08$, Interaktionseffekte „Fehlerblock“ x „Gruppe“ $F(2, 44) = 1.25, p = .30$, alle übrigen $F < 1$.

Wurden die Gruppen nach *commission* Fehler ja bzw. nein eingeteilt, zeigte sich jedoch ein Unterschied zwischen fehlerfreien Phasen und Fehlerphasen: In ersteren wichen die Probanden durchschnittlich um 2.27 s ($SD = 1.62$) vom geforderten Eintragszeitpunkt ab, während in den Fehlerphasen etwas größere Abweichungen, im Mittel 2.66 s ($SD = 2.01$), festgestellt wurden. Die statistische Auswertung ergab entsprechend einen signifikanten Haupteffekt für „Fehlerfrei vs. Fehler“, $F(1, 22) = 4.32, p = .05$. Kein anderer Effekt wurde signifikant, Interaktionseffekte „Fehlerfrei vs. Fehler“ x „*commission* Fehler ja/nein“ $F(1, 22) = 1.19, p = .29$, „Fehlerblock“ x „Fehlerfrei vs. Fehler“ x „*commission* Fehler ja/nein“ $F(2, 44) = 1.76, p = .19$, alle übrigen $F < 1$.

Bei der Reaktionszeitaufgabe quittierten die Probanden das Verbindungssymbol durchschnittlich 1.46 s ($SD = 0.57$) nach dessen Erscheinen. Basierend auf der experimentellen Gruppeneinteilung ergab die Datenanalyse einen signifikanten Unterschied zwischen fehlerfreien Phasen und Fehlerphasen: In den Fehlerphasen fielen die Reaktionszeiten der Probanden durchweg höher aus ($M = 1.75$, $SD = 0.80$) als in den fehlerfreien Phasen ($M = 1.36$, $SD = 0.49$), Haupteffekt „Fehlerfrei vs. Fehler“ $F(1, 22) = 7.66$, $p = .01$. Darüber hinaus wurden keine statistisch bedeutsamen Effekte gefunden, Haupteffekt „Fehlerblock“ $F(2, 44) = 1.20$, $p = .31$, Interaktionseffekte „Fehlerblock“ x „Gruppe“, $F(2, 44) = 1.52$, $p = .23$, „Fehlerblock“ x „Fehlerfrei vs. Fehler“ x „Gruppe“ $F(2, 44) = 1.12$, $p = .334$, alle übrigen $F < 1$.

Bei einer Gruppeneinteilung nach „*commission* Fehler ja/nein“ zeigten Probanden die einen *commission* Fehler begingen, sobald eine Fehlfunktion von AFIRA eintrat, einen Leistungseinbruch bei der Reaktionszeitaufgabe. Demgegenüber gelang es Probanden, die keinen *commission* Fehler begingen, sowohl in fehlerfreien wie auch in Fehlerphasen ein annähernd konstantes Leistungsniveau beizubehalten (siehe Abbildung 21). Dies manifestierte sich in einem signifikanten Interaktionseffekt für „Fehlerfrei vs. Fehler“ x „*commission* Fehler ja/nein“ aus, $F(1, 22) = 4.42$, $p = .05$. Kein anderer Effekt wurde signifikant, Haupteffekt „Fehlerfrei vs. Fehler“ $F(1, 22) = 1.51$, $p = .23$, alle übrigen $F < 1$.

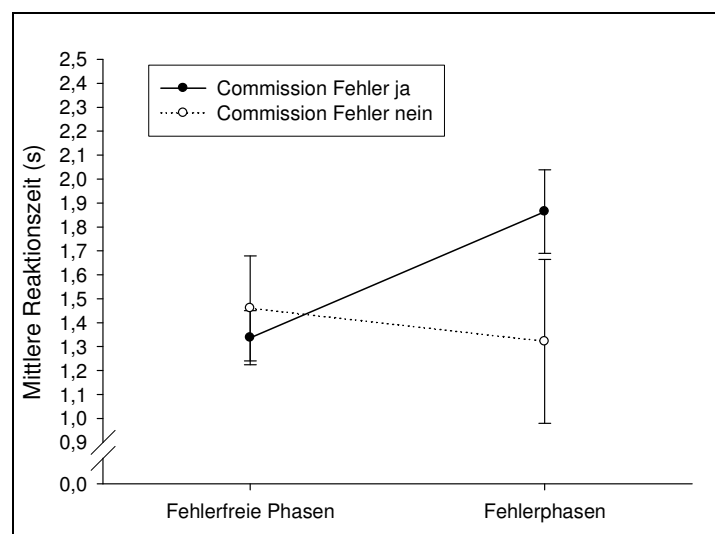


Abb. 21: Exp. I: Leistung in der Reaktionszeitaufgabe in fehlerfreien Phasen und Fehlerphasen (*commission* Fehler ja/nein)

3.4.7 Vertrauen in Automation und Selbstvertrauen

Das Vertrauen in Automation (CAMS bzw. AFIRA) sowie das Selbstvertrauen der Probanden wurde zu vier Zeitpunkten abgefragt: Die erste Erhebung erfolgte nach dem manuellen Training, die zweite Erhebung unmittelbar vor der Bedingungsmanipulation, die dritte nach der Bedingungsmanipulation und die vierte Erhebung nach Abschluss der Testphase.

Das Vertrauen in Automation wurde, wie der zeitliche Verlauf in Abbildung 22 erkennen lässt, von den Probanden der Erfahrungsgruppe zwar nach der Bedingungsmanipulation (Training 3) geringfügig niedriger eingeschätzt ($M = 3.50$, $SD = 1.17$) als von Probanden der Informationsgruppe ($M = 4.08$, $SD = 0.67$). Eine Auswertung der Daten mithilfe einer 2(Gruppe) x 4(Abfragezeitpunkt) ANOVA mit Messwiederholung zeigte aber weder einen Haupteffekt „Abfragezeitpunkt“, $F(3, 66) = 1.63$, $p = .19$, noch einen Haupteffekt „Gruppe“ oder einen Interaktionseffekt, jeweils $F < 1$.

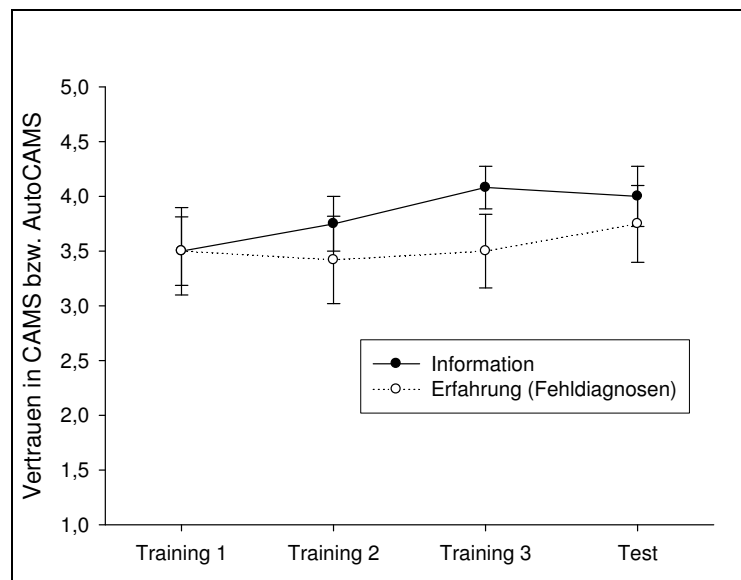


Abb. 22: Exp. I: Vertrauen in Automation (experimentelle Gruppen)

Mit Blick auf das Selbstvertrauen (siehe Abb. 23) ergab die Auswertung folgendes Bild: Während in der Erfahrungsgruppe das Selbstvertrauen über die ersten drei Messzeitpunkte, also bis nach der Bedingungsmanipulation abnahm, um dann zum Ende der Testphase wieder anzusteigen, zeigte sich in der Informationsgruppe ein entgegengesetztes Muster: Dort stieg das Selbstvertrauen bis nach der Bedingungsmanipulation an, um dann zum Ende der Testphase einen Einbruch zu erfahren. Eine 2 (Gruppe) x 4 (Abfragezeitpunkt) ANOVA mit Messwiederholung wies einen signifikanten Interaktionseffekt „Gruppe“ x „Abfragezeitpunkt“ aus, $F(3, 66) = 3.71$, $p = .02$, jedoch wieder keine Haupteffekte, „Gruppe“ $F < 1$, „Abfragezeitpunkt“ $F(3, 66) = 1.61$, $p = .95$.

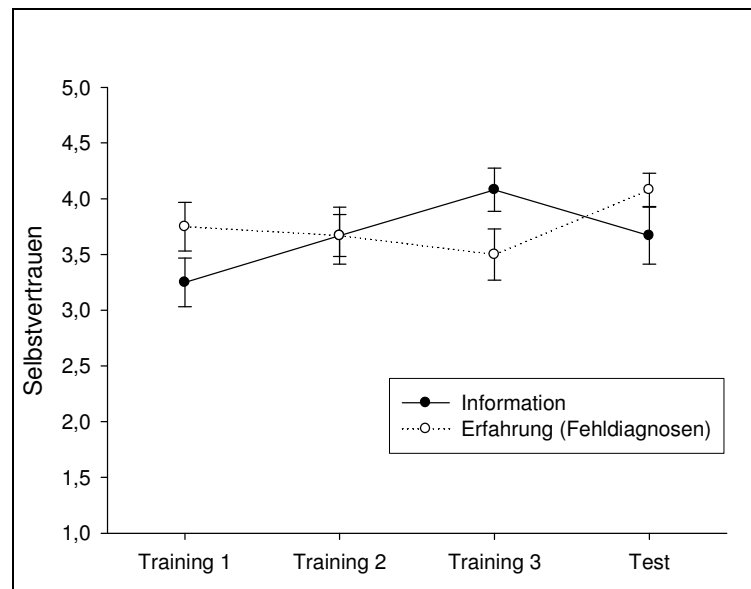


Abb. 23: Exp. I: Selbstvertrauen (experimentelle Gruppen)

Auch bezüglich der subjektiven Daten wurde geprüft, inwiefern sich Probanden, die einen *commission* Fehler begingen von jenen, die die AFIRA-Fehldiagnose entdeckten, unterscheiden. Die Auswertung erfolgte wiederum über 2(*commission* Fehler ja/nein) x 4(Abfragezeitpunkt) A-NOVAs mit Messwiederholung. Für das Vertrauen in Automation wurden weder die beiden Haupteffekte (für „*commission* Fehler ja/nein“ und „Abfragezeitpunkt“ $F < 1$) noch der Interaktionseffekt, $F(3, 66) = 1.35, p = .27$, signifikant. In Bezug auf das Selbstvertrauen zeigte sich das selbe Ergebnismuster, auch hier wurde kein Effekt signifikant, Haupteffekt „*commission* Fehler ja/nein“ $F(1, 22) = 3.55, p = .07$, Haupteffekt „Abfragezeitpunkt“ und Interaktionseffekt $F < 1$.

3.5 DISKUSSION EXPERIMENT I

Die Ergebnisse der vorgestellten Studie zeigen, dass *complacency* nicht nur in klassischen *monitoring*-Settings (Parasuraman et al., 1993) ein ernst zu nehmendes Problem darstellt, sondern auch im Kontext der Assistenzsystem-Nutzung. Alle Probanden zeigten in gewissem Umfang *complacency* gegenüber der Diagnosefunktion des Assistenzsystems, riefen also nicht alle Informationen, die zur Überprüfung derselben notwendig gewesen wären, ab. Zur Operationalisierung von *complacency* wurde, der Kritik von Moray (2003; Moray & Inagaki, 2000) Rechnung tragend, ein normatives Modell optimalen Informationssuchverhaltens erstellt und *complacency* erst bei einem – im Vergleich zu diesem – suboptimalen Informationssuchverhalten attestiert. Diese Operationalisierung ermöglichte es im Gegensatz zu früheren Studien (z.B. Bagheri & Jamieson, 2004a, b; Parasuraman et al., 1993; Prinzel et al., 2001), verschiedene Ausprägungen von *complacency* differenzieren zu können. Zudem wurde durch die direkte Erfassung die Grundvoraussetzung dafür erfüllt, *complacency* unabhängig von möglichen Verhaltenskonsequenzen im Sinne von *commission*

Fehlern untersuchen und mögliche Querbezüge der Phänomene analysieren zu können. Bislang findet sich ein ähnlich direkter Ansatz zur Erfassung des Automationsverifikationsverhaltens nur bei Skitka, Mosier und Burdick (2000), wenngleich dieser aufgrund der Künstlichkeit des Verifikationsprozederes letztlich nicht zu überzeugen vermag.

Im Folgenden werden die Ergebnisse im Detail nach den Hypothesen gegliedert vorgestellt und diskutiert. Gemäß Hypothese 1 wurde erwartet, dass Fehldiagnosen im Training *complacency* im späteren Umgang mit dem System reduzieren. Die Ergebnisse bestätigen diese Annahme. Probanden der Erfahrungsgruppe riefen einen deutlich höheren Anteil der zur Überprüfung erforderlichen Parameter ab. Damit zeigten sie eine geringere *complacency*-Ausprägung als Probanden der Informationsgruppe und das, obwohl auch diese darauf hingewiesen worden waren, dass Automationsfehler auftreten können. Dieser Befund passt zur Beobachtung von Skitka, Mosier, Burdick & Rosenblatt (2000), dass die bloße Information, dass Fehler auftreten könnten, *automation bias*-Effekte nicht in erwünschtem Maße zu reduzieren vermag. Dagegen vermochten nur 2 fehlerhafte von insgesamt 10 AFIRA-Diagnosen das Verhalten der Probanden für den weiteren Versuchsablauf maßgeblich zu verändern. Hier kann eine Parallele zu den Befunden von Lee und Moray (1992) sowie Dzindolet et al. (2003) gesehen werden: Selbst dann, wenn Automationsfehler nur seltene Ereignisse darstellen, kommt ihnen doch eine erhebliche Wirkung zu.

Keine Unterschiede zeigten sich dagegen zwischen den Gruppen mit Blick auf die Gesamtanzahl von Informationsabrufen und den Anteil relevanter Informationsabrufe während der Fehlerdiagnosephasen. Das bedeutet, dass alle Probanden in etwa gleichem Maße Informationen abriefen. Darüber, warum der Anteil relevanter Informationsquellen im Laufe der Zeit abnahm, kann nur gemutmaßt werden. Möglich wäre, dass sich darin die zunehmende Routine und evtl. auch Langeweile der Probanden bei der Fehlerdiagnose widerspiegelte: Während zu Beginn der Testphase derselbe relevante Parameter eventuell mehrfach abgerufen werden musste, genügte zum Ende hin eine einmalige Rückversicherung, sodass mehr Kapazität für den Abruf „irrelevanter“ Parameter zur Verfügung stand.

Festzuhalten bleibt jedoch, dass sich die Erfahrung von Fehldiagnosen positiv auf die gezielte Überprüfung der vorgeschlagenen Diagnosen auswirkte. Wichtig dabei ist, dass diese Leistungsverbesserung nicht auf Kosten der übrigen Aufgaben geschah. So wurde die Leistung in der prospektiven Gedächtnisaufgabe weder durch die experimentelle Manipulation noch durch das Vorliegen einer Fehlfunktion beeinflusst. Auch bei der Reaktionszeitaufgabe zeigte sich kein Effekt der experimentellen Gruppenzuteilung. Allerdings erwies sich die Leistung der Probanden in Fehlerphasen als schlechter im Vergleich zu den fehlerfreien Phasen, was der zusätzlichen Beanspruchung durch Fehlerdiagnose und -management geschuldet sein dürfte. Die Fehlererfahrung

im Training hatte jedoch auch hier keine Prioritätenverschiebung zur Folge und schlug sich nicht in einer Beanspruchungserhöhung nieder.

Zusammenfassend kann damit festgehalten werden, dass die Erfahrung von seltenen Fehldiagnosen im Training eine effektive Maßnahme zur Reduktion von *complacency*-Effekten gegenüber der Diagnosefunktion darstellen kann. Allerdings kann die Erfahrung von Fehlern *complacency*-Effekte nicht gänzlich vermeiden, sondern vermag diese nur zu reduzieren. Zudem demonstrieren die Ergebnisse, dass *complacency* nicht nur mit negativen Konsequenzen verbunden ist. Im Gegenteil, *complacency* kann, solange das betreffende automatisierte System zuverlässig arbeitet und korrekte Empfehlungen generiert, die Leistung sogar verbessern. In der vorliegenden Studie zeigt sich ein solcher Effekt darin, dass die Probanden der Informationsgruppe, die höhere *complacency*-Ausprägungen aufwiesen, zugleich signifikant weniger Zeit für die Fehlerdiagnose benötigten. Im Umkehrschluss legt dies die Vermutung nahe, dass Zeitdruck *complacency* erhöhen könnte. Dieses Ergebnis passt zu den früheren Befunden, dass *complacency* insbesondere in stark beanspruchenden *multi-task*-Settings auftritt, in denen mehrere Aufgaben parallel zu bearbeiten sind, und der Zeitdruck daher vergleichsweise hoch ausfällt (Parasuraman et al., 1993).

Gemäß Hypothese 2 wurde, basierend auf früheren Befunden von Muir und Moray (1996) erwartet, dass die Fehlererfahrung im Training sich in einer generellen Skepsis gegenüber dem System niederschlägt, und damit Probanden der Erfahrungsgruppe auch in den von AFIRA als fehlerfrei indizierten Phasen mehr Informationen über den Systemzustand abrufen als Probanden der Informationsgruppe. Entgegen der Erwartung zeigten sich jedoch keine Unterschiede zwischen den Gruppen, unabhängig davon, ob die Gesamtanzahl von Informationsabrufen oder der Anteil relevanter Informationsabrufe betrachtet wurde. Allerdings nahm letztgenannter – wie auch schon in den Fehlerdiagnosephasen – über die Zeit hinweg ab. Möglicherweise kann dies auch hier als ein Hinweis auf die zunehmende Routine bzw. Langeweile der Probanden gewertet werden.

Festzuhalten ist damit, dass sich die Erfahrung von Fehldiagnosen ausschließlich auf das Verhalten der Probanden gegenüber dieser funktionalen Komponente von AFIRA auswirkte; das Überwachungsverhalten in Bezug auf die Alarmfunktion blieb dagegen unbeeinflusst. Offenbar differenzieren Probanden zwischen den einzelnen funktionalen Aspekten eines Systems und können somit ihrem Überprüfungsverhalten sehr wohl eine hohe funktionale Spezifität im Sinne von Lee und See (2004) verleihen. Der Befund von Muir und Moray (1996), dass Fehler einer bestimmten automatisierten Funktion sich auch auf die Haltung von Operateuren gegenüber anderen Funktionen desselben Systems auswirken, konnte somit nicht repliziert werden. Stattdessen zeitigten Automationsfehler in der hier berichteten Studie nur einen sehr spezifischen Effekt.

Ein vergleichbar differenzieller Effekt von Systemfehlern wird in der Literatur für einfache, binäre Alarmsysteme diskutiert. Meyer (2004) unterscheidet zwei Verhaltensweisen gegenüber binären Alarmsystemen: Während *reliance* beschreibt, inwiefern ein Operateur sich auf das Alarmsystem in Abwesenheit eines Alarms verlässt, bezieht sich *compliance* darauf, in welchem Umfang der Operateur sich auf das System verlässt, wenn ein Alarm vorliegt. Der Versuch, die Befunde der hier berichteten Studie in diese Begriffsstruktur zu integrieren, führt zu dem Ergebnis, dass falsche Diagnosen eines Assistenzsystems nur die *compliance* gegenüber der Diagnosefunktion (als Sonderform der *compliance*), nicht aber die *reliance* der Operateure gegenüber der Alarmfunktion des Systems mindern.

Die dritte Hypothese bezog sich auf die Fehldiagnose von AFIRA: Erwartet wurde, dass Probanden, die Fehldiagnosen bereits im Training erfahren hatten, bei der Fehldiagnose in der Testphase weniger *commission* Fehler begehen als Probanden, die nur über die Möglichkeit von Automationsfehlern informiert worden waren. Die Ergebnisse bestätigen diese Annahme nicht. Insgesamt begingen nur 5 der 24 Probanden einen *commission* Fehler, folgten also der falschen Diagnose, wobei diese zu gleichen Teilen aus den beiden experimentellen Gruppen stammten. Der geringe Anteil von *commission* Fehlern überrascht, zumal in früheren Studien sehr viel höhere Fehlerraten (54 bis 100 %, Mosier et al., 1998; Mosier et al., 2001; Skitka, Mosier, Burdick & Rosenblatt, 2000) beobachtet wurden. Eine mögliche Erklärung hierfür wäre, dass die Fehldiagnose relativ einfach zu entdecken war – bereits ein Parameter genügte, um die Diagnose als falsch zu überführen. Zudem besteht ein genereller Unterschied zwischen der Verifikation und Falsifikation von Diagnosen: Während eine gegebene Diagnose ggf. bereits durch das Auffinden einer widersprechenden Information falsifiziert werden kann, ist eine vollständige Verifikation für den Fall, dass die Diagnose korrekt ist, deutlich aufwändiger. Wäre die Fehldiagnose dagegen schwieriger zu entdecken gewesen, wäre der Anteil von *commission* Fehlern mutmaßlich höher ausgefallen. Allerdings wurden von allen Probanden, also auch jenen fünf, die der falschen Diagnose folgten, Informationen abgerufen, die mit der Diagnose nicht vereinbar waren. Umso mehr stellt sich die Frage, wodurch sich jene fünf Probanden dann gegenüber den übrigen Probanden auszeichneten. Diese Frage war Gegenstand der vierten Hypothese:

Davon ausgehend, dass *commission* Fehler eine mögliche Folge von *complacency* darstellen, wurde erwartet, dass Probanden, die einen solchen Fehler begehen, höhere *complacency*-Ausprägungen zeigen im Vergleich zu Probanden, die keinen *commission* Fehler begehen. Mit Blick auf die ersten neun Fehlfunktionen, bei denen AFIRA korrekte Diagnosen lieferte, bestätigten die Ergebnisse diese Annahme: Probanden, die der Fehldiagnose folgten, wiesen deutlich höhere *complacency*-Ausprägungen auf, überprüften die AFIRA-Diagnosen also weit weniger umfänglich.

Zudem zeichneten sich die Probanden, die einen *commission* Fehler vermeiden konnten, generell durch ein aktiveres Informationssuchverhalten in den Fehlerdiagnosephasen aus: Im Vergleich zu Probanden, die einen *commission* Fehler begingen, nahmen sie sich fast doppelt so viel Zeit für die Fehlerdiagnose und riefen fast doppelt so oft Informationen vor dem Abschicken eines Reparaturauftrages ab. In den fehlerfreien Phasen unterschied sich das Informationssuchverhalten dagegen nicht zwischen Probanden, die einen *commission* Fehler begingen und solchen, die ihn vermeiden konnten.

Mit Blick auf die Zusatzaufgaben zeigte sich folgendes Bild, wenn die Gruppeneinteilung nach „*commission* Fehler ja/nein“ erfolgte: Bei der prospektiven Gedächtnisaufgabe fielen die Leistungen aller Probanden in den Fehlerphasen schlechter aus als in den fehlerfreien Phasen. Bei der Reaktionszeitaufgabe zeigten in den Fehlerphasen dagegen nur Probanden, die einen *commission* Fehler begingen, schlechtere Leistungen. Probanden, die den Fehler vermeiden konnten, wurden durch das Auftreten der Fehlfunktion offenbar weniger stark beeinträchtigt. Sie zeigten unabhängig davon, ob ein Fehler vorlag oder nicht, nahezu konstante Reaktionszeiten. Dieser Effekt könnte als ein Hinweis darauf gewertet werden, dass Probanden, die *commission* Fehler begingen, durch die Fehlfunktionen insgesamt mehr beansprucht wurden und dies durch eine Leistungsabnahme bei den Zusatzaufgaben kompensierten. Gegen diese Interpretation spricht jedoch die Tatsache, dass die Probanden, die einen solchen Fehler begingen, weniger Aufwand in die Informationssuche in Fehlerdiagnosen investierten, was allenfalls mit einer geringeren nicht aber höheren Beanspruchung in Verbindung gebracht werden könnte.

Zusammenfassend lässt sich festhalten, dass Probanden, die einen *commission* Fehler begehen, sich durch höhere *complacency*-Ausprägungen und eine weniger aktive Informationssuche in Fehlerdiagnosephasen auszeichnen.

Gemäß Hypothese 5 wurde angenommen, dass *complacency* zu einem Verlust der manuellen Fehlerdiagnose-Kompetenz beitragen kann, was durch Befunde von Lorenz et al. (2002) nahe gelegt wird. Erwartet wurde, dass Probanden, die im Training Fehldiagnosen erfahren hatten, beim Ausfall des Assistenzsystems am Ende der Testphase, geringere Fertigkeitsverluste im Vergleich zum Training aufweisen. Allerdings zeigen die Ergebnisse keinen solchen Effekt. Im Gegenteil, statt eines Fertigkeitsverlust wurde für beide Gruppen ein deutlicher Trainingseffekt beobachtet. Sowohl der Informations- als auch der Erfahrungsgruppe gelang es, ihre Fehlerdiagnose- und Fehlermanagementfertigkeiten über den Verlauf der Testphase hinweg im Vergleich zum Training am ersten Tag deutlich zu verbessern. Offenbar war der begrenzte Umfang des Trainings in der vorliegenden Studie nicht hinreichend, um weitere Trainingszugewinne in der Testphase zu vermeiden und mögliche Fertigkeitsverluste zu Tage treten zu lassen. Alternativ könnte

auch argumentiert werden, dass die Testphase letztlich zu kurz war, um überhaupt einen Fertigungsverlust entstehen zu lassen. Ergebnisse von Sauer et al. (2000) legen eine solche Interpretation nahe: Sie berichten, dass Probanden selbst nach acht Monaten ohne zwischenzeitliche Praxis mit CAMS Fehlfunktionen nahezu genauso schnell diagnostizieren konnten wie in der ersten Trainingseinheit.

Hypothese 6 zum Vertrauen in Automation konnte nicht bestätigt werden. Probanden, die im Training Fehldiagnosen erfahren hatten, unterschieden sich diesbezüglich nicht von Probanden, die im Training nur über mögliche Automationsfehler informiert worden waren. Erwartet worden war, dass sich die Erfahrung von Automationsfehlern mindernd auf das Vertrauen in Automation auswirkt. Die Gruppenmittelwerte weisen zwar tendenziell in diese Richtung, es wurde jedoch kein signifikanter Unterschied festgestellt. Dieser Befund legt zunächst nahe, dass die Erfahrung von Fehldiagnosen im Training das Vertrauen in Automation nicht oder nur marginal beeinflusst. Dies verwundert, da die gewählten Fehler „einfache“, nicht nachvollziehbare Fehler darstellten und zudem mit Kosten verbunden waren. Die zentralen Voraussetzungen für den mehrfach replizierten vertrauensmindernden Effekt von Automationsfehlern (Dzindolet et al., 2003; Madhavan et al., 2006; Masalonis et al., 1998) waren folglich gegeben. Vor diesem Hintergrund stellt sich die Frage, ob die gewählte Operationalisierung geeignet war, um eventuell bestehende Vertrauensunterschiede angemessen abbilden zu können. Möglich wäre, dass die abgeforderte Einschätzung von „Vertrauen in AFIRA“ zu unspezifisch mit Blick auf die unterschiedlichen funktionalen Komponenten von AFIRA gewählt wurde – zumal sich im Verhalten der Probanden eine hohe diesbezügliche Spezifität dokumentierte. Während die Alarmfunktion von AFIRA stets zuverlässig arbeitete und auch die generierten Handlungsempfehlungen immer zur gestellten Diagnose korrespondierten, war es nur die Diagnosefunktion, die für die Erfahrungsgruppe in 2 von 10 Fällen fehlerhafte Ergebnisse produzierte. Nicht auszuschließen ist, dass in die Vertrauensurteile der Probanden die hohe Zuverlässigkeit des Systems in den übrigen Funktionen einfluss und zwei Fehldiagnosen im Training das allgemeine Vertrauen in AutoCAMS nicht zu mindern vermochten.

Zusammenfassend lassen sich die folgenden Schlussfolgerungen aus Experiment I ziehen: Erstens erwies sich der experimentelle Ansatz, *complacency* im Umgang mit einem Entscheidungsassistenzsystem unabhängig von möglichen Verhaltenskonsequenzen zu operationalisieren, als erfolgreich und kann in dieser Form für weitere Untersuchungen des Phänomens als geeignet erachtet werden.

Zweitens zeigen die Ergebnisse, dass Kosten- wie auch Nutzeneffekte zu berücksichtigen sind, wenn es gilt, die zugrunde liegenden Determinanten von *complacency* zu verstehen. Eine hö-

here *complacency*-Ausprägung spiegelt sich in geringeren Diagnosezeiten wider, was aus Perspektive des Operators insbesondere in Situationen mit hohem Zeitdruck als klarer Leistungsvorteil gesehen werden mag. Das Blatt wendet sich jedoch, sobald die vom Assistenzsystem generierten Direktiven nicht mehr korrekt sind: Dann sind mit hohen *complacency*-Ausprägungen deutlich Nachteile verbunden, die sich in einem erhöhten Risiko für *commission* Fehler niederschlagen. Allerdings zeigten alle Probanden *complacency* in gewissem Umfang, während nur fünf von ihnen ein *commission* Fehler unterlief. Dies weist darauf hin, dass *commission* Fehler eine mögliche, jedoch keinesfalls notwendige Folge von *complacency* darstellen. Offen bleibt jedoch, ob dies auch für, im Vergleich zu dieser Studie, komplexere und weniger leicht zu entdeckende Diagnosefehler gilt.

Drittens stellen Trainings, die Fehldiagnosen eines Assistenzsystems einschließen und diese für Operateure unerwartet und direkt erfahrbar machen, eine geeignete Maßnahme zur Reduktion von *complacency* gegenüber der Diagnosefunktion dar. Nicht zu erwarten ist jedoch, dass sich damit *complacency*-Effekte gänzlich vermeiden lassen. Zum einen bewirkte in der vorliegenden Studie die Erfahrung von Fehldiagnosen zwar eine Intensivierung des Überprüfungsverhaltens, führte aber nicht zu einer vollständigen „optimalen“ Überprüfung der Diagnosen. Zum anderen wirkte sich die Erfahrung von Fehldiagnosen spezifisch aus: Nur das Überprüfungsverhalten gegenüber der Diagnosefunktion, nicht aber die Überwachung der Alarmfunktion, konnte durch die experimentelle Manipulation verbessert werden. Dies wirft die Frage auf, ob sich auch andere Fehlertypen nur spezifisch auf das Verhalten des Operators auswirken bzw. ob es, ähnlich der Unterscheidung von *reliance* und *compliance* bei der Nutzung binärer Alarmsysteme (Meyer, 2004), auch bei der Nutzung von Assistenzsystemen qualitativ verschiedene Formen des „sich Verlassens“ auf das System zu unterscheiden gilt. Möglicherweise fördern Automationsfehler zwar die Wachsamkeit des Operators bezüglich der als fehlerhaft kennen gelernten Teilfunktion, aber nicht für andere Teilfunktionen der Automation, für die jedoch ebenso wenig eine 100-prozentige Zuverlässigkeit garantiert werden kann. Sollte dem so sein, müssten sich etwa im Training erfahrene Systemausfälle eines Assistenzsystems zwar in einer vermehrten Überprüfung der Alarmfunktion niederschlagen, nicht aber auf die Überprüfung der Diagnosefunktion auswirken.

Dies zu untersuchen, war unter anderem Ziel des zweiten, im Folgenden beschriebenen Experiments.

4 EXPERIMENT II: ZUM EINFLUSS VON SYSTEMAUSFÄLLEN

4.1 VORBEMERKUNG

Die im Folgenden berichtete Studie wurde an der Technischen Universität Berlin durchgeführt. Die Datenerhebung erfolgte dabei im Rahmen einer von der Autorin dieser Arbeit betreuten Diplomarbeit, wobei ein Teil der erhobenen Daten bereits für die Diplomarbeit herangezogen wurden (vgl. auch Elepfandt, Bahner & Manzey, 2007). Die hier berichteten Fragestellungen umfassen in Teilen jene der Diplomarbeit, gehen jedoch deutlich über diese hinaus.

4.2 EINLEITUNG

In Experiment I wurde gezeigt, dass durch die Erfahrung von Fehldiagnosen eines Assistenzsystems im Training *complacency*-Effekte reduziert werden können: Probanden, die Fehler der Diagnosefunktion erfahren hatten, überprüften im späteren Umgang mit dem System die vom Assistenzsystem vorgeschlagenen Diagnosen sorgfältiger als Probanden, die nur über die Eventualität von Automationsfehlern in Kenntnis gesetzt worden waren. Zudem wurde gezeigt, dass *automation bias* in Form eines *commission* Fehlers eine mögliche, wenn auch nicht notwendige, Folge von *complacency* gegenüber der Diagnosefunktion darstellt. Dagegen zeigten sich keine Unterschiede im Informationssuchverhalten der beiden experimentellen Gruppen in jenen Phasen, in denen vom Assistenzsystem kein zu behebender Fehler angezeigt wurde. Die Fehlererfahrung wirkte somit spezifisch und beeinflusste nur das Verhalten der Probanden gegenüber der Diagnosefunktion des Assistenzsystems, nicht aber gegenüber der Alarmfunktion. In dieser Spezifität kann eine Analogie zu der von Meyer (2004) vorgeschlagenen Unterscheidung von *reliance* und *compliance* im Kontext binärer Alarmsysteme gesehen werden. Wie bereits erwähnt, kennzeichnet *reliance* das Ausmaß in dem sich ein Operateur bei Abwesenheit eines Alarms darauf verlässt, dass tatsächlich kein kritischer Systemzustand vorliegt. *Compliance* beschreibt dagegen das Ausmaß in dem sich der Operateur bei Vorliegen eines Alarms darauf verlässt, dass auch faktisch ein kritischer Zustand gegeben ist. Dabei wird konzeptuell davon ausgegangen, dass *reliance* insbesondere durch Auslassungsfehler des Alarmsystems (*automation misses*) beeinflusst wird: Wenn ein Alarmsystem wiederholt kritische Situationen nicht meldet, wird sich dies in einer *reliance*-Minderung niederschlagen, d.h. der Operateur wird in Phasen, in denen kein Alarm gegeben ist, vermehrt prüfen, ob tatsächlich keine kritische Situation vorliegt. Demgegenüber wird angenommen, dass *compliance* vor allem durch eine hohe Rate falscher Alarme beeinträchtigt wird und Operateure Alarme weniger beachten bzw. im Extremfall diese sogar gänzlich ignorieren (*cry wolf*-Phänomen, vgl. auch Dixon &

Wickens, 2006). Gegenstand mehrerer aktueller Studien ist die Frage, inwiefern die beiden Phänomene tatsächlich, wie konzeptuell angenommen, voneinander unabhängig sind, und sich somit falsche Alarme nur auf *compliance* und Auslassungsfehler ausschließlich auf *reliance* auswirken. Die bisher heterogene Befundlage lässt hier jedoch noch keinen eindeutigen Schluss zu (Dixon & Wickens, 2006; Dixon, Wickens & McCarley, 2007). Auch ist die Frage möglicher negativer Folgen von unangemessener *reliance* bzw. *compliance* empirisch bislang ungeklärt. Theoretisch liegt es jedoch nahe, anzunehmen, dass sich diese, sofern die Automation fehlerhaft arbeitet, in Form der von Mosier und Skitka (1996) als *automation bias* zusammengefassten Fehlertypen manifestieren könnten: Im Falle von *overreliance* können *omission* Fehler die Folge sein, im Falle von *overcompliance* wären dagegen *commission* Fehler zu erwarten.

Übertragen auf die vorliegende Arbeit bzw. den Kontext der Assistenzsystemnutzung, kann *overreliance* als Analogon zu *complacency* gegenüber der Alarmfunktion (mangelnde Überwachung) und *overcompliance* als Analogon zu *complacency* gegenüber der Diagnosefunktion (mangelnde Überprüfung) interpretiert werden. Die Ergebnisse des ersten Experiments zeigen, dass es auch bei der Nutzung von Assistenzsystemen beide Aspekte zu unterscheiden gilt, wobei Fehldiagnosen zwar den *compliance*-Aspekt, nicht jedoch den *reliance*-Aspekt beeinflussten. Dieser Befund kann als erster Hinweis darauf gedeutet werden, dass die beiden Aspekte voneinander unabhängig sein könnten.

Im Falle der Unabhängigkeit müssten Auslassungsfehler des Assistenzsystems zu einer *reliance*-Minderung und einer Minderung von *omission* Fehlern führen, während sich kein Einfluss auf die *compliance*-Ausprägung gegenüber der Diagnosefunktion und die Anzahl von *commission* Fehlern zeigen dürfte. Dies zu prüfen ist Ziel des zweiten Experiments. Untersucht werden soll, welche Effekte mit einem Training verbunden sind, das auf die Erfahrung von Systemausfällen, und damit von Auslassungsfehlern der Alarmfunktion des Systems, anstelle von Fehlern der Diagnosefunktion ausgerichtet ist. Geprüft werden soll, ob auch ein solches Training wiederum ausschließlich spezifisch wirkt. In diesem Fall würde das Training zwar zu einem verbesserten Informationssuchverhalten in vermeintlich fehlerfreien Phasen führen und möglicherweise das Risiko von *omission* Fehlern reduzieren. Nicht reduziert würden jedoch *complacency*-Effekte bezüglich der vom Assistenzsystem gestellten Diagnosen sowie mögliche *commission* Fehler.

Die Untersuchung soll damit einen Beitrag zur weiteren Aufklärung des Zusammenspiels von allgemeinem Informationssuchverhalten, *complacency* und *automation bias* leisten. Das experimentelle Design orientiert sich dabei mit dem Ziel einer möglichst direkten Vergleichbarkeit weitgehend am ersten Experiment: Wieder sollen zwei experimentelle Gruppen miteinander verglichen werden: Eine Erfahrungsgruppe, die im Training die Erfahrung von Ausfällen des

Assistenzsystems macht und eine Informationsgruppe, die nur über die Möglichkeit, dass Automationsfehler auftreten können, informiert wird. In einer nachfolgenden Testphase sollen dann wieder zunächst neun CAMS-Fehlfunktionen auftreten, die durch das Assistenzsystem korrekt gemeldet und diagnostiziert werden. Am Ende der Testphase werden dann zwei Systemausfälle simuliert. Im Unterschied zu Experiment I, in dem nur die Diagnosefunktion ausfiel, das Vorliegen einer Fehlfunktion jedoch durch den Masteralarm korrekt angezeigt wurde, soll jetzt weder AFIRA noch der Masteralarm einen Hinweis auf das Vorliegen einer Fehlfunktion liefern, sodass die Probanden nur durch eigene aktive Informationssuche die beiden Fehler entdecken können. Werden diese nicht rechtzeitig, d.h. vor Überschreiten eines kritischen Wertebereichs, entdeckt, wird dies als *omission* Fehler gewertet. Nach den beiden Systemausfällen soll auch in Experiment II Gelegenheit für potenzielle *commission* Fehler gegeben werden. Hierzu wird vor Abschluss des Experiments eine Fehldiagnose vom Assistenzsystem generiert. Allerdings soll diese, im Vergleich zu Experiment I, schwieriger zu falsifizieren sein, um zu prüfen, ob in diesem Fall die *commission*-Fehlerrate höher ausfällt und sich im Übrigen auch dann die Befunde aus Experiment I replizieren lassen. Ein letzter wesentlicher Unterschied zu Experiment I besteht in der Operationalisierung des Vertrauens in Automation: Dieses soll differenziert für die einzelnen Teilfunktionen des Assistenzsystems (Alarm, Diagnose, Vorschläge für das Fehlermanagement) erfasst werden, um so einem potenziell spezifischen Effekt der Fehlererfahrung Rechnung tragen zu können. Zudem soll die Abfrage indirekt über eine Einschätzung der Zuverlässigkeit der Teilfunktionen erfolgen, da anzunehmen ist, dass diese für die Probanden einfacher, da konkreter, zu bewerten ist als das Vertrauen in die jeweilige Teilfunktion. Die Einschätzung der Zuverlässigkeit ist zwar nicht eins zu eins mit dem Vertrauen gleichzusetzen, es kann jedoch davon ausgegangen werden, dass letzteres maßgeblich durch ersteres bestimmt wird (z.B. Muir & Moray, 1996).

4.3 HYPOTHESEN EXPERIMENT II

Mit Blick auf die experimentelle Bedingungsmanipulation werden folgende Effekte erwartet:

Hypothese 1: Probanden, die im Training Ausfälle eines Assistenzsystems im Sinne von Auslassungsfehlern erfahren, rufen in vom Assistenzsystem als fehlerfrei angezeigten Phasen **mehr Informationen über den Systemzustand ab** als Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert worden sind.

Hypothese 2: Probanden, die im Training Ausfälle eines Assistenzsystems im Sinne von Auslassungsfehlern erfahren haben, begehen in der späteren Arbeit mit dem System **weniger *omission* Fehler** als Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert worden sind.

Hypothese 3: Probanden, die im Training Ausfälle eines Assistenzsystems im Sinne von Auslassungsfehlern erfahren haben, **unterscheiden sich nicht bezüglich ihrer *complacency*-Ausprägung gegenüber der Diagnosefunktion** von Probanden, die nur über die Möglichkeit von Automationsfehlern informiert worden sind.

Hypothese 4: Probanden, die im Training Ausfälle eines Assistenzsystems im Sinne von Auslassungsfehlern erfahren haben, **unterscheiden sich nicht bezüglich des Begehens von *commission* Fehlern** von Probanden, die nur über die Möglichkeit von Automationsfehlern informiert worden sind.

Zum Zusammenhang von *automation bias*, *complacency* und dem Informationssuchverhalten in vermeintlich fehlerfreien Phasen werden die Hypothesen 5 und 6 aufgestellt:

Hypothese 5: Probanden, die einen ***commission* Fehler** begehen, zeigen **höhere *complacency*-Ausprägungen gegenüber der Diagnosefunktion** im Vergleich zu Probanden, die keinen *commission* Fehler begehen.

Hypothese 6: Probanden, die ***omission* Fehler** begehen, rufen in den vom Assistenzsystem als fehlerfrei angezeigten Phasen **weniger Informationen über den Systemzustand** ab als Probanden, die keine(n) *omission* Fehler begehen.

Als subjektive Komponente von *complacency* und *automation bias* wird die wahrgenommene Zuverlässigkeit der verschiedenen Teilfunktionen des Assistenzsystems erfasst und folgende Hypothese aufgestellt:

Hypothese 7: Probanden, die im Training Ausfälle eines Assistenzsystems im Sinne von Auslassungsfehlern erfahren haben, schätzen die **Zuverlässigkeit der Alarmfunktion** des Assistenzsystems geringer ein als Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert worden sind.

4.4 METHODE EXPERIMENT II

4.4.1 Stichprobe

Insgesamt nahmen $N = 24$ Studierende der Ingenieurwissenschaften der Technischen Universität Berlin an der Studie teil. Die Probanden (4 weiblich, 20 männlich) wurden über Bekanntgabe in Vorlesungen und E-Mail-Verteilern gewonnen. Ein männlicher Proband musste von der Datenauswertung ausgeschlossen werden, da dieser mehrfach gegen die Instruktionen verstoßen hatte (trotz wiederholter Aufforderung fast durchgängiger Verzicht auf die Aktivierung des Feh-

lerbuttons trotz erkannter Fehlfunktion; vgl. hierzu Kap. 4.4.3 in dem die Aufgaben der Probanden genauer beschrieben sind). Die verbleibenden 23 Probanden waren zwischen 21 und 29 Jahren alt ($M = 24.33$, $SD = 1.85$) und verfügten nicht über Vorerfahrungen mit der genutzten Versuchsumgebung. Als Aufwandsentschädigung erhielten sie nach der Teilnahme 40 Euro.

4.4.2 Versuchsplan

Auch dieser Studie lag ein einfaktorieller (Fehlererfahrung vs. Fehlerinformation) Versuchsplan zugrunde, wobei die beiden Faktorstufen wiederum in verschiedenen Probandengruppen realisiert wurden (*between Design*).

Um sicherzustellen, dass die experimentellen Gruppen sich nicht von vornherein in zentralen abhängigen Variablen unterscheiden, erfolgte die Gruppeneinteilung nach dem ersten Untersuchungstag basierend auf den folgenden Variablen:

- Manuelle Fehler*detektionszeit*: Berechnet wurde erstens über alle zehn Fehlfunktionen hinweg, zweitens über die beiden auftretenden N₂ Lecks und drittens über die beiden auftretenden O₂ Blockierungen des Ventils die mittlere Dauer (s) zwischen Auftreten einer Fehlfunktion und Aktivierung des Fehlermodus. N₂ Lecks und O₂ Blockierungen wurden gesondert betrachtet, da für diese in der Testphase Systemausfälle simuliert wurden, sodass dort die Fehlerdetektion, -identifikation und -behebung erneut rein manuell erfolgen musste.
- Manuelle Fehler*identifikationszeit*: Ermittelt wurde diese über alle zehn Fehlfunktionen als mittlere Dauer (s) zwischen Aktivierung des Fehlermodus und Absenden des richtigen Reparaturauftrags. Zudem wurde auch hier die mittlere Diagnosezeit für N₂ Lecks und O₂ Blockierungen des Ventils gesondert ermittelt.
- Informationssuchverhalten in den grünen Phasen: Analysiert wurde, wie viele Informationsabrufe in den jeweils letzten 120 s vor Auftreten einer Fehlfunktion („fehlerfreie Phasen“) durchschnittlich pro Minute abgerufen wurden; diese Eingrenzung der fehlerfreien Phase wurde gewählt, da die Probanden unterschiedlich lange brauchten um das System nach auftreten eines Fehlers wieder in einen fehlerfreien Zustand zu bringen. Für die letzten 120 s vor Auftreten des nächsten Fehlers gelang es jedoch allen Probanden einen fehlerfreien Zustand hergestellt zu haben, sodass dieses Zeitfenster eine bestmögliche interindividuelle Vergleichbarkeit gewährleistete.
- Selbsteinschätzung der Fehlerdetektions-, Fehlerdiagnose- und Fehlermanagementleistung (Fragebogenitems 6-8, vgl. Kap. 4.4.4.1).

- Wahrgenommene Zuverlässigkeit des Stickstoff- und Sauerstoffsubsystems (Fragebogenitems 9-10, vgl. Kap. 4.4.4.1) als Indikator für das Vertrauen in Automation.

Die Mittelwerte der beiden experimentellen Gruppen nach der Einteilung sind Anhang H zu entnehmen.

Die Bedingungsmanipulation wurde am zweiten Untersuchungstag während des AutoCAMS-Trainings (siehe Kap. 4.4.5.2) realisiert. Alle Probanden erhielten zu Beginn des Trainings die Information, dass Fehler des Assistenzsystems vorkommen können und daher die Parameter weiterhin überwacht werden sollten. Im Rahmen der Trainingseinheit wurde für die Hälfte der Probanden ($n = 12$; da eine Person im nachhinein ausgeschlossen werden musste verblieben für die Datenauswertung $n = 11$) bei zwei der zehn Fehlfunktionen ein Ausfall von AFIRA simuliert, sodass die Fehlerdetektion, -diagnose und das Fehlermanagement rein manuell erfolgen mussten (**Erfahrungsgruppe**). Auch diese Automationsfehler waren wie in Experiment I für die Probanden weder vorhersehbar noch nachvollziehbar und zudem mit hohen Kosten verbunden, da die Fehlerdetektion, -diagnose und -behebung rein manuell erfolgen musste. Auch hier wurden somit die Voraussetzungen für einen vertrauensmindernden Effekt erfüllt. Die andere Hälfte der Probanden ($n = 12$) arbeitete durchgehend mit einem zuverlässig arbeitenden Assistenzsystem, das alle Fehlfunktion entdeckte und korrekt diagnostizierte. Damit diente dieser Gruppe die Information im Rahmen der Instruktion als einziger Hinweis auf die nicht 100-prozentige Reliabilität des Systems (**Informationsgruppe**).

Die abhängigen Variablen wurden in einer auf das AutoCAMS-Training folgenden Testphase erfasst und werden – der besseren Nachvollziehbarkeit wegen – im Anschluss an die Beschreibung der Durchführung in einem gesonderten Kapitel (4.4.6) beschrieben.

4.4.3 Anpassungen der Experimentalumgebung AutoCAMS

Da auch in dieser Studie dieselbe Versuchsumgebung eingesetzt wurde, die bereits in Experiment I Verwendung fand, soll im Folgenden ausschließlich auf abweichende Aspekte eingegangen werden.

Verbunden mit dem Ziel, den Einfluss von Ausfällen der Alarmfunktion des Assistenzsystems AFIRA zu untersuchen, galt es, unabhängige Hinweise auf vorliegende Fehlfunktionen zu vermeiden. Dies war in Studie I nicht der Fall. Dort wurden auftretende Fehlfunktionen stets und unabhängig von AFIRA auch über den Masteralarm angezeigt. Auch bei den beiden letzten Fehlern der Testphase, für die in Studie I ein Ausfall von AFIRA simuliert wurde, erhielten die Probanden somit einen Hinweis auf das Vorliegen einer Fehlfunktion. Eine solche Realisierung er-

laubt es jedoch nicht, die erwarteten bedingungsabhängigen Veränderungen in der manuellen Fehlerdetektion abzubilden, da diese immer systemseitig unterstützt bleibt. Um dies dennoch zu ermöglichen, wurden für Studie II zwei Änderungen der Versuchsumgebung vorgenommen:

Zum einen wurden Masteralarm und AFIRA gekoppelt, sodass der Masteralarm nur dann auf rot wechselte, wenn auch AFIRA eine Fehlermeldung anzeigte. Bei simulierten Systemausfällen dagegen erschien weder eine AFIRA-Meldung noch wurde eine auftretende Fehlfunktion durch den Masteralarm indiziert. Analog stand während des CAMS-Trainings (ohne AFIRA) kein Masteralarm zur Verfügung (Anhang I zeigt eine Abbildung der geänderten Benutzeroberfläche), sodass neben der Fehlerdiagnose und -behebung auch die Fehlerdetektion manuell erfolgen musste.

Um die Fehlerdetektionsleistung quantifizieren zu können, wurde zum anderen ein Fehlermodus-Button eingeführt (siehe Abb. 24). In fehlerfreien Phasen erschien dieser in grüner Farbe mit der Beschriftung „OK“. In fehlerfreien Phasen erschien dieser in grüner Farbe mit der Beschriftung „OK“.

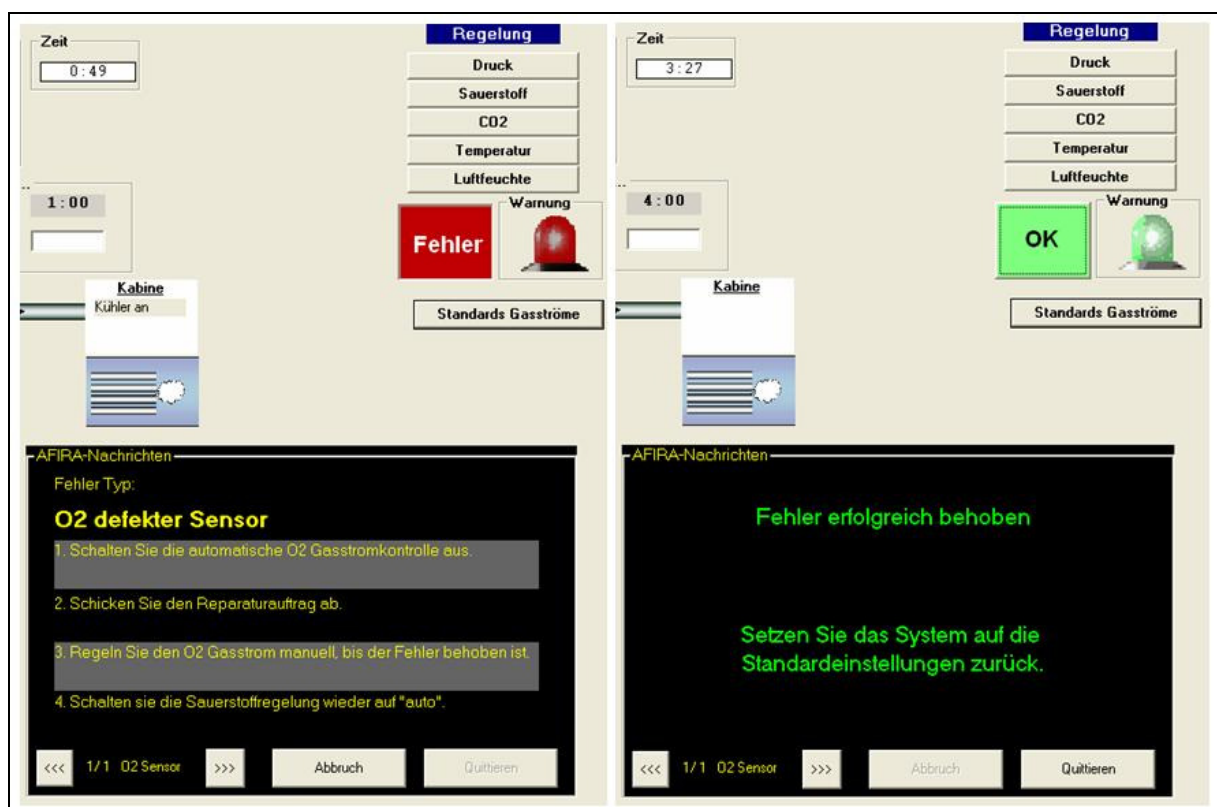


Abb. 24: AutoCAMS mit Fehlermodus-Button (links: Fehlermodus aktiviert – „Fehler“, rechts: Fehlermodus deaktiviert – „OK“)

Sobald ein Proband jedoch eine Fehlfunktion im System festgestellt hatte, war es seine Aufgabe, so schnell wie möglich per Mausklick auf den Button den Fehlermodus zu aktivieren. Der Button wechselte dann von grün auf rot und zeigte die Beschriftung „Fehler“. Erst danach sollte die eigentliche Fehleridentifikation stattfinden. Durch diese Trennung wurde es bei Systemausfällen

len möglich, die Fehlerdetektionszeit, aus der sich ggf. *omission* Fehler ablesen lassen, und die Fehleridentifikationszeit als Indikator für Fertigkeitsverluste empirisch zu trennen. Nach erfolgreicher Reparatur der Fehlfunktion war es Aufgabe der Probanden, den Fehlermodus wiederum per Mausklick auf den Button zu deaktivieren, sodass selbiger wieder in grüner Farbe „OK“ anzeigte.

Eine weitere Änderung betraf die Bestimmung der Auftretenszeitpunkte von O₂-Sensordefekten. In Studie I wurde im Vorfeld ausschließlich festgelegt, zu welchem Zeitpunkt der jeweilige Fehler auftritt. Offen blieb damit jedoch, ob der Fehler bei steigendem oder sinkendem Parameterverlauf eintrat, wobei ein Auftreten bei steigendem Verlauf eine aufwändigere Testprozedur bedeutete als bei sinkendem Verlauf (vier versus zwei Prüfelemente). Um den für die Diagnoseprüfung erforderlichen Aufwand über alle Probanden konstant halten zu können, sollten in Experiment II O₂ Sensordefekte sowohl im AutoCAMS-Training als auch in der Testphase nur bei steigenden Werten auftreten. Hierzu wurde zwar der Auftretenszeitpunkt (s nach Programmbeginn) dennoch festgelegt, dies aber systemseitig als frühest möglicher Zeitpunkt interpretiert, sodass die Fehlfunktion erst bei Ansteigen des Sauerstoffwertes auftrat. N₂ Sensordefekte, mit denen eine ähnliche Problematik verbunden ist, traten nur im CAMS-Training auf.

4.4.4 *Material*

4.4.4.1 Fragebogen zur Zuverlässigkeit des Assistenzsystems

In Studie I erwies sich die in Anlehnung an Lee und Moray (1992) gewählte allgemeine Abfrage des Vertrauens in Automation (jeweils nur ein Item bezogen auf CAMS und AFIRA) als ungeeignet, um möglicherweise nur geringe oder spezifische Vertrauensunterschiede abbilden zu können. Daher wurde für Experiment II ein neuer Fragebogen entwickelt, der Rückschlüsse auf das Vertrauen in Automation differenziert für die verschiedenen Systemfunktionen erlauben sollte. Von zentraler Bedeutung waren hierbei die Alarmfunktion von AFIRA, die Diagnosefunktion von AFIRA und die von AFIRA vorgeschlagenen Fehlermanagementschritte.

Um jedoch den Untersuchungsgegenstand nicht transparent werden zu lassen und um den Probanden die Einschätzung zu erleichtern, wurde nicht direkt nach dem Vertrauen in die einzelnen Teilfunktionen von AFIRA gefragt, sondern gebeten, die Zuverlässigkeit von AFIRA diesbezüglich einzuschätzen. Aus demselben Grund wurden neben den drei zentralen Teilfunktionen auch weitere, für die Fragestellung irrelevante Systemfunktionen abgefragt. Auch mit Blick auf das Selbstvertrauen wurden die Probanden nicht direkt befragt, sondern gebeten, ihre eigene Leistung einzuschätzen. Im Zentrum standen dabei, analog zu den Teilfunktionen von AFIRA, die Fehlererkennung und -meldung, die Fehlerdiagnose und das Fehlermanagement. Auch hier

wurden zusätzlich andere Aufgaben im Umgang mit dem System abgefragt, die für die Fragestellung der Studie jedoch zu vernachlässigen sind. Zudem wurde in den Fragebogen ein dritter Abschnitt mit allgemeinen Fragen zur Benutzerfreundlichkeit des Systems eingebaut. Auch dies diente der Kaschierung der eigentlichen Fragestellung der Untersuchung. Items, Antwortformate und Zielvariablen sind in Tabelle 6 zusammengestellt.

Tab. 6: Exp. II: AutoCAMS-Fragebogen

Item		Zielvariable
Bitte schätzen Sie Ihre eigene Leistung im Umgang mit dem System ein:		
1	Ressourcen sparen	Füllitem
2	Manuelle Steuerung O ₂	Füllitem
3	Manuelle Steuerung N ₂	Füllitem
4	Verbindungsprüfung	Füllitem
5	CO ₂ -Check	Füllitem
6	Fehlererkennung und -meldung	Selbsteinschätzung Fehlerdetektionsleistung
7	Fehlerdiagnose	Selbsteinschätzung Fehlerdiagnoseleistung
8	Fehlermanagement	Selbsteinschätzung Fehlermanagementleistung
Bitte schätzen Sie die Zuverlässigkeit der einzelnen Teilsysteme ein:		
9	Sauerstoff (O ₂)	Wahrgenommene Zuverlässigkeit O ₂ -Subsystem (relevant für exp. Gruppeneinteilung)
10	Stickstoff (N ₂)	Wahrgenommene Zuverlässigkeit N ₂ -Subsystem (relevant für exp. Gruppeneinteilung)
11	Kohlenstoffdioxid (CO ₂)	Füllitem
12	Temperatur	Füllitem
13	Luftfeuchtigkeit	Füllitem
14	Ausführung Reparaturauftrag	Füllitem
15	AFIRA Fehlererkennung	Wahrgenommene Zuverlässigkeit Alarmfunktion
16	AFIRA Fehlerdiagnose	Wahrgenommene Zuverlässigkeit Diagnosefunktion
17	AFIRA Fehlermanagement	Wahrgenommene Zuverlässigkeit Handlungsempfehlungen

Item	Zielvariable
Bitte geben Sie an, inwiefern Sie den folgenden Aussagen zustimmen:	
18 Alle Meldungen des Systems waren für mich sofort verständlich.	Füllitem
19 Alle im System verwendeten Begriffe waren für mich sofort verständlich.	Füllitem
20 Es war für mich stets unmittelbar ersichtlich, was Befehle im System bewirken.	Füllitem
21 Das System erschwerte meine Aufgabenbearbeitung durch eine uneinheitliche Gestaltung.	Füllitem
22 Das System hat mich zu überflüssigen Arbeitsschritten gezwungen.	Füllitem
23 Die vom System erzeugten Ausgaben enthielten keine überflüssigen, zu knappen oder unverständlich formulierten Informationen.	Füllitem
24 Die Darstellung der Informationen auf dem Bildschirm unterstützte mich bei der Bearbeitung meiner Aufgaben optimal.	Füllitem

Für die Items 1-17 diente eine 5-stufige Ratingskala mit den Endpolen „gering“ und „hoch“ der Answerterfassung, während für die Items 18-24 eine ebenfalls 5-stufige Ratingskala mit den Endpolen „stimme nicht zu“ und „stimme zu“ Verwendung fand.

Neben diesem „AutoCAMS-Fragebogen“ fand zusätzlich eine Fragebogenversion Einsatz, die sich rein auf CAMS bezog (im Weiteren: „CAMS-Fragebogen“). In dieser wurden die Items 15-17 ersatzlos gestrichen. Beide Fragebogenversionen wurden rechnergestützt dargeboten.

4.4.4.2 Checkliste für Fehlerdetektion, -diagnose und -behebung

Im Unterschied zu Experiment I fand in dieser Studie eine Checkliste für die Fehlerdetektion, -diagnose und -behebung Einsatz. Dies liegt darin begründet, dass es in dieser Studie, anders als zuvor, Aufgabe der Probanden war, bei jeder Fehlerbehandlung zunächst per Mausklick das System in den so bezeichneten Fehlermodus zu versetzen, und diesen nach erfolgter Fehlerbehebung wieder zu deaktivieren. Zur Einübung dieser Handlungssequenz diente die Checkliste, die dementsprechend ausschließlich während des Trainings eingesetzt wurde. Sie sah es für jede auftretende Fehlfunktion vor, handschriftlich zum einen den aufgetretenen Fehler einzutragen, zudem aber auch die manuell erfolgten Handlungsschritte zu quittieren. Diese waren im Detail „Fehlermodus ein“, „Gasstrom/Steuerung ändern“, „Reparaturauftrag absenden“, „Gas-

strom/Steuerung zurücksetzen“ sowie „Fehlermodus aus“. Ein Ausdruck der Checkliste findet sich in Anhang J.

4.4.4.3 Handout zur Fehlerdiagnose und -behebung

Wie in Studie I stand den Probanden auch in Studie II in bestimmten Untersuchungsphasen (siehe nachfolgendes Kap. 4.4.5) ein Handout zur Verfügung. Dieses wurde für Studie II mit Blick auf die Fehlerbehebung um die Aktivierung bzw. Deaktivierung des Fehlermodus ergänzt. Ein Ausdruck des Handouts findet sich in Anhang K.

4.4.5 Durchführung

Die Studie wurde im Frühjahr 2007 an der Technischen Universität Berlin durchgeführt. Jeweils vier Probanden nahmen gleichzeitig teil und kamen an zwei verschiedenen Tagen für jeweils ca. vier Stunden zur Versuchsdurchführung. Da der Untersuchungsablauf in weiten Teilen dem der Studie I ähnelte, wird der Ablauf im Folgenden nur knapp beschrieben und insbesondere auf abweichende Elemente eingegangen.

4.4.5.1 Erster Untersuchungstag

Am ersten Tag fand zunächst eine präsentationsgestützte Einführung in CAMS und die manuelle Fehlerdetektion, -diagnose und -behebung statt (Präsentation siehe Anhang L).

In Abweichung zu Studie I wurden die Probanden daraufhin trainiert, den Fehlermodus zu aktivieren, sobald eine Fehlfunktion, etwa durch die Beobachtung abnormaler Gasdurchflussraten bzw. Parameterverläufe, festgestellt wurde, und zwar noch vor einer gezielten Fehlerdiagnose. Teil der Instruktion war es auch, den Fehlermodus nach der Behebung der Fehlfunktion und der Rückführung der Parameter in den Normalbereich wieder zu deaktivieren. Zur Unterstützung der Einübung dieser Handlungssequenz fand die Checkliste für Fehlerdetektion, -diagnose und -behebung (siehe Anhang J) während des gesamten ersten Untersuchungstages Einsatz. Zusätzlich stand den Probanden für den ersten Tag ein Handout zur Verfügung, das die erforderlichen Fehlerdiagnose- und Behebungsschritte übersichtlich zusammenfasste (siehe Anhang K).

Im Anschluss an die Einführung in das System folgte das CAMS-Training, in dem die Probanden vier mal 16 Minuten mit dem System arbeiteten. Über die vier Übungsblöcke verteilt, traten zehn, weder bezüglich des Fehlertyps noch des Auftretenszeitpunkts vorhersehbare, Fehlfunktionen auf. Tabelle 7 zeigt den Ablauf des Trainings mit Blick auf die einzelnen Blöcke und die darin auftretenden Fehlfunktionen.

Wie ersichtlich, trat jede Fehlfunktion einmal auf, mit Ausnahme der Fehlfunktionen „N₂ Leck“ und „O₂ Blockierung des Ventils“, die es jeweils zweimal zu diagnostizieren und beheben galt. Mit Blick auf die Verteilung der Fehler wurde eine möglichst große Bandbreite der Länge fehlerfreier Phasen angestrebt (zwischen vier und sieben Minuten). Damit sollte verhindert werden, dass Probanden im späteren Versuch bei langen, vermeintlich fehlerfreien Phasen (als solche würde ein AFIRA-Systemausfall zunächst anmuten) skeptisch werden und ihr Informationssuchverhalten entsprechend intensivieren.

Tab. 7: Exp. II: Ablauf des CAMS-Trainings mit Blick auf auftretende CAMS-Fehlfunktionen

Block (jeweils 16 Min)	Nr.	Fehlfunktion	Auftretenszeitpunkt (Min. nach Blockbeginn)
Block 1	1	O ₂ Blockierung des Ventils	3:00
	2	N ₂ offen steckendes Ventil	10:00
Block 2	3	N ₂ Leck	2:00
	4	O ₂ offen steckendes Ventil	7:00
	5	N ₂ Blockierung des Ventils	11:00
Block 3	6	O ₂ Sensordefekt	2:30
	7	N ₂ Leck	9:00
Block 4	8	O ₂ Leck	2:00
	9	O ₂ Blockierung des Ventils	6:00
	10	N ₂ Sensordefekt	12:00

Im Anschluss an das CAMS-Training und somit als Abschluss des ersten Untersuchungstages füllten die Probanden den in Kapitel 4.4.4.1 beschriebenen CAMS-Fragebogen aus.

4.4.5.2 Zweiter Untersuchungstag

Zwischen erstem und zweitem Untersuchungstag lagen für die Probanden im Mittel 7.39 Tage ($SD = 0.59$). Bei der Terminplanung wurde darauf geachtet, dass jeweils vier Probanden derselben Bedingung teilnahmen, um so in der Informationsgruppe „stellvertretenden“ Fehlererfahrungen durch Berichte von Probanden der Erfahrungsgruppe vorzubeugen. Zu Beginn des zweiten Tages wurden kurz die Funktionsweise von CAMS und die Aufgaben der Probanden wiederholt. Es folgte eine präsentationsgestützte Einführung in AutoCAMS (vgl. Anhang M). Mit Blick auf die Reliabilität des Assistenzsystems erhielten alle Probanden die folgende Information:

„Da die Möglichkeit besteht, dass Fehler des Assistenzsystems vorkommen, sollten die Parameter weiterhin überwacht werden.“

Nach dieser Einführung folgte das Training mit Assistenzsystemunterstützung (AutoCAMS-Training), bei dem die Probanden wie schon am ersten Untersuchungstag parallel die Checkliste (vgl. Anhang J) ausfüllten, das Handout (vgl. Anhang K) stand Ihnen jedoch nicht mehr zur Verfügung. Über drei 22-minütige Blöcke traten insgesamt zehn Fehlfunktionen auf. Bei der Hälfte der Probanden wurde für zwei dieser zehn Fehlfunktionen ein Ausfall von AFIRA simuliert (**Erfahrungsgruppe**), während die andere Hälfte der Probanden mit einem stets zuverlässigen Assistenzsystem arbeitete (**Informationsgruppe**). Um sicherzustellen, dass die Probanden der Erfahrungsgruppe den Ausfall des Assistenzsystems auch tatsächlich wahrnahmen, kontrollierte die Versuchsleiterin nach jedem Block die vom Probanden geführte Checkliste und meldete ggf. zurück, welche Fehlfunktion fehlte und wann diese aufgetreten war.

In der nachfolgenden Tabelle (Tab. 8) ist der Ablauf der Trainingsphase mit Blick auf die auftretenden Fehlfunktionen festgehalten, die in unregelmäßigen, für die Probanden nicht antizipierbaren Zeitabständen auftraten.

Tab. 8: Exp. III: Ablauf des AutoCAMS-Trainings mit Blick auf auftretende CAMS-Fehlfunktionen

Block (jeweils 22 Min)	Nr.	Fehlfunktion (tatsächlich vorliegend)	Auftretenszeitpunkt (Min. nach Blockbe- ginn)	Ggf. AFIRA System- ausfall (nur Erfahrungsgruppe)
Block 1	1	O ₂ Sensordefekt	3:00	
	2	O ₂ Leck	11:30	
	3	N ₂ offen steckendes Ventil	17:00	
Block 2	4	N ₂ Sensordefekt	2:30	
	5	O ₂ Leck	6:30	AFIRA-Ausfall 1
	6	N ₂ Blockierung des Ventils	11:00	
	7	N ₂ Leck	19:00	
Block 3	8	O ₂ offen steckendes Ventil	5:00	
	9	N ₂ Sensordefekt	9:00	AFIRA-Ausfall 2
	10	O ₂ Blockierung des Ventils	13:30	

Alle Probanden füllten nach dem AutoCAMS-Training zum ersten Mal den AutoCAMS-Fragebogen (vgl. Kap. 4.4.4.1) aus.

Daran schloss sich die Testphase an, unterteilt in vier Blöcke mit jeweils 25-minütiger Dauer. Zwischen diesen Blöcken erhielten die Probanden die Gelegenheit zu 5-minütigen Pausen. In dieser Phase arbeiteten die Probanden ohne Checkliste. Tabelle 9 zeigt den Ablauf der Testphase mit Blick auf die CAMS-Fehlfunktionen und deren Auftretenszeitpunkt im jeweiligen Block.

Tab. 9: Exp. II: Ablauf der Testphase mit Blick auf auftretende CAMS-Fehlfunktionen

Block (jeweils 25 Min.)	Nr.	Fehlfunktion (tatsächlich vorliegend)	Auftretenszeit- punkt (Min. nach Blockbeginn)	AFIRA-Fehlfunktionen (alle Probanden)
Block 1	1	O ₂ Blockierung des Ventils	3:00	
	2	N ₂ Sensordefekt	8:00	
	3	O ₂ offen steckendes Ventil	16:30	
	4	N ₂ Blockierung des Ventils	21:00	
Block 2	5	O ₂ Sensordefekt	4:00	
	6	O ₂ Leck	12:00	
Block 3	7	O ₂ Sensordefekt	2:30	
	8	N ₂ Leck	11:00	
	9	N ₂ offen steckendes Ventil	16:30	
	10	O ₂ Blockierung des Ventils	21:00	AFIRA-Ausfall
Block 4	11	N ₂ Blockierung des Ventils	2:30	
	12	N ₂ offen steckendes Ventil	8:00	
	13	N ₂ Leck	12:30	AFIRA-Ausfall
	14	O ₂ offen steckendes Ventil	16:30	AFIRA-Fehldiagnose: O ₂ Sensordefekt

Für die ersten neun CAMS-Fehlfunktionen arbeitete AFIRA zuverlässig, beim zehnten Fehler wurde jedoch ein Ausfall des Systems simuliert, der eine rein manuelle Fehlerdetektion, -diagnose und -behebung erforderte. Darauf folgten zwei Fehlfunktionen, bei denen AFIRA wieder zuverlässig arbeitete, bevor bei Fehler 13 erneut ein Systemausfall simuliert wurde. Als Fehler-typen für die Systemausfälle wurden hier jene gewählt, die bereits im manuellen CAMS-Training zweifach und damit in besonderem Maße geübt worden waren: Eine O₂ Blockierung des Ventils und ein N₂ Leck. Bei Fehler 14 generierte AFIRA schließlich eine Fehldiagnose. Statt des tatsächlich vorliegenden offen steckenden Sauerstoffventils diagnostizierte das Assistenzsystem einen O₂ Sensordefekt. Sowohl die tatsächlich vorliegende Fehlfunktion als auch der fälschlich von AFIRA vorgeschlagene Fehler waren vorher wiederholt aufgetreten (2 bzw. 3 Mal), sodass davon auszu-

gehen ist, dass alle Probanden hinreichend geübt waren, um die AFIRA-Diagnose falsifizieren zu können. Im Vergleich zu Studie I war diese Fehldiagnose jedoch nur über eine umfangreiche Prüfprozedur (4 Prüfelemente im Vergleich zu 1 Prüfelement) als solche zu erkennen.

Nach der Testphase füllten die Probanden ein zweites und damit letztes Mal den Auto-CAMS-Fragebogen aus.

4.4.6 Abhängige Variablen

Soweit nicht anders vermerkt, wurden die abhängigen Variablen wieder aus den in Logfiles protokollierten Mausklicks und Tastatureingaben der Probanden unter Berücksichtigung des jeweils vorliegenden Systemzustands abgeleitet.

4.4.6.1 Omission Fehler

Zur Operationalisierung etwaiger *omission* Fehler dienten die AFIRA-Systemausfälle bei Fehler 10 und 13. Für beide Fehler wurde als Zeitraum, innerhalb dessen die Fehlfunktion vom Probanden erkannt werden musste, die Zeit bis zum Verlassen des Sollwertbereichs, also die Zeit bis zum Eintreten eines kritischen Systemzustandes, definiert. Diese beträgt durchschnittlich für Fehler 10 (O₂ Blockierung des Ventils) 20 s und für Fehler 13 (N₂ Leck) 35 s nach Auftreten der Fehlfunktion. Sofern die Probanden innerhalb dieses Zeitraums den Fehlermodus nicht aktivierten, um damit die Detektion einer vorliegenden Fehlfunktion zu signalisieren, wurde dies als *omission* Fehler gewertet.

4.4.6.2 Commission Fehler

Commission Fehler wurden analog zu Experiment I operationalisiert (vgl. Kap. 3.3.5.5), wobei die Fehldiagnose von AFIRA in Experiment II bei Fehlfunktion 14 auftrat.

4.4.6.3 Informationssuchverhalten in fehlerfreien Phasen

Auch das Informationssuchverhalten wurde analog zu Experiment I operationalisiert, allerdings bestand hier ein Unterschied in der Definition der „fehlerfreien“ Phasen: Während in Experiment I der potenzielle Beginn einer fehlerfreien Phase mit der vollständigen Behebung des Parameters gleichgesetzt wurde, zählte in Experiment II erst die Zeit ab Deaktivierung des Fehlermodus als „fehlerfrei“. Das Ende der fehlerfreien Phase wurde in beiden Studien mit dem Auftreten der Fehlfunktion (Beginn der Fehlerdiagnose- bzw. Fehlerdetektionszeit) gleichgesetzt und

jeweils das Informationssuchverhalten in den vorhergehenden 120 s (sofern diese fehlerfrei waren) analysiert. Die Definitionen der einzelnen Variablen entsprechen denen in Experiment I (Kap. 3.3.5.4).

4.4.6.4 Fehlerdiagnosezeit

Die Fehlerdiagnosezeit wurde genauso wie in Experiment I (vgl. Kap. 3.3.5.1) über die Zeit (s) zwischen Auftreten einer Fehlfunktion (potenzielle Detektierbarkeit) und dem Absenden des richtigen Reparaturauftrages operationalisiert.

4.4.6.5 Informationssuchverhalten allgemein in Fehlerdiagnosephasen

Das Informationssuchverhalten in Fehlerdiagnosephasen wurde analog zu Studie I ausgewertet (siehe Kap. 3.3.5.2) und entsprechend über die Gesamtzahl an Informationsabrufen und den Anteil relevanter Informationsabrufe operationalisiert.

4.4.6.6 Complacency bezüglich der Diagnosefunktion

Complacency, erfasst über die Überprüfung der AFIRA-Diagnosen, wurde ebenfalls analog zu Experiment I (vgl. Kap. 3.3.5.3) operationalisiert. Der einzige Unterschied bestand darin, dass O₂ Sensordefekte nur bei steigenden Parameterwerten auftraten und somit jeder O₂ Sensordefekt vier Prüfelemente erforderte. Dies wurde so realisiert, um interindividuell besser vergleichbare Bedingungen zu schaffen.

4.4.6.7 „Return-to-manual“-Leistung

Für die Operationalisierung der *return-to-manual*-Leistung dienten Fehler 10 und 13, die aufgrund des simulierten AFIRA-Ausfalls eine manuelle Fehlerdetektion, -diagnose und -behebung erforderten.

Analog zu Experiment I (vgl. Kap. 3.3.5.6) wurde wiederum die Anzahl der auf Anhieb korrekt abgeschickten Reparaturaufträge analysiert (0, 1 oder 2). Allerdings wurde zudem, nicht wie in Experiment I die Fehlerdiagnosezeit, sondern die Fehleridentifikationszeit, d.h. die Zeit zwischen Entdecken des Fehlers (Aktivierung des Fehlermodus) und Absenden des richtigen Reparaturauftrags (unabhängig davon, ob vorher bereits falsche Aufträge abgeschickt worden waren) erfasst. Dieses Maß wurde gewählt, um eine Konfundierung von *omission* Fehlern und Fertigkeitsverlusten vermeiden zu können. Dies wäre in Experiment II bei einer Betrachtung der Fehlerdi-

agnosezeit letztlich nicht gegeben, da eine hohe Zeitdauer zwischen Auftreten einer Fehlfunktion und Abschicken des richtigen Reparaturauftrags sowohl einer sehr späten Fehlerdetektion (*omission* Fehler) als auch einer langwierigen Fehleridentifikation (Fertigkeitsverlust) geschuldet sein kann. In Experiment I wurde den Probanden dagegen die Fehlerdetektion immer durch den stets korrekt funktionierenden Masteralarm abgenommen.

4.4.6.8 Leistung in der prospektiven Gedächtnis- und Reaktionszeitaufgabe

Die Leistung in der prospektiven Gedächtnisaufgabe sowie in der Reaktionsaufgabe wurde wie in Experiment I beschrieben operationalisiert (vgl. Kap. 3.3.5.7). Der einzige Unterschied bestand in der Definition der „fehlerfreien“ Phasen, wobei der Beginn einer fehlerfreien Phase mit der Deaktivierung des Fehlermodus gleichgesetzt wurde.

4.4.6.9 Wahrgenommene Zuverlässigkeit des Assistenzsystems und Einschätzung der eigenen Leistung

Sowohl die wahrgenommene Zuverlässigkeit des Assistenzsystems als auch die Einschätzung der eigenen Leistung wurde mit dem in Kapitel 4.4.4.1 beschriebenen AutoCAMS-Fragebogen erfasst. Die Items 6-8 dienten zur Erfassung der Selbsteinschätzung der Fehlerdetektions-, Fehlerdiagnose- und Fehlermanagementleistung. Die Items 15-17 wurden zur Erfassung der wahrgenommenen Zuverlässigkeit von AFIRA, genauer der Alarmfunktion, der Diagnosefunktion und der vorgeschlagenen Handlungsempfehlungen herangezogen.

Für alle sechs Variablen wurden die Selbsteinschätzungen in eine numerische Skala von 1 (gering) bis 5 (hoch) übertragen. Insgesamt erfolgten drei Erhebungen der subjektiven Daten:

- Nach dem CAMS-Training (vor der Bedingungsmanipulation): CAMS-Fragebogen
- Nach dem AutoCAMS-Training (nach der Bedingungsmanipulation): AutoCAMS-Fragebogen
- Nach der Testphase: AutoCAMS-Fragebogen

Die statistische Auswertung der Daten von Experiment II erfolgte analog zu Experiment I (vgl. Kap. 3.3.6).

4.5 ERGEBNISSE EXPERIMENT II

4.5.1 Omission Fehler

Für den zehnten Fehler wurde der erste von zwei Systemausfällen von AFIRA simuliert. 80 % der Informationsgruppe, aber nur 18 % der Erfahrungsgruppe entdeckten diesen Fehler nicht oder erst nach Überschreitung des kritischen Wertebereichs, begingen also einen *omission* Fehler, $p < .01$, Fisher-Yates Test (einseitig). Beim zweiten Systemausfall (Fehler 13) verschwand dieser Gruppenunterschied. Wiederum übersahen 18 % der Erfahrungsgruppe, aber diesmal nur 22 % der Informationsgruppe die Fehlfunktionen (siehe Abb. 25).

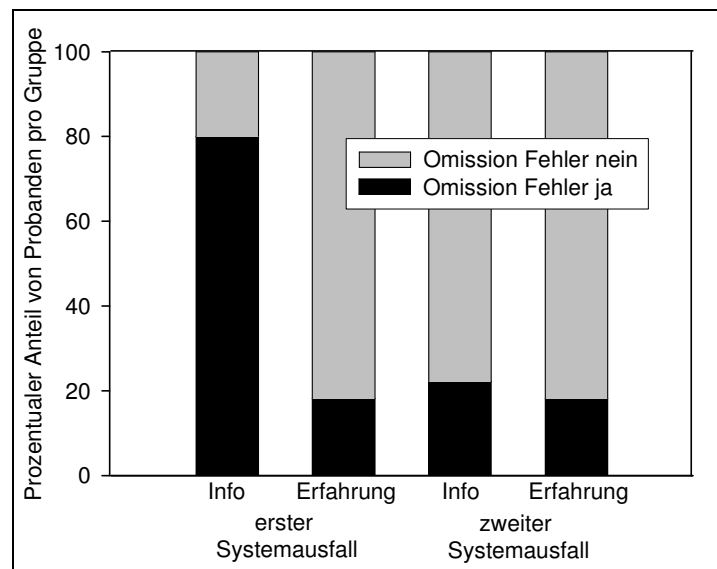


Abb. 25: Exp. II: *Omission* Fehler beim ersten und zweiten Systemausfall (experimentelle Gruppen)

Für die im Folgenden beschriebenen abhängigen Variablen interessierte – neben potenziellen Unterschieden aufgrund der experimentellen Manipulation – ob und inwiefern sich Probanden, die einen *omission* Fehler begingen, von den übrigen Probanden unterscheiden. Hierzu wurden die Probanden in zwei Gruppen eingeteilt: Die Gruppe „*omission* Fehler ja“ wurde durch Probanden gebildet, die bei mindestens einem der beiden Systemausfälle den Fehler übersehen hatten ($n = 12$), während die Gruppe „*omission* Fehler nein“ alle übrigen Probanden umfasste ($n = 10$). Ein Proband konnte keiner der beiden Gruppen zugeteilt werden, da er in beiden Fällen vor einer Aktivierung des Fehlermodus einen Reparaturauftrag abschickte, dies aber erst kurz nach Überschreiten des kritischen Wertebereichs geschah. Ob dieser Proband die CAMS-Fehler noch innerhalb des kritischen Zeitfensters entdeckt hatte oder erst unmittelbar vor Abschicken des Reparaturauftrags, ließ sich somit nicht entscheiden.

4.5.2 Commission Fehler

Bei der 14. Fehlfunktion von CAMS generierte AFIRA eine Fehldiagnose. 74 % der Probanden folgten dieser fälschlicher Weise und begingen somit einen *commission* Fehler. Die Probanden verteilten sich dabei gleich auf die beiden experimentellen Gruppen (Erfahrungsgruppe: $n = 9$; Informationsgruppe: $n = 8$). Die Bedingungsmanipulation hatte somit keinen Einfluss auf das Auftreten von *commission* Fehlern, $p = .64$, Fisher-Yates-Test (zweiseitig). Eine Analyse der Diagnoseprüfung unmittelbar vor Abschicken des falschen Reparaturauftrages zeigte, dass jene Probanden, die der Fehldiagnose nicht folgten, ausnahmslos alle zur Überprüfung notwendigen Informationen vorher abgerufen hatten. Von den übrigen Probanden (*commission* Fehler ja) folgten 82 % der Automation aufgrund variierender Ausprägungen von *complacency* (keine oder unvollständige Überprüfung, siehe Abb. 26). 18 % der Probanden vollzogen jedoch alle zur Diagnoseprüfung erforderlichen Schritte und entschieden sich dennoch dafür, trotz der eindeutig widersprechenden Informationen, der automatisch generierten Direktive zu folgen.

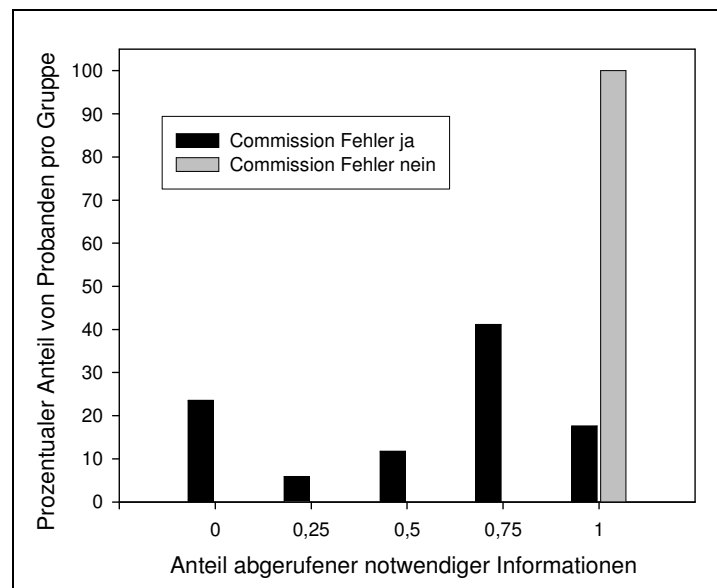


Abb. 26: Exp. II: Überprüfung der Fehldiagnose (Fehler 14) und *commission* Fehler

Probanden, die einen *omission* Fehler begingen, wiesen keine erhöhte Rate von *commission* Fehlern auf, $p = .635$, Fischer-Yates-Test (zweiseitig). Die Verteilung ist Tabelle 10 zu entnehmen.

Tab. 10: Exp. II: Häufigkeitsverteilung von *commission* und *omission* Fehlern

		Omission Fehler		
		Ja	Nein	Gesamt
Commission Fehler	Ja	9	7	16
	Nein	3	3	6
	Gesamt	12	10	22

Auch bezüglich des *commission* Fehlers interessierte, inwiefern sich die Probanden, die einen solchen Fehler begingen, von jenen, die der Fehldiagnose nicht folgten, unterscheiden. Entsprechend werden im Folgenden jeweils auch die Ergebnisse für diese Gruppeneinteilung berichtet (*commission* Fehler ja: $n = 17$; *commission* Fehler: nein $n = 6$).

4.5.3 Überwachung der Alarmfunktion: Informationssuchverhalten in fehlerfreien Phasen

Die Informationssuchdaten der Probanden in den fehlerfreien Phasen vor den Fehlern 1-3, 4-6 und 7-9 wurden zu drei Blöcken zusammengefasst, um intraindividuelle Varianz zu reduzieren. Die Auswertung erfolgte mithilfe von 2(Gruppe) x 3(Fehlerblock) Varianzanalysen (ANOVAs) mit „Fehlerblock“ als Messwiederholungsfaktor. Als globaler Indikator für das Informationssuchverhalten diente die durchschnittliche Anzahl aller Informationsabrufe pro Minute, wobei diese jeweils bezogen auf die letzten 120 s vor Beginn einer Fehlfunktion ermittelt wurde. Wie aus Abbildung 27 ersichtlich wird, riefen Probanden der Erfahrungsgruppe signifikant mehr Informationen ab ($M = 19.06$, $SD = 5.75$) als Probanden der Informationsgruppe ($M = 14.43$, $SD = 4.86$), Haupteffekt „Gruppe“ $F(1, 21) = 4.37$, $p = .05$. Weder ein Haupteffekt „Fehlerblock“ noch ein Interaktionseffekt wurde beobachtet, $F < 1$.

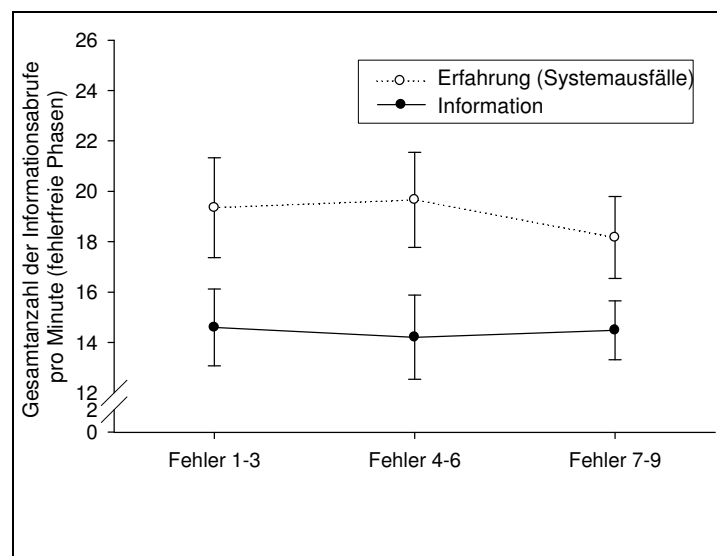


Abb. 27: Exp. II: Informationsabrufe in fehlerfreien Phasen (experimentelle Gruppen)

Im Mittel betrug der Anteil relevanter Informationsabrufe 86 % ($SD = 10$) an der Gesamtanzahl von Informationsabrufen. Dabei zeigten sich weder Unterschiede zwischen den Gruppen noch im zeitlichen Verlauf, Interaktionseffekt „Fehlerblock“ x „Gruppe“ $F(2, 20) = 1.02$, $p = .37$, alle übrigen $F < 1$.

Die Auswertung basierend auf der Gruppeneinteilung „*omission* Fehler ja/nein“, mithilfe von $2(\text{omission Fehler ja/nein}) \times 3(\text{Fehlerblock})$ ANOVAs mit Messwiederholung für den zweiten Faktor, ergab keine signifikanten Effekte, weder für die Gesamtzahl an Informationsabrufen noch für den Anteil relevanter Informationsabrufe, alle $F < 1.30$

Ein etwas anderes Bild zeigte sich, wenn die Einteilung auf Basis der Gruppeneinteilung „*commission* Fehler ja/nein“ erfolgte. Auch hier zeigten sich keine signifikanten Effekte mit Blick auf die Gesamtzahl von Informationsabrufen, $F < 1$, jedoch bezüglich des Anteils relevanter Informationsabrufe (siehe Abb. 28):

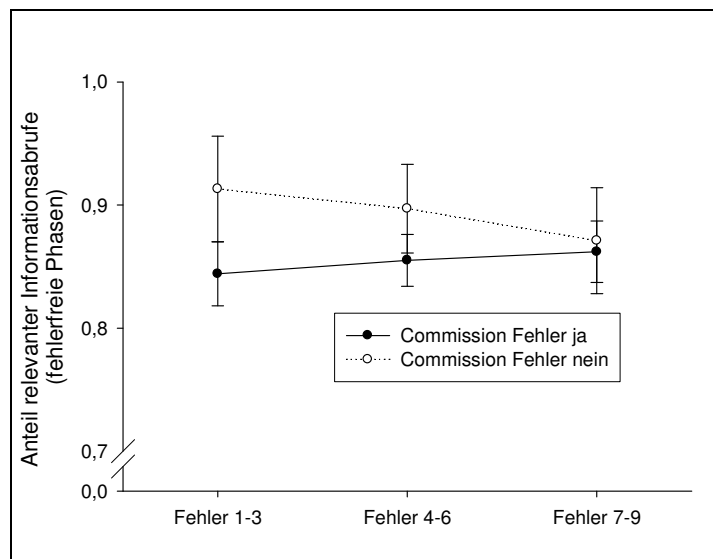


Abb. 28: Exp. II: Anteil relevanter Informationsabrufe in fehlerfreien Phasen (*commission* Fehler ja/nein)

Probanden, die einen *commission* Fehler vermeiden konnten, zeigten über den zeitlichen Verlauf hinweg eine leichte Abnahme des Anteils relevanter Parameter, während Probanden, die einen solchen Fehler begingen, eine leichte Zunahme beobachten ließen. Signifikant wurde entsprechend der Interaktionseffekt „Fehlerblock“ x „Gruppe“, $F(2, 42) = 3.18, p = .05$, alle übrigen $F < 1$.

4.5.4 Fehlerdiagnosezeit

Analog zu Experiment I diente die Zeit vom Auftreten einer Fehlfunktion bis zum Absenden des richtigen Reparaturauftrages als allgemeiner Leistungsparameter für die Fehlerdiagnose. Die Probanden benötigten hierfür bei den ersten neun, von AFIRA korrekt diagnostizierten, Fehlfunktionen durchschnittlich 23.27 s ($SD = 11.49$).

Erwartungsgemäß trat, im Unterschied zu den Befunden von Experiment I, diesmal kein Effekt der Trainingsintervention auf. Eine auf der experimentellen Gruppeneinteilung basierende

2(Gruppe) x 3(Fehlerblock) ANOVA mit Messwiederholung ergab keine Gruppenunterschiede, Haupteffekt „Gruppe“ $F(1, 21) = 1.08, p = .31$. Zudem wurde weder ein Messwiederholungseffekt (Faktor „Fehlerblock“) noch ein Interaktionseffekt beobachtet, $F(2, 42) = 2.63; p = .09$ bzw. $F(2, 42) = 1.35, p = .27$.

Auch bei einer Einteilung der Gruppen nach „*omission* Fehler ja/nein“ zeigte sich dasselbe Bild, Haupteffekt „*omission* Fehler ja/nein“ $F(1, 20) = 1.82, p = .19$, Haupteffekt „Fehlerblock“ $F(2, 40) = 2.30, p = .11$, Interaktionseffekt $F < 1$.

Allerdings wurde, wie auch schon in Experiment I, ein signifikanter Unterschied zwischen Probanden, die einen *commission* Fehler begingen und jenen, die diesen vermeiden konnten, gefunden. Letztere nahmen sich signifikant mehr Zeit für die Fehlerdiagnose ($M = 32.83, SD = 19.85$) als die übrigen Probanden ($M = 19.90, SD = 11.80$), Haupteffekt „*commission* Fehler ja/nein“ $F(1, 21) = 7.21, p = .01$. Kein anderer Effekt wurde signifikant, Haupteffekt „Fehlerblock“ $F(2, 42) = 2.79, p = .07$, Interaktionseffekt $F < 1$.

4.5.5 Überprüfung der Diagnosefunktion

4.5.5.1 Informationssuchverhalten allgemein in Fehlerdiagnosephasen

Im Mittel tätigten die Probanden bei den ersten neun Fehlfunktionen zwischen Auftreten des Fehlers und dem Abschicken des ersten Reparaturauftrages 4.45 ($SD = 2.77$) Informationsabrufe. Der Anteil relevanter Informationsabrufe betrug dabei im Mittel 79 % ($M = .79, SD = .12$). Die statistische Auswertung erfolgte wiederum zunächst basierend auf der experimentellen Gruppeneinteilung über 2(Gruppe) x 3(Fehlerblock) Varianzanalysen mit Messwiederholung (ANOVAs). Bezüglich der Gesamtanzahl von Informationsabrufen ergaben sich keine signifikanten Effekte, Haupteffekt „Fehlerblock“ $F(2, 42) = 3.12, p = .06$, Haupteffekt „Gruppe“ $F < 1$, Interaktionseffekt $F(2, 42) = 1.10, p = .34$. Auch der Anteil relevanter Informationsabrufe unterschied sich nicht zwischen den Gruppen und zeigte keine Variation über den zeitlichen Verlauf, Haupteffekt „Fehlerblock“ $F(2, 42) = 1.67, p = .20$, Haupteffekt „Gruppe“ $F < 1$, Interaktionseffekt $F(2, 42) = 1.42, p = .25$.

Erfolgte die Gruppeneinteilung gemäß „*omission* Fehler ja/nein“, ergab die Auswertung dasselbe Bild. Kein Effekt wurde signifikant, für die Gesamtanzahl an Informationsabrufen Haupteffekt „*omission* Fehler ja/nein“ $F(1, 20) = 2.98, p = .10$, Haupteffekt „Fehlerblock“ $F(2, 40) = 2.91, p = .07$, Interaktionseffekt $F < 1$; für den Anteil relevanter Informationsabrufe: Haupteffekt „Fehlerblock“ $F(2, 36) = 1.27, p = .29$, alle übrigen $F < 1$.

Wurden die Daten jedoch basierend auf der Gruppeneinteilung „*commission* Fehler ja/nein“ analysiert, ergab sich ein anderes Befundmuster (siehe Abb. 29):

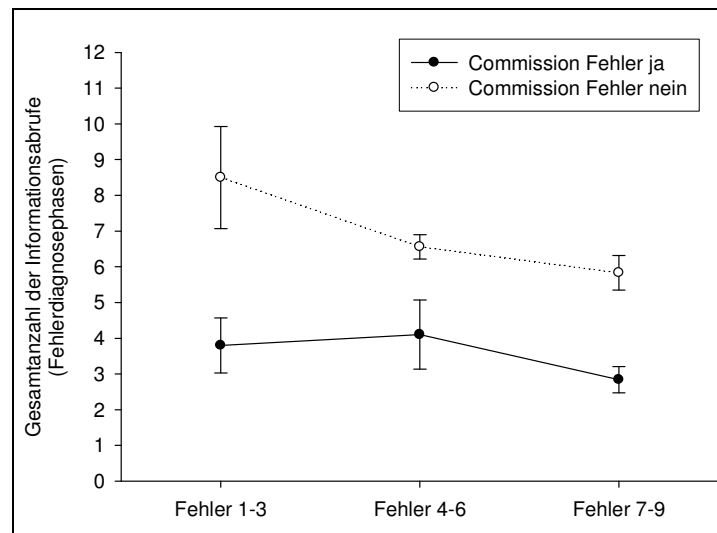


Abb. 29: Exp. II: Informationsabrufe in Fehlerdiagnosephasen (*commission* Fehler ja/nein)

Probanden, die einen *commission* Fehler vermeiden konnten, riefen insgesamt mehr Informationen ($M = 6.96$, $SD = 1.46$) ab als Probanden, die der falschen Diagnose folgten ($M = 3.57$, $SD = 2.59$). Zudem nahm die Anzahl von Informationsabrufen über die Zeit hinweg ab. Entsprechend ergab die 2(*commission* Fehler ja/nein) x 3(Fehlerblock) ANOVA mit Messwiederholung einen hoch signifikanten Haupteffekt „*commission* Fehler ja/nein“ $F(1, 21) = 9.12$, $p = .007$ und einen signifikanten Messwiederholungseffekt (Faktor: „Fehlerblock“), $F(2, 42) = 3.54$, $p = .04$. Ein Interaktionseffekt wurde nicht beobachtet, $F < 1$.

Der beschriebene Gruppenunterschied für alle Informationsabrufe fand sich jedoch nicht für den Anteil relevanter Informationsabrufe, wie auch kein anderer Effekt bei diesem Maß statistische Signifikanz erreichte, Haupteffekt „Fehlerblock“ $F(2, 44) = 1.26$, $p = .30$, alle übrigen $F < 1$.

4.5.5.2 Complacency bezüglich der Diagnosefunktion

Der Umfang, in dem die ersten neun AFIRA-Diagnosen von den Probanden einer Überprüfung unterzogen wurden, ist in Abb. 30, basierend auf der experimentellen Gruppeneinteilung dargestellt. Wie in Experiment I überprüfte keine der beiden Gruppen die vorgeschlagenen Diagnosen vollständig, bevor sie den Diagnosen folgten und entsprechende Reparaturaufträge abschickten ($M = .68$, $SD = .24$). Abweichend von Experiment I unterschieden sich die experimentellen Gruppen jedoch nicht. Offenbar vermochte es die Erfahrung von Systemausfällen im Training nicht, *complacency* mit Blick auf die Diagnosefunktion von AFIRA zu reduzieren, Haupt-

effekt „Gruppe“ $F < 1$. Auch der Faktor „Fehlerblock“ und die Interaktion der beiden Faktoren beeinflusste das Überprüfungsverhalten nicht, beide $F < 1$.

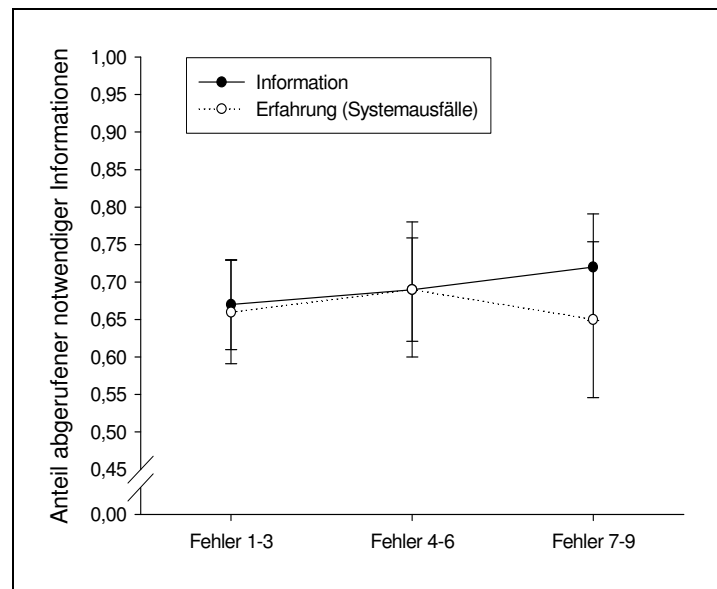


Abb. 30: Exp. II: *Complacency* (experimentelle Gruppen)

Die Datenanalyse mit dem Gruppenfaktor „*omission* Fehler ja/nein“ zeigte ebenfalls keine Unterschiede zwischen den Gruppen. Eine 2(*omission* Fehler ja/nein) x 3(Fehlerblock) ANOVA mit Messwiederholung ergab Haupteffekt „*omission* Fehler ja/nein“ $F(1, 20) = 1.66, p = .21$, alle übrigen Effekte $F < 1$.

Allerdings zeigten sich Gruppenunterschiede, wenn die Einteilung gemäß „*commission* Fehler ja/nein“ erfolgte (siehe Abb. 31). Wie in Experiment I zeichnten sich auch in Experiment II Probanden, die einen *commission* Fehler begingen, durch eine höhere *complacency*-Ausprägung aus.

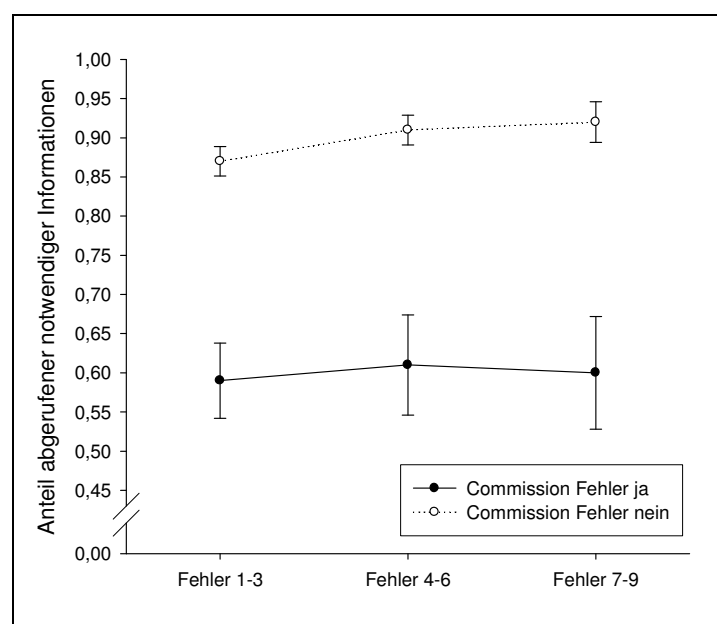


Abb. 31: Exp. II: *Complacency* (*commission* Fehler ja/nein)

Probanden, die die Fehldiagnose entdeckten, riefen bei den vorhergehenden neun Fehlfunktionen einen deutlich höheren Anteil der erforderlichen Prüfelemente ab ($M = 0.90$, $SD = 0.04$) als Probanden, die die Fehldiagnose übersahen ($M = 0.60$, $SD = 0.24$), Haupteffekt „*commission* Fehler ja/nein“, $F(1, 21) = 8.78$, $p = .007$. Kein anderer Effekt wurde signifikant ($F < 1$).

4.5.6 „Return-to-manual“-Leistung

Bei den beiden Systemausfällen von AFIRA (Fehler 10 und 13) mussten die Probanden die Fehler rein manuell entdecken, diagnostizieren und beheben. Ausgezählt wurde zunächst, wie viele der Probanden bei den beiden Fehlern als erstes einen falschen Reparaturauftrag abschickten. Insgesamt unterlief $n = 8$ der Probanden ein solcher Fehler, allerdings jeweils nur bei einer der beiden Fehlfunktionen. Davon entstammten 6 Probanden der Erfahrungsgruppe und 2 der Informationsgruppe, der Unterschied wurde jedoch nicht signifikant, $p = .19$, zweiseitiger Fisher-Yates-Test. Auch zeigte sich kein Zusammenhang mit dem Begehen eines *commission* Fehlers bzw. eines *omission* Fehlers, in beiden Fällen $p = 1.00$, zweiseitiger Fischer-Yates-Test.

Zur Ermittlung eines potenziellen Fertigkeitsverlusts diente die durchschnittliche Fehleridentifikationszeit, operationalisiert über die Zeit zwischen Entdecken der Fehlfunktion (Aktivierung des Fehlermodus) und Absenden des richtigen Reparaturauftrags. Diese wurde für die beiden Fehlfunktionen 10 und 13 gemittelt und mit den mittleren Fehleridentifikationszeiten der korrespondierenden Fehlfunktionen im CAMS-Training verglichen. Die Ergebnisse sind in Abbildung 32 dargestellt. Ein Fertigkeitsverlust ist für keine der beiden Gruppen festzustellen. Stattdessen zeigt sich, wie auch schon in Experiment I, ein deutlicher Trainingseffekt für beide Gruppen.

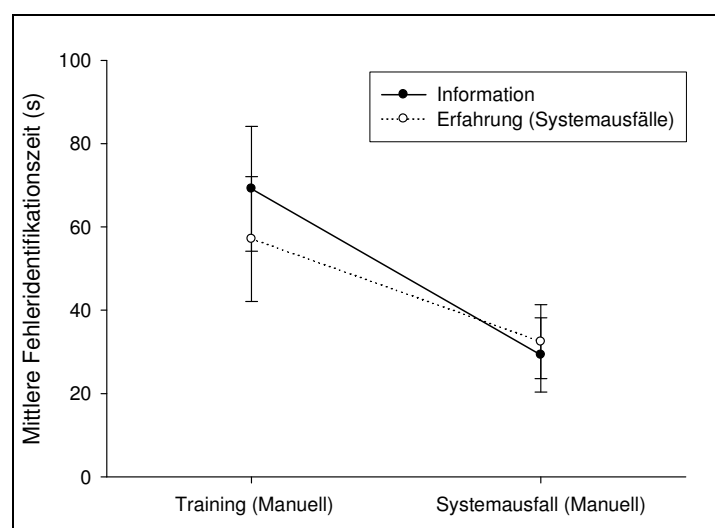


Abb. 32: Exp. II: Return-to-manual-Leistung (experimentelle Gruppen)

Während die Probanden im Training im Mittel 62.14 s ($SD = 48.04$) für die Fehleridentifikation benötigten, gelang ihnen dies am Ende der Testphase bereits nach der Hälfte der Zeit ($M = 30.86$, $SD = 28.90$). Die statistische Auswertung erfolgte über eine 2(Gruppe) x 2(Training vs. Systemausfall) ANOVA mit Messwiederholung für den zweiten Faktor und zeigte entsprechend einen hochsignifikanten Haupteffekt „Training vs. Systemausfall“, $F(1, 20) = 8.40$, $p = .009$, für alle anderen $F < 1$.

Auch bei den beiden anderen Gruppeneinteilungen – *omission* Fehler ja/nein und *commission* Fehler ja/nein – zeigten sich keine Unterschiede zwischen den Gruppen. Signifikant wurden erneut nur die Haupteffekte für den Faktor „Training vs. Systemausfall“, $F(1, 20) = 8.10$, $p = .01$ bzw. $F(1, 20) = 7.44$, $p = .01$, für alle übrigen Effekte $F < 1$.

4.5.7 Leistung in der prospektiven Gedächtnis- und Reaktionszeitaufgabe

Wie bereits bei Experiment I erfolgte die Auswertung mit Hilfe von 2(Gruppe) x 3(Fehlerblock) x 2(Fehlerfrei versus Fehler) ANOVAs mit den letzten beiden Faktoren als *within*-Faktoren. Einbezogen wurden dabei jeweils die ersten neun Fehlfunktion bei denen AFIRA korrekte Diagnosen lieferte.

Im Folgenden wird erst auf die Ergebnisse für die prospektive Gedächtnisaufgabe und dann für die Reaktionszeitaufgabe eingegangen. Geschildert werden jeweils die Ergebnisse der Auswertung für die drei verschiedenen Gruppeneinteilungen experimentell, *omission* Fehler ja/nein und *commission* Fehler ja/nein.

Mit Blick auf die durchschnittliche Abweichung vom geforderten Eintragezeitpunkt bei der prospektiven Gedächtnisaufgabe zeichnete die Auswertung, basierend auf der experimentellen Gruppeneinteilung, folgendes Bild (vgl. Abb. 33): Zum einen zeigte sich über die Fehlerblöcke hinweg eine Abnahme der Abweichungen und somit eine Leistungsverbesserung im Sinne eines Trainingseffekts, Haupteffekt „Fehlerblock“ $F(2, 42) = 4.63$, $p = .02$. Zudem fielen die durchschnittlichen Abweichungszeiten in Fehlerphasen höher aus als in den fehlerfreien Phasen, Haupteffekt „Fehlerfrei vs. Fehler“ $F(1, 21) = 9.24$, $p = .006$. Weitere signifikante Effekte wurden nicht beobachtet, Interaktionseffekt „Fehlerblock“ x „Fehlerfrei vs. Fehler“ $F(2, 42) = 2.94$, $p = .06$, für alle übrigen $F < 1$.

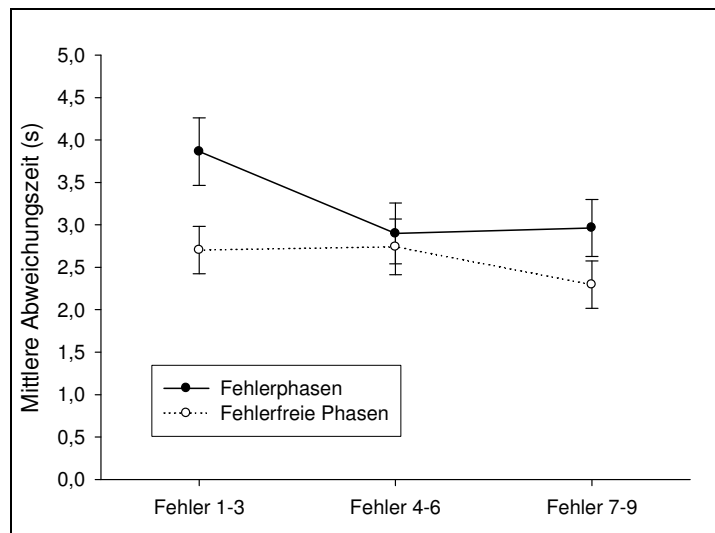


Abb. 33: Exp. II: Leistung in der prospektiven Gedächtnisaufgabe (fehlerfreie vs. Fehlerphasen)

Bei einer Einteilung der Gruppen danach, ob mindestens ein *omission* Fehler begangen worden war oder nicht, ergab die varianzanalytische Auswertung dasselbe Bild: Beobachtet wurden signifikante Haupteffekte für die Faktoren „Fehlerblock“, $F(2, 40) = 3.68, p = .03$ und „Fehlerfrei vs. Fehler“, $F(1, 20) = 7.52, p = .01$. Alle übrigen Effekte erreichten keine statistische Signifikanz, Interaktionseffekte „Fehlerblock“ x „*omission* Fehler ja/nein“ $F(2, 40) = 2.23, p = .12$, „Fehlerblock“ x „Fehlerfrei vs. Fehler“ $F(2, 40) = 2.02, p = .15$, alle übrigen $F < 1$.

Ein etwas anderes Bild ergab sich dagegen, wenn die Gruppeneinteilung auf der Basis von „*commission* Fehler ja/nein“ vorgenommen wurde (siehe Abb. 34):

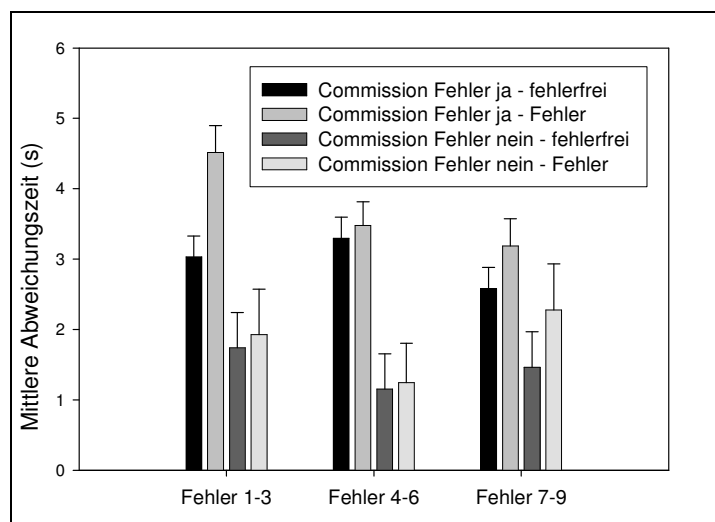


Abb. 34: Exp. II: Leistung in der prospektiven Gedächtnisaufgabe für fehlerfreie und Fehlerphasen (*commission* Fehler ja/nein)

Probanden, die einen *commission* Fehler begingen, zeigten von Beginn an deutlich höhere Abweichungszeiten, konnten diese aber über den Verlauf der Testphase hinweg im Sinne eines Trainingseffekts verbessern. Probanden, die die falsche Diagnose erkannten, nahmen die Eintra-

gungen dagegen insgesamt näher am geforderten Eintragezeitpunkt vor und verbesserten sich nicht über den zeitlichen Verlauf. Darüber hinaus fielen in beiden Probandengruppen die Abweichungszeiten bei Vorliegen einer Fehlfunktion höher aus als in fehlerfreien Phasen. Die statistische Auswertung ergab einen Haupteffekt „*commission* Fehler ja/nein“, $F(1, 21) = 12.99, p = .002$, einen Haupteffekt für den Faktor „Fehlerfrei vs. Fehler“, $F(1, 21) = 5.16, p = .03$, sowie einen signifikanten Interaktionseffekt „Fehlerblock“ x „*commission* Fehler“, $F(2, 42) = 3.39, p = .04$. Alle weiteren geprüften Effekte wurden nicht signifikant, Haupteffekt „Fehlerblock“ $F(2, 42) = 2.63, p = .08$, Interaktionseffekte „Fehlerblock“ x „Fehlerfrei vs. Fehler“ $F(2, 42) = 1.34, p = .27$, „Fehlerblock“ x „Fehlerfrei vs. Fehler“ x „*commission* Fehler ja/nein“ $F(2, 42) = 1.53, p = .23$, alle übrigen $F < 1$.

Bezüglich der Reaktionszeitaufgabe zeigte die Analyse, basierend auf der experimentellen Gruppeneinteilung, folgende Ergebnisse: Wie bereits in Experiment I fielen die Reaktionszeiten der Probanden in den Fehlerphasen durchweg höher aus ($M = 1.33, SD = 0.41$) als in den fehlerfreien Phasen ($M = 1.06, SD = 0.40$), Haupteffekt „Fehlerfrei vs. Fehler“ $F(1, 22) = 10.60, p = .004$. Darüber hinaus wurden keine statistisch bedeutsamen Effekte gefunden, Interaktionseffekte „Fehlerblock“ x „Fehlerfrei vs. Fehler“ $F(2, 42) = 1.82, p = .17$, „Fehlerblock“ x „Fehlerfrei vs. Fehler“ x „Gruppe“ $F(2, 42) = 1.17, p = .32$, alle übrigen $F < 1$.

Eine Einteilung der Gruppen danach, ob mindestens ein *omission* Fehler begangen worden war, ergab ebenfalls einen Leistungsunterschied zwischen fehlerfreien Phasen und Fehlerphasen (siehe Abb. 35).

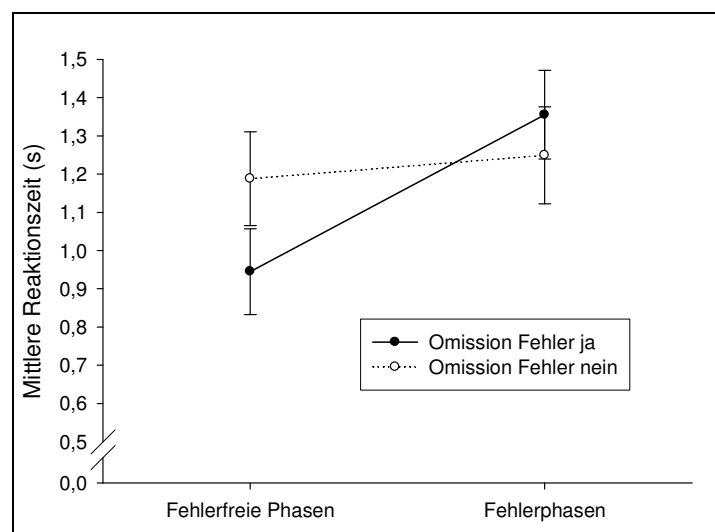


Abb. 35: Exp. II Leistung in der Reaktionszeitaufgabe in fehlerfreien vs. Fehlerphasen (*omission* Fehler ja/nein)

Für Probanden, die einen *omission* Fehler begingen, fiel dieser extremer aus: Sie zeigten im Vergleich zu Probanden, die den Fehler vermeiden konnten, bessere Leistungen in den fehlerfrei-

en Phasen, aber auch schlechtere Leistungen bei Vorliegen einer Fehlfunktion. Entsprechend ergab die statistische Analyse einen signifikanten Haupteffekt für den Faktor „Fehlerfrei vs. Fehler“, $F(1, 20) = 10.0, p = .005$, und einen Interaktionseffekt „Fehlerfrei vs. Fehler“ x „*commission* Fehler ja/nein“, $F(1, 20) = 5.45, p = .03$. Darüber hinaus wurden keine Effekte beobachtet, Interaktionseffekt „Fehlerblock“ x „Fehlerfrei vs. Fehler“ $F(2, 40) = 1.80, p = .18$, für alle übrigen Effekte $F < 1$.

Bei einer Gruppeneinteilung nach „*commission* Fehler ja/nein“ zeigte die Auswertung wiederum den Leistungsunterschied zwischen fehlerfreien Phasen und Fehlerphasen (siehe Abb. 36). Zudem aber zeichneten sich die Probanden, die einen *commission* Fehler begingen, durch höhere durchschnittliche Reaktionszeiten und damit schlechtere Leistungen aus, Haupteffekt „*commission* Fehler ja/nein“ $F(1, 21) = 4.28, p = .05$, Haupteffekt „Fehlerfrei vs. Fehler“ $F(1, 22) = 7.84, p = .01$. Alle übrigen geprüften Effekte erwiesen sich als statistisch nicht bedeutsam, Interaktionseffekt „Fehlerblock“ x „Fehlerfrei vs. Fehler“ x „*commission* Fehler ja/nein“ $F(2, 42) = 1.81, p = .18$, alle übrigen $F < 1$.

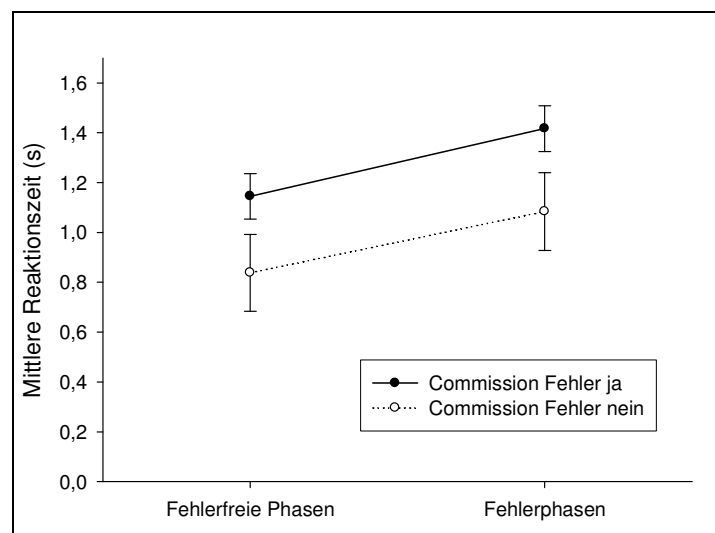


Abb. 36: Exp. II: Leistung in der Reaktionszeitaufgabe (*commission* Fehler ja/nein)

4.5.8 Wahrgenommene Zuverlässigkeit des Assistenzsystems und Einschätzung der eigenen Leistung

4.5.8.1 Informationsgruppe vs. Erfahrungsgruppe

Im Folgenden soll zunächst die von den Probanden wahrgenommene Zuverlässigkeit der Automationsfunktionen berichtet und anschließend auf die leistungsbezogenen Selbsteinschätzungen eingegangen werden.

Die Zuverlässigkeit der Automation wurde bezüglich dreier Aspekte von AFIRA (Alarmfunktion, Diagnosefunktion, Handlungsempfehlungen) einmal nach der Bedingungsmanipulation und einmal nach der Testphase erhoben. Die Auswertung der Daten erfolgte über $2(\text{Gruppe}) \times 2(\text{Messzeitpunkt})$ faktorielle ANOVAs mit Messwiederholung.

Die wahrgenommene Zuverlässigkeit der Alarmfunktion ist in Abhängigkeit der experimentellen Manipulation in Abbildung 37 dargestellt. Probanden der Erfahrungsgruppe schätzten die Zuverlässigkeit der Alarmfunktion deutlich geringer ($M = 2.73$, $SD = .90$) ein als Probanden der Informationsgruppe ($M = 3.79$, $SD = .62$). Nach der Testphase fielen die Zuverlässigkeitseinschätzungen der Informationsgruppe jedoch deutlich ab und erreichten in etwa das Niveau der Erfahrungsgruppe. Statistisch manifestierte sich dies in einem Haupteffekt „Gruppe“ $F(1, 21) = 11.00$, $p = .003$, einem Messwiederholungseffekt, $F(1, 21) = 24.52$, $p < .001$, sowie einem signifikanten Interaktionseffekt, $F(1, 21) = 33.14$, $p < .001$.

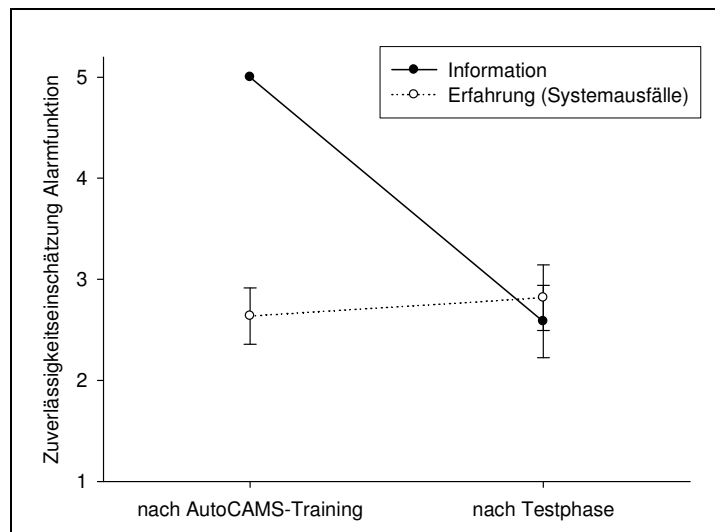


Abb. 37: Exp. II: Wahrgenommene Zuverlässigkeit der Alarmfunktion von AFIRA (experimentelle Gruppen)

Die wahrgenommene Zuverlässigkeit der Diagnosefunktion von AFIRA ist in Abbildung 38 dargestellt. Zu Beginn der Testphase schätzten alle Probanden deren Zuverlässigkeit höher ein ($M = 4.57$, $SD = 0.84$) als nach der Testphase, in der alle Probanden die Erfahrung einer Fehldiagnose von AFIRA gemacht hatten ($M = 3.04$, $SD = 0.88$). Zudem fiel die Zuverlässigkeitseinschätzung der Informationsgruppe zu Beginn höher aus um dann eine stärkere Abnahme zu zeigen. Neben einem signifikanten Messwiederholungseffekt, $F(1, 21) = 60.26$, $p < .001$, ergab die varianzanalytische Auswertung folglich einen signifikanten Interaktionseffekt, $F(1, 21) = 6.70$, $p = .02$. Ein Haupteffekt für den Faktor „Gruppe“ wurde nicht beobachtet, $F(1, 21) = 2.14$, $p = .16$.

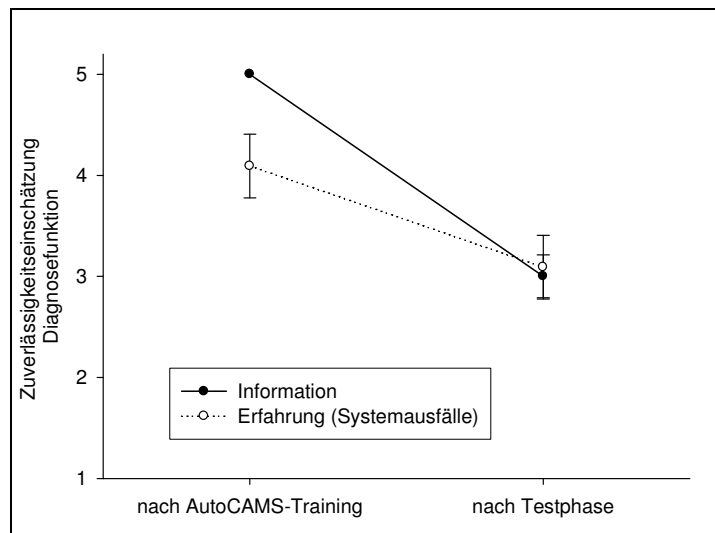


Abb. 38: Exp. II: Wahrgenommene Zuverlässigkeit der Diagnosefunktion von AFIRA (experimentelle Gruppen)

Mit Blick auf die wahrgenommene Zuverlässigkeit der Handlungsempfehlungen von AFIRA zeigt sich ein ähnliches Befundmuster (siehe Abb. 39): Auch hier ergab die Auswertung signifikante Effekte für den Faktor „Messzeitpunkt“, $F(1, 21) = 10.77, p = .004$, und die Interaktion „Gruppe“ x „Messzeitpunkt“, $F(1, 21) = 5.16, p = .03$, wobei kein Haupteffekt „Gruppe“ beobachtet wurde, $F < 1$.

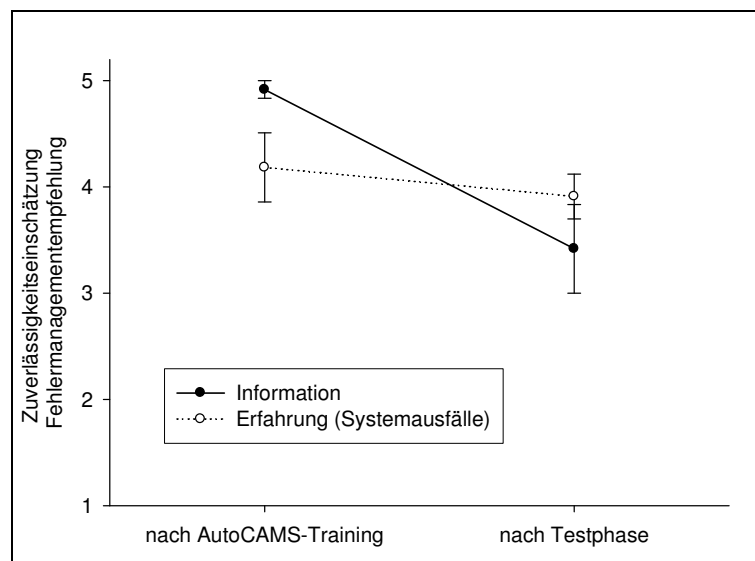


Abb. 39: Exp. II: Wahrgenommene Zuverlässigkeit der Fehlermanagementempfehlungen von AFIRA (experimentelle Gruppen).

Die leistungsbezogene Selbsteinschätzung der Probanden wurde insgesamt dreimal, vor der Bedingungsmanipulation, nach der Bedingungsmanipulation und nach der Testphase, erfasst. Dem Rechnung tragend erfolgte die statistische Auswertung hier über 2(Gruppe) x 3(Messzeitpunkt) ANOVAs mit Messwiederholung. Abbildung 40 zeigt die Selbsteinschätzung der Probanden bezüglich der manuellen Fehlerdetektion.

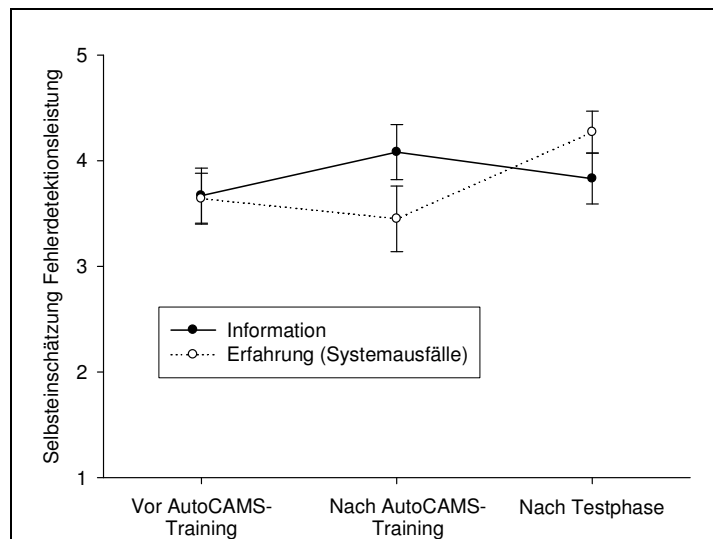


Abb. 40: Exp. II: Selbsteinschätzung der Fehlerdetektionsleistung (experimentelle Gruppen)

Es zeigten sich weder Unterschiede zwischen den Gruppen noch im zeitlichen Verlauf, Haupteffekt „Gruppe“ $F < 1$, Haupteffekt „Messzeitpunkt“ $F(2, 42) = 1.53, p = .23$, Interaktionseffekt $F(2, 42) = 2.57, p = .09$.

In Abbildung 41 ist die Selbsteinschätzung der Probanden der manuellen Fehlerdiagnoseleistung dargestellt. Während in der Erfahrungsgruppe die Einschätzung der eigenen Diagnoseleistung nach der Bedingungsmanipulation abnahm, zeigte sich für die Informationsgruppe ein gegenteiliger Effekt. Nach der Testphase fielen die Einschätzungen beider Gruppen jedoch wieder nahezu identisch aus. Signifikant wurde der Interaktionseffekt „Gruppe“ x „Messzeitpunkt“, $F(2, 42) = 4.15, p = .02$. Haupteffekte wurden für die Faktoren „Gruppe“ und „Messzeitpunkt“ jedoch nicht beobachtet, $F(1, 21) = 1.14, p = .30$ und $F < 1$.

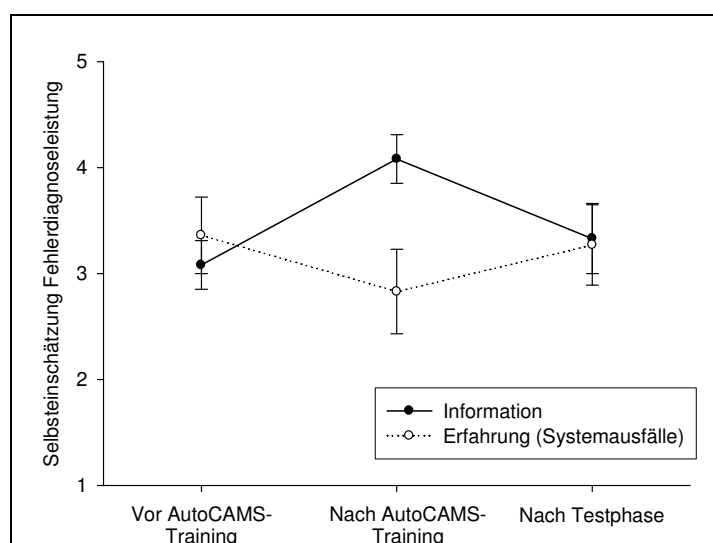


Abb. 41: Exp. II: Selbsteinschätzung der Fehlerdiagnoseleistung (experimentelle Gruppen)

Die Leistungseinschätzung bezüglich des Fehlermanagements ist in Abbildung 42 ersichtlich: Hier unterschieden sich die Informations- und die Erfahrungsgruppe, wobei die Informationsgruppe ihre Leistung höher einschätzte (3.83, $SD = 0.54$) als die Erfahrungsgruppe (3.24, $SD = 0.68$) aufwies. Entsprechend weist die statistische Analyse einen signifikanten Haupteffekt „Gruppe“ aus, $F(1, 21) = 5.32, p = .03$. Weder der Messwiederholungseffekt noch der Interaktionseffekt wurden signifikant, $F(2, 42) = 2.24, p = .12$ bzw. $F(2, 42) = 1.82, p = .17$.

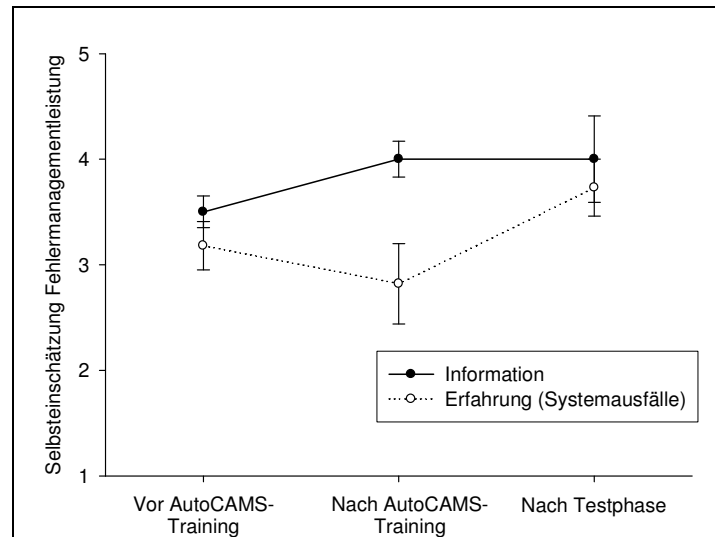


Abb. 42: Exp. II: Selbsteinschätzung der Fehlermanagementleistung (experimentelle Gruppen)

4.5.8.2 Omission Fehler ja vs. nein

Probanden, die mindestens einen *omission* Fehler begingen, unterschieden sich in keiner der subjektiven Variablen von Probanden, die diesen Fehler vermeiden konnten, für alle Haupteffekte „*omission* Fehler ja/nein“ $F < 1.51$. Signifikant wurden lediglich wiederum die bereits beschriebenen Messwiederholungseffekte, auf eine erneute Beschreibung wird – mangels Erkenntniszugewinns – an dieser Stelle verzichtet.

4.5.8.3 Commission Fehler ja versus nein

Probanden, die einen *commission* Fehler begingen, schätzten die Zuverlässigkeit der Alarmfunktion von AFIRA geringer ein ($M = 3.06, SD = 0.97$) als Probanden, die die Fehldiagnose rechtzeitig entdeckten ($M = 3.92, SD = 0.38$), Haupteffekt „*commission* Fehler ja/nein“ $F(1, 21) = 4.38, p = .05$ (siehe Abb. 43).

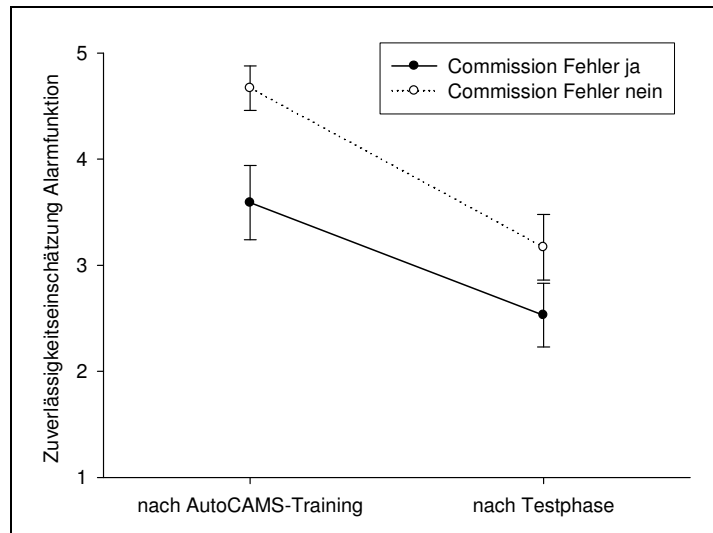


Abb. 43: Exp. II: Wahrgenommene Zuverlässigkeit der Alarmfunktion von AFIRA (*commission* Fehler ja/nein)

Zudem zeigte sich, wie schon bei der experimentellen Gruppeneinteilung, auch hier ein signifikanter Messwiederholungseffekt, $F(1, 21) = 9.77, p = .005$, ein Interaktionseffekt wurde jedoch nicht beobachtet, $F < 1$

In Bezug auf die wahrgenommene Zuverlässigkeit der Diagnosefunktion als auch der Fehlermanagementempfehlungen von AFIRA wurde weder ein Unterschied zwischen den Gruppen, $F < 1$ bzw. $F(1, 21) = 1.51, p = .23$, noch ein Interaktionseffekt, beide $F < 1$, beobachtet. Signifikant wurden für beide funktionalen Aspekte von AFIRA, wie auch bei der experimentellen Gruppeneinteilung, der Haupteffekt „Messzeitpunkt“, der sich in einer Abnahme der Zuverlässigkeitseinschätzungen zum Ende der Testphase ausdrückt, für die Diagnosefunktion $F(1, 21) = 32.46, p < .001$, für die Fehlermanagementempfehlung $F(1, 21) = 5.33, p = .03$.

Mit Blick auf die wahrgenommene Zuverlässigkeit der Fehlerdetektion ergab die Auswertung keine signifikanten Effekte, Messwiederholungseffekt $F(2, 42) = 1.17, p = .32$, Haupteffekt „*commission* Fehler ja/nein“ $F(1, 21) = 1.45, p = .24$, Interaktionseffekt $F < 1$.

Ein anderes Bild zeigte die Auswertung jedoch bei der Einschätzung der eigenen Fehlerdiagnoseleistung (siehe Abb. 44): Probanden, die einen *commission* Fehler begingen, schätzten ihre Leistung durchweg niedriger ein ($M = 3.08, SD = 0.70$) als Probanden, die die Fehldiagnose von AFIRA rechtzeitig erkannten ($M = 4.06, SD = 0.53$), Haupteffekt „*commission* Fehler ja/nein“ $F(1, 21) = 9.54, p = .006$. Weder ein Messwiederholungseffekt noch ein Interaktionseffekt wurden beobachtet, $F(2, 42) = 1.35, p = .27$ bzw. $F(2, 42) = 2.54, p = .09$.

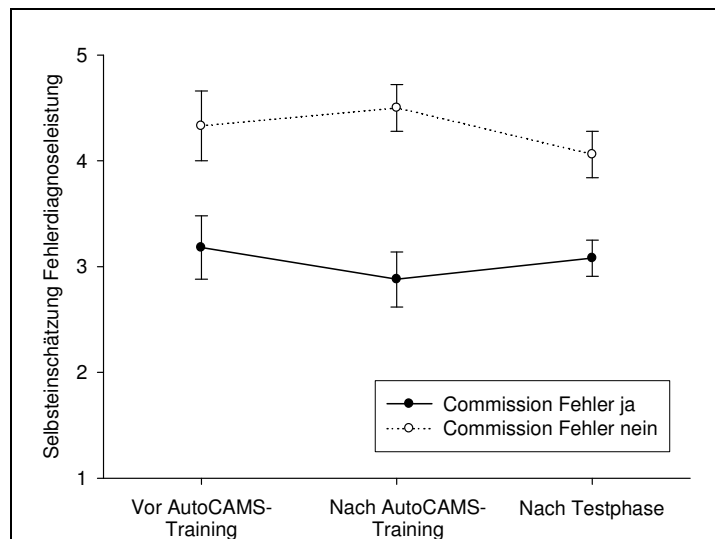


Abb. 44: Exp. II: Selbsteinschätzung der Fehlerdiagnoseleistung (*commission* Fehler ja/nein)

Die Einschätzung der eigenen Fehlermanagementleistung war in den beiden Gruppen ebenfalls unterschiedlich ausgeprägt (siehe Abb. 45): Wieder vertrauten Probanden, die einen *commission* Fehler begingen, ihren eigenen Fertigkeiten signifikant weniger ($M = 3.33$, $SD = 0.62$) als Probanden, die den Fehler vermeiden konnten ($M = 4.17$, $SD = 0.35$), Haupteffekt „*commission* Fehler ja/nein“ $F(2, 42) = 9.46$, $p = .006$. Signifikant wurde zudem der Messwiederholungseffekt, $F(2, 42) = 3.38$, $p = .04$. Ein Interaktionseffekt wurde nicht gefunden, $F(2, 42) = 1.70$, $p = .19$.

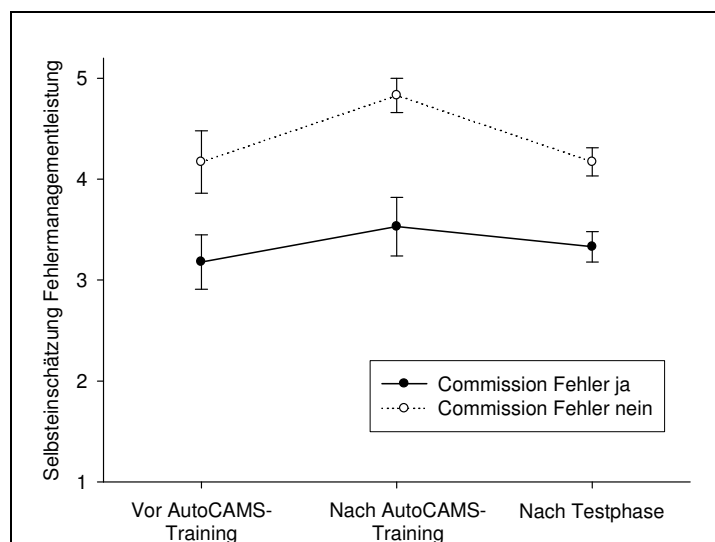


Abb. 45: Exp. II: Selbsteinschätzung der Fehlermanagementleistung (*commission* Fehler ja/nein)

4.6 DISKUSSION EXPERIMENT II

Die Ergebnisse des zweiten Experiments unterstreichen einmal mehr die Bedeutung eines übersteigerten Vertrauens und daraus resultierender Verhaltensweisen als ernstzunehmendes Problem im Umgang mit automatisierten Assistenzsystemen. 80 % der Probanden, die nur über die Eventualität von Automationsfehlern informiert worden waren, und immerhin 18 % der Probanden, die bereits im Training Systemausfälle erfahren hatten, begingen beim ersten Systemausfall in der Testphase einen *omission* Fehler. Zudem folgten 74 % der Probanden fälschlicher Weise der vom Assistenzsystem vorgeschlagenen Fehldiagnose und begingen somit einen *commission* Fehler. Darüber hinaus zeigten, wie bereits in Experiment I, alle Probanden *complacency* in gewissem Umfang gegenüber der Diagnosefunktion. Die Fehlererfahrung im Training konnte dabei nur das Überwachungsverhalten in fehlerfreien Phasen verbessern und *omission* Fehler reduzieren. *Complacency* gegenüber der Diagnosefunktion des Assistenzsystems und das Risiko von *commission* Fehlern blieben durch die Fehlererfahrung dagegen völlig unbeeinflusst.

Im Folgenden werden die Ergebnisse im Detail diskutiert, wobei die aufgestellten Hypothesen der Gliederung der Befunde dienen sollen.

Mit Hypothese 1 wurde erwartet, dass Probanden, die im Training Ausfälle eines Assistenzsystems erfahren, in fehlerfreien Phasen mehr Informationen über den Systemzustand abrufen als Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert werden. Die Ergebnisse stützen diese Annahme: Probanden der Erfahrungsgruppe tätigten im Mittel 5 Informationsabrufe mehr als Probanden der Informationsgruppe, die pro Minute durchschnittlich 14 Mal Informationen abriefen. Mit Blick auf den Anteil relevanter Informationsabrufe unterschieden sich die Gruppen dagegen nicht.

Auch Hypothese 2 wird durch die Ergebnisse gestützt: Beim ersten Systemausfall begingen die Probanden der Erfahrungsgruppe signifikant weniger *omission* Fehler als Probanden der Informationsgruppe. Letztere übersahen die Fehlfunktion beim ersten Ausfall des Assistenzsystems mehr als viermal so oft. Der Unterschied zwischen den Gruppen verschwand jedoch beim zweiten Systemausfall: Bei diesem konnte die Informationsgruppe, die mittlerweile auch einen Automationsfehler erfahren hatte, ihre Detektionsleistung deutlich verbessern, und erreichte mit einem prozentualen Anteil von ca. 22 *omission* Fehlern das Niveau der Erfahrungsgruppe. Dies zeigt einmal mehr, dass bereits einzelne Fehlererfahrungen in starkem Maße und unmittelbar verhaltenswirksam werden können.

Die Erfahrung von Systemausfällen im Training vermochte es also, sowohl die Überwachung des Systems in vermeintlich fehlerfreien Phasen zu fördern als auch die Anzahl von *omissi-*

on Fehlern zu reduzieren. Damit zeigt sich abermals, dass Automationsfehler auch dann, wenn sie seltene Ereignisse darstellen, einen deutlichen Effekt zeitigen können (vgl. auch Lee & Moray, 1992; Dzindolet et al., 2003).

Im Falle eines spezifischen Effekts der Fehlererfahrung müsste jedoch das Verhalten in den Fehlerdiagnosephasen unbeeinflusst bleiben (Hypothese 3). Die Ergebnisse zeigen, dass dem auch so ist: Die beiden experimentellen Gruppen unterscheiden sich nicht bezüglich der Überprüfung der automatisch generierten Diagnosevorschläge, weisen also vergleichbare *complacency*-Ausprägungen gegenüber der Diagnosefunktion auf. Auch die Anzahl von *commission* Fehlern unterscheidet sich nicht zwischen den Gruppen, womit auch Hypothese 4 als bestätigt zu sehen ist.

Im Unterschied zur sehr geringen Anzahl von *commission* Fehlern in Experiment I folgten diesmal 74 % der Probanden der Fehldiagnose und schickten den vorgeschlagenen, aber falschen, Reparaturauftrag ab. Dies zeigt, dass die geringe Anzahl von *commission* Fehlern in Experiment I offenbar der Tatsache geschuldet war, dass die Fehldiagnose sehr einfach als falsch zu überführen war. Während in Experiment I bereits der Abruf einer Information genügte, um die Diagnose zu falsifizieren, waren in Experiment II vier Prüfelemente erforderlich, um den Diagnosefehler zu erkennen. Der in Experiment II gefundene höhere Prozentsatz von *commission* Fehlern ist dabei im Einklang mit früheren Studien, in denen durchweg *commission*-Fehlerraten von über 50 % gefunden wurden (Mosier et al., 1998; Mosier et al., 2001; Skitka, Mosier, Burdick & Rosenblatt, 2000).

Nach Mosier et al. (2001) können *commission* Fehler die Folge von *complacency* im Sinne einer unzureichenden Verifikation automatisch generierter Direktiven sein, aber auch trotz der zur Kenntnisnahme widersprechender Informationen auftreten. Eine explorative Datenauswertung liefert Belege für beide Fehlerarten: Die Mehrzahl der Probanden (82 %) beging einen *commission* Fehler aufgrund variierender *complacency*-Ausprägungen, rief also keine oder nicht genügend Informationen zur Überprüfung der Diagnose ab. 18 % folgten jedoch der Fehldiagnose, obwohl sie alle Informationen abgerufen hatten, um die Diagnose als fehlerhaft zu überführen.

Mit dem Ziel zu prüfen, ob die Befunde aus Experiment I einer Replikation standhalten, wurde auch in Experiment II analysiert, inwiefern sich die Probanden, die einen *commission* Fehler begingen, von den übrigen Probanden unterscheiden. Wie in Experiment I zeigten auch in Experiment II Probanden, die einen *commission* Fehler begingen, eine höhere *complacency*-Ausprägung im Vergleich zu Probanden, die den Diagnosefehler erkannten. Die Befunde stützen somit auch Hypothese 5. Darüber hinaus riefen Probanden, die einen *commission* Fehler vermeiden konnten, insgesamt fast doppelt so oft Informationen in der Fehlerdiagnose ab, wobei keine Unterschiede

bezüglich des Anteils relevanter Informationsabrufe beobachtet wurden. Zudem investierten sie über 50 % mehr Zeit in die Fehlerdiagnose als Probanden, die der Fehlerdiagnose fälschlicherweise folgten. Zusammenfassend kann also festgehalten werden, dass in Experiment II die Befunde aus Experiment I in Bezug auf den Zusammenhang von *commission* Fehlern und einem allgemein reduzierten Informationssuchverhalten in Fehlerdiagnosephasen repliziert werden konnten. Die Bedeutung dieser Befunde wird zusätzlich dadurch unterstrichen, dass sich weder zwischen den experimentellen Gruppen noch zwischen den Gruppen „*omission* Fehler ja“ vs. „*omission* Fehler nein“ Unterschiede in den erwähnten abhängigen Variablen – *complacency*, Informationsabrufe gesamt, Anteil relevanter Informationsabrufe und Fehlerdiagnosezeit – zeigten. Somit ist die unzureichende Bereitschaft, Aufwand in die Überprüfung von automatisch generierten Direktiven zu investieren, als spezifischer Wegbereiter für *commission* Fehler zu sehen. *Omission* Fehler scheinen dagegen nicht mit einem reduzierten Informationssuchverhalten in Fehlerdiagnosephasen verbunden zu sein. In dieselbe Richtung weist auch der Befund, dass die Wahrscheinlichkeit, einen *commission* Fehler zu begehen, nicht dadurch erhöht wird, bereits einen *omission* Fehler begangen zu haben.

Zusammengenommen deuten diese Ergebnisse darauf hin, dass das Verhalten in vermeintlich fehlerfreien Phasen unabhängig ist vom Verhalten in den Fehlerdiagnosephasen. Davon ausgehend, dass die Konzepte *reliance* und *compliance* (Meyer, 2001) auch auf den Kontext der Assistenzsystemnutzung übertragen werden können, lässt sich daraus Folgendes ableiten: Erstens wirken sich Auslassungsfehler mindernd auf *reliance* gegenüber der Alarmfunktion aus, was in Einklang mit den Befunden von Dixon et al. (2007; auch Dixon & Wickens, 2006) ist. Zweitens wirken sich Systemausfälle im Sinne von Auslassungsfehlern nicht auf die *compliance*-Ausprägung gegenüber der Diagnosefunktion aus. Auch dies passt zu den eben erwähnten Studien, wenngleich *compliance* dort ausschließlich gegenüber der Alarmfunktion betrachtet wurde.

Analog zur bestätigten Annahme, dass Probanden, die *commission* Fehler begehen, höhere *complacency* Ausprägungen gegenüber der Diagnosefunktion aufweisen, wurde auch erwartet, dass Probanden, die einen *omission* Fehler begehen, in fehlerfreien Phasen häufiger Informationen über den Systemzustand abrufen (Hypothese 6). Die Befunde sprechen jedoch gegen diese Hypothese. Zwischen Probanden, die einen *omission* Fehler begingen, und jenen, die die Fehlfunktionen rechtzeitig erkannten, obwohl sie von AFIRA nicht gemeldet wurden, zeigten sich keine Unterschiede, weder bezüglich der Anzahl von Informationsabrufen noch mit Blick auf den Anteil relevanter Informationsabrufe. Im Übrigen wurden auch keine Unterschiede zwischen Probanden, die einen *commission* Fehler vermeiden konnten, und Probanden, die einen solchen Fehler begingen, gefunden. Allerdings riefen erstere (*commission* Fehler nein) zu Beginn der Testphase

einen etwas höheren Anteil relevanter Informationen ab, näherten sich aber über den Verlauf der Testphase hinweg an das etwas niedrigere Ausgangsniveau der Probanden, die einen *commission* Fehler begingen, an. Somit handelt es sich auch dabei offenbar nicht um einen grundsätzlichen Unterschied zwischen Operateuren, die *commission* Fehler begehen und solchen, die sie vermeiden können.

Hypothese 7 bezog sich auf die wahrgenommene Zuverlässigkeit der Alarmfunktion des Assistenzsystems und wurde durch die Ergebnisse bestätigt: Probanden, die im Training Ausfälle des Assistenzsystems erfahren hatten, schätzen die Zuverlässigkeit der Alarmfunktion signifikant geringer ein als Probanden, die im Training nur über die Möglichkeit von Automationsfehlern informiert worden waren. Sie unterschieden sich jedoch nicht von der Informationsgruppe mit Blick auf die Zuverlässigkeitseinschätzung der Diagnosefunktion und der Handlungsempfehlungen des Assistenzsystems. Dies zeigt, dass Probanden tatsächlich über eine hohe funktionale Spezifität in der Wahrnehmung der Zuverlässigkeit verschiedener Automationsfunktionen verfügen. Aufgrund des engen Zusammenhangs von wahrgenommener Zuverlässigkeit und Vertrauen in Automation ist anzunehmen, dass diese Spezifität auch in dem den verschiedenen Teilfunktionen entgegengebrachten Vertrauen reflektiert wird.

Die Erfahrung der Systemausfälle im Training wirkte sich auch auf die Einschätzung der eigenen Leistung aus: Nach der Bedingungsmanipulation schätzen Probanden der Erfahrungsgruppe ihre eigene Kompetenz in Fehlerdiagnose und -management geringer ein. Möglicherweise spiegelt sich darin eine an der Realität relativierte Einschätzung der eigenen Leistung wider, zumal Probanden dieser Gruppe Fehlerdiagnose und -management kurz zuvor ohne Automationsunterstützung komplett manuell ausführen mussten. Dagegen nahm in der Informationsgruppe die Einschätzung der eigenen Fehlerdiagnoseleistung deutlich und die Einschätzung der eigenen Fehlermanagementleistung zumindest tendenziell zu. Vor dem Hintergrund, dass die Probanden die AFIRA-Diagnosen allerdings nicht vollständig überprüften und somit auch die Diagnoseentscheidung letztlich nicht im Detail nachvollzogen, mag sich darin eher eine Überschätzung der eigenen Leistung, denn ein realer Kompetenzzuwachs widerspiegeln.

Bei einem Vergleich der Probanden, die mindestens einen *omission* Fehler begingen, mit jenen, die beide CAMS-Fehlfunktionen rechtzeitig entdeckten, zeigten sich keine Unterschiede, auch nicht bezüglich der subjektiv wahrgenommenen Zuverlässigkeit der Alarmfunktion von AFIRA. Nahegelegen hätte hier, dass *omission* Fehler mit einer als vergleichsweise hoch wahrgenommenen Zuverlässigkeit der Alarmfunktion einhergehen. Ebenso überrascht, dass Probanden, die einen *commission* Fehler begingen, sich mit Blick auf die Zuverlässigkeitseinschätzung der Diagnosefunktion von AFIRA nicht von jenen unterschieden, die die Fehldiagnose entdeckten.

Stattdessen schätzten erstere (*commission* Fehler: ja) die Zuverlässigkeit der Alarmfunktion geringer ein, wenngleich sich dies nicht, wie bereits beschrieben, in einem umfangreicheren Informations-suchverhalten in fehlerfreien Phasen widerspiegelte. Allerdings waren *commission* Fehler mit einer geringeren Einschätzung der eigenen Fehlerdiagnose- und Fehlermanagementkompetenz verbunden. Das legt den Schluss nahe, dass die Probanden nicht etwa aufgrund eines – absolut gesehen – überhöhten Vertrauens in die Automation der Fehldiagnose folgten, sondern weil sie ihre eigene Urteilskompetenz für geringer erachteten. Dieses Ergebnis steht im Widerspruch zu einem Befund von Skitka et al. (1999), dem zufolge die Schwierigkeit der infrage stehenden Teilaufgabe von den Probanden umso geringer eingeschätzt wurde, je höher die *commission* Fehlerrate ausfiel. Allerdings schätzten die Probanden in der hier vorliegenden Arbeit nicht die Schwierigkeit der manuellen Fehlerdiagnose, sondern ihre eigene Leistung dabei ein, sodass ein direkter Vergleich letztlich nicht möglich ist.

Von Interesse war auch bei Experiment II die *return-to-manual*-Leistung. Untersucht wurde, ob sich die Probanden bei den beiden Fehlfunktionen, für die ein Systemausfall simuliert wurde, in ihrer Fehleridentifikationsleistung unterscheiden. Im Vergleich zum Training, in dem die Probanden wie auch bei den Systemausfällen die Fehlfunktionen manuell entdecken, identifizieren und beheben mussten, zeigte sich für beide experimentellen Gruppen ein deutlicher Trainingsgewinn. Somit trat, wie schon in Experiment I, weder für die Erfahrungsgruppe noch für die Informationsgruppe ein Fertigkeitsverlust auf. Ferner fanden sich auch keine Unterschiede zwischen Probanden, die einen *omission* bzw. *commission* Fehler begingen, und jenen, die den jeweiligen Fehler vermeiden konnten.

Für die beiden Zusatzaufgaben, die prospektive Gedächtnis- und die Reaktionszeitaufgabe, ergab die Auswertung folgende Ergebnisse: Für beide Aufgaben fiel die Leistung in fehlerfreien Phasen besser aus als bei Vorliegen einer Fehlfunktion. Zudem zeigte sich bei der prospektiven Gedächtnisaufgabe ein Trainingseffekt über die Zeit, der jedoch bei der Reaktionszeitaufgabe nicht beobachtet wurde. In keiner der beiden Aufgaben wurde ein Unterschied zwischen den experimentellen Gruppen festgestellt. Auch zeigten sich Probanden, die einen *omission* Fehler begingen, weitgehend jenen vergleichbar, die den Fehler vermieden: Nur bei der Reaktionszeitaufgabe zeigten sich *omission* Fehler verbunden mit einem stärkeren Leistungskontrast zwischen fehlerfreien und Fehlerphasen. Anders sah das Befundmuster dagegen bei Probanden aus, die einen *commission* Fehler begingen: Sie zeigten in beiden Zusatzaufgaben durchweg schlechtere Leistungen als die Probanden, die die Fehldiagnose rechtzeitig erkannten. Im Zusammenhang mit dem geringeren Selbstvertrauen in die manuelle Fehlerdiagnose und -managementkompetenz legt

dies die Vermutung nahe, dass *commission* Fehler insbesondere von Probanden begangen werden, die im Vergleich zu den übrigen Probanden leistungsschwächer und/oder unsicherer sind.

Zusammenfassend lassen sich die folgenden vier Schlussfolgerungen aus Experiment II ziehen: Erstens liefern die Ergebnisse einen klaren Beleg für die spezifische Wirkung von Automationsfehlern. Wie erwartet, konnte durch die Erfahrung von Systemausfällen im Training das Informationssuchverhalten der Probanden in fehlerfreien Phasen intensiviert und damit eine verstärkte Überwachung der Alarmfunktion des Assistenzsystems erreicht werden. Zudem wurde der prozentuale Anteil von *omission* Fehlern beim ersten Systemausfall im Vergleich zur Informationsgruppe um über 60 % reduziert. Die Tatsache, dass dieser Gruppenunterschied beim zweiten Systemausfall verschwand, zeigt einmal mehr den unmittelbaren Effekt von Automationsfehlern. Dagegen beeinflusste die Fehlererfahrung im Training weder das Überprüfungsverhalten gegenüber der Diagnosefunktion des Assistenzsystems noch das Auftreten von *commission* Fehlern. In ähnlicher Weise war in Experiment I der analoge Effekt für Fehldiagnosen im Training gezeigt worden: Sie verbesserten zwar die Überprüfung der Diagnosefunktion des Assistenzsystems, beeinflussten jedoch nicht die Überwachung der Alarmfunktion des Systems. Offenbar wird durch die beeinträchtigte Reliabilität einer Systemfunktion nicht die Reliabilität des Gesamtsystems infrage gestellt. Dies belegt, was bereits aufgrund der Ergebnisse von Experiment I vermutet worden war: Probanden können sehr gut zwischen Systemkomponenten mit variierenden Reliabilitätsgraden differenzieren und sind somit zu einer hohen funktionalen Spezifität im Sinne von Lee und See (2004) in der Lage. Eine solche hohe funktionale Spezifität stellt eine Grundvoraussetzung für ein angemessenes Vertrauen in Automation dar und ist somit durchaus wünschenswert für die Mensch-Automation-Interaktion. Die Kehrseite besteht jedoch darin, dass mit der Erfahrung von exemplarischen Automationsfehlern nur eng umschriebene Effekte erzielt werden können.

Zweitens legen die Befunde nahe, dass *omission* und *commission* Fehler voneinander unabhängige Phänomene verkörpern. Darauf verweisen zum einen die soeben zusammengefassten differenziellen Effekte von Systemausfällen auf das Überprüfungsverhalten gegenüber der Alarm- bzw. der Diagnosefunktion des automatisierten Systems. Zum anderen war das Begehen eines der beiden Fehlertypen nicht mit einem höheren Risiko für den jeweils anderen Fehlertyp verbunden.

Drittens unterstreichen die Ergebnisse den bereits in Experiment I gefundenen Befund, dass das Begehen eines *commission* Fehlers mit einer hohen *complacency*-Ausprägung gegenüber der Diagnosefunktion einhergeht: Probanden, die einen solchen Fehler begingen, fielen bereits bei früheren Diagnosen durch eine unvollständigere Überprüfung derselben sowie ein insgesamt weniger intensives Informationssuchverhalten in Fehlerdiagnosephasen auf. Somit wurde erneut

gezeigt, dass *complacency* eine mögliche Ursache für *commission* Fehler darstellt. Zudem liefern die Ergebnisse einen ersten empirischen Beleg für die Annahme, dass es zwei verschiedene Formen von *commission* Fehlern zu unterscheiden gilt (Mosier et al., 2001): Zum einen können sie die Folge einer unzureichenden Diagnoseprüfung (*complacency*) sein, zum anderen aber auch trotz der zur Kenntnisnahme widersprechender Informationen auftreten. Die Ergebnisse weisen darauf hin, dass *commission* Fehler sehr viel häufiger (hier viermal so oft) als Folge von *complacency* auftreten, wenngleich immerhin knapp 20 % der Probanden der Fehldiagnose folgten, obwohl sie widersprechende Informationen abgerufen hatten. Darüber hinaus deuten die Befunde darauf hin, dass *commission* Fehler auch mit einem geringeren Selbstvertrauen in die eigene Fehlerdiagnose- und -managementkompetenz verbunden sein können. Dies könnte einerseits eine mögliche Erklärung dafür sein, warum der Fehldiagnose von einem Teil der Probanden trotz Kenntnis eindeutig widersprechender Informationen gefolgt wurde. Andererseits kann dies auch als Hinweis auf eine dritte mögliche Ursache von *commission* Fehlern gewertet werden. Das geringere Selbstvertrauen, verbunden mit den beobachteten schlechteren Leistungen in der prospektiven Gedächtnisaufgabe sowie der Reaktionszeitaufgabe, könnte auch Ausdruck eines insgesamt geringeren Leistungsniveaus der Probanden, die einen *commission* Fehler begingen, sein.

Viertens kann mit Blick auf die praktische Bedeutung des Experiments Folgendes festgehalten werden: Bei der Gestaltung von Trainingsprogrammen, die es zum Ziel haben, Effekte eines übersteigerten Vertrauens in Automation zu reduzieren, ist die spezifische Wirkung von Automationsfehlern zu berücksichtigen: Die Ergebnisse zeigen, dass die Erfahrung von Automationsfehlern im Training geeignet ist, solche Probleme zu reduzieren. Allerdings ist dabei nur ein Effekt für die jeweils als fehlerhaft kennen gelernte Automationsfunktion zu erwarten. Daraus folgt, dass Trainings, die Probleme eines übersteigerten Vertrauens in Automation umfassend reduzieren wollen, alle automatisierten Funktionen und die korrespondierenden antizipierbaren Fehler einbeziehen sollten.

5 GESAMTDISKUSSION

Die zentralen Anliegen der vorliegenden Arbeit waren es, einen Beitrag zu einer empirisch fundierten Klärung der Konzepte *complacency* und *automation bias* und des Bezugs der Konzepte untereinander zu leisten sowie mögliche Gegenmaßnahmen im Sinne von Automationsfehlern im Training zu untersuchen. Hierzu wurden zwei experimentelle Studien durchgeführt, wobei eine Prozesssteuerungsmikrowelt, in der die Probanden bei der Fehlerdetektion, -diagnose und -behebung durch ein Assistenzsystem unterstützt wurden, als Versuchsumgebung diente. In der ersten Studie wurde der Einfluss von Fehldiagnosen im Training auf das Überwachungs- und Überprüfungsverhalten der Probanden untersucht, während im zweiten Experiment der Einfluss von Systemausfällen im Sinne von Auslassungsfehlern während des Trainings im Fokus stand.

Im Folgenden sollen die Ergebnisse der beiden Experimente im Lichte der zu Beginn (Kap. 1.7) definierten Ziele und Fragestellungen zusammenfassend diskutiert werden.

(1) Übertragung auf einen neuen und gemeinsamen Anwendungsbereich

Die Tatsache, dass in beiden Experimenten sowohl *complacency*- als auch *automation bias*-Effekte auftraten, zeigt, dass die Phänomene auch auf den Anwendungskontext der Prozesskontrolle übertragbar sind und somit auch außerhalb des Luftfahrtkontexts Relevanz besitzen. Zudem konnte gezeigt werden, dass *complacency* – nicht nur, wie bisher angenommen, im Kontext von *monitoring*-Aufgaben – sondern auch im Umgang mit Entscheidungsassistenzsystemen ein ernst zu nehmendes Problem darstellt.

(2) Operationalisierung von *complacency* unabhängig von möglichen Verhaltenskonsequenzen

Complacency wurde definiert als unzureichende Überwachung bzw. Überprüfung automatisierter Systeme, was eng mit einem übersteigerten Vertrauen in die Automation zusammenhängt. Beide Aspekte – Überprüfung und Überwachung – wurden in der vorliegenden Studie adressiert. Ersterer wurde mit Blick auf die Diagnosefunktion des Assistenzsystems betrachtet und diente der Operationalisierung von *complacency* im engeren Sinne. Analysiert wurde, wie viele der für eine vollständige Diagnoseprüfung erforderlichen Informationsquellen von den Probanden abgerufen wurden, bevor sie einer vom Assistenzsystem vorgeschlagenen Diagnose folgten. Dieser Operationalisierungsansatz erwies sich als geeignet, um methodische Schwächen bisheriger Studien zu überwinden: *Complacency* wurde damit unabhängig von möglichen Folgen direkt über das Informationssuchverhalten erfasst. Zudem wurde so ein Vergleich mit einem „optimalen“ Informati-

onssuchverhalten (vollständige Überprüfung) ermöglicht und das Verhalten der Probanden sowohl als „unzureichend“ klassifizierbar als auch quantifizierbar.

Bezüglich des zweiten Aspekts von *complacency*, des Überwachungsverhaltens gegenüber der Alarmfunktion, konnte in der vorliegenden Arbeit kein normatives Modell „optimalen“ Informationssuchverhaltens formuliert werden, sodass keine Aussagen über eine absolut „unzureichende“ Überwachung im Sinne von *complacency* zulässig sind. Ein Vergleich der Probanden untereinander gewährte jedoch zumindest Aufschluss über die relative Ausprägung der Überwachungsintensität. Künftigen Studien muss es somit überlassen bleiben, Wege zu finden, um auch bei kontinuierlichen Überwachungsaufgaben „optimale“ Abruffrequenzen zu definieren.

Im Gegensatz zu den meisten bisherigen Studien wurde in der vorliegenden Arbeit auch das Vertrauen in Automation bzw. die Zuverlässigkeit der Automationsfunktionen als Indikator hierfür erfasst. Hierbei erwies sich eine differenzierte Abfrage der Zuverlässigkeit der verschiedenen Teilfunktionen der Automation (Experiment II) gegenüber einer Abfrage des allgemeinen Vertrauens in Automation (Experiment I) als überlegen. Während in Experiment I keine Unterschiede zwischen den experimentellen Gruppen festzustellen waren, konnte mithilfe der spezifischen Abfrage in Experiment II ein Effekt der Bedingungsmanipulation festgestellt werden: Probanden, die Systemausfälle im Training erfahren hatten, schätzten die Zuverlässigkeit der Alarmfunktion des Assistenzsystems geringer ein als Probanden, die keine solche Erfahrung gemacht hatten. Hinsichtlich der wahrgenommenen Zuverlässigkeit der Diagnosefunktion und der Empfehlungen zur Fehlerbehebung zeigten sich dagegen erwartungsgemäß keine Unterschiede zwischen den Gruppen. Dies weist darauf hin, dass Zuverlässigkeitsurteile – und damit sehr wahrscheinlich auch die Vertrauensurteile – gegenüber automatisierten Systemen durch eine hohe funktionale Spezifität gekennzeichnet sind.

(3) Analyse möglicher Maßnahmen zur Vorbeugung von complacency

Untersucht wurde einerseits der Einfluss von Fehldiagnosen, und andererseits der Einfluss von Systemausfällen auf beide Aspekte von *complacency* – die Überprüfung der Diagnosefunktion und die Überwachung der Alarmfunktion.

Wie erwartet, zeigten die Ergebnisse des ersten Experiments, dass die Erfahrung von Fehldiagnosen im Training zu einer Reduktion von *complacency* gegenüber der Diagnosefunktion führt: Im Vergleich zu Probanden, die nur über die Eventualität von Automationsfehlern informiert wurden, überprüften jene Probanden, die Fehldiagnosen erfahren hatten, spätere Diagnosen des Assistenzsystems weitaus gründlicher. Allerdings wurden keine Unterschiede zwischen den Gruppen im Informationssuchverhalten in vom Assistenzsystem als fehlerfrei indizierten Phasen fest-

gestellt. Während die Fehldiagnosen im Training also die Überprüfung der Diagnosefunktion des Assistenzsystems signifikant verbessern konnten, hatten sie keinen Einfluss auf die Überwachung der Alarmfunktion desselben. Dies steht im Widerspruch zu früheren Ergebnissen von Muir und Moray (1996), denen zufolge erwartet worden war, dass sich die Fehlererfahrung in einem Teilsystem auch auf den Umgang mit anderen Teilsystemen derselben Automation generalisiert. Der überraschende Befund wurde zum Anlass genommen, in einem zweiten Experiment zu prüfen, wie sich im Training erfahrene Fehler der Alarmfunktion im Sinne von Ausfällen des Assistenzsystems auswirken. Dabei zeigte sich, dass auch diese nur einen spezifischen Effekt zeitigten: Sie beeinflussten zwar positiv das spätere Überwachungsverhalten gegenüber der Alarmfunktion (mit Blick auf die Anzahl von Informationsabrufen pro Minute), vermochten es jedoch nicht, die Überprüfung der Diagnosefunktion zu verbessern. Gezeigt wurde damit, dass Automationsfehler zwar durchaus geeignet sind, um im Rahmen von Trainings *misuse*-Effekte zu reduzieren, doch beeinflussen sie das Verhalten von Operateuren nur gegenüber der als fehlerhaft erlebten Teilfunktion der Automation. Eine Generalisierung auf andere Teilfunktionen ist nicht zu erwarten. Dies legt den Schluss nahe, dass sich nicht nur das Vertrauen in Automation, sondern auch das Überwachungs- bzw. Überprüfungsverhalten durch eine hohe funktionale Spezifität auszeichnet. Was einerseits wünschenswert mit Blick auf die Entwicklung eines angemessenen Vertrauens in Automation (Lee & See, 2004) ist, bedeutet zugleich mehr Aufwand mit Blick auf die Gestaltung von Trainings: Dabei genügt es nicht, einzelne Automationsfehler zu adressieren, was jedoch im Gegensatz zur gängigen Praxis von Operateur-Trainings steht. Diese fokussieren oftmals ausschließlich auf den Umgang mit kompletten Systemausfällen, beziehen jedoch nicht die Möglichkeit, dass ein System zwar prinzipiell verfügbar sein, aber falsche Direktiven liefern kann, mit ein. Folglich wird Operateuren das Bild vermittelt, dass automatisierte Assistenzsysteme „Alles oder Nichts“-Systeme repräsentieren, die entweder korrekt oder aber überhaupt nicht arbeiten. Ein solcher Trainingsansatz wurde im Luftfahrtkontext bereits als mögliches *complacency*-begünstigendes Problem vermutet (Aeronautica Civil of the Republic of Columbia, 1996). Die Befunde der vorliegenden Arbeit liefern hierfür einen empirischen Beleg.

(4) Analyse möglicher Folgen von *complacency*

Als mögliche Folgen von *complacency* wurden in der vorliegenden Arbeit zum einen *automation bias*-Effekte im Sinne von *omission* und *commission* Fehlern sowie potenzielle Fertigkeitsverluste untersucht.

4a) *Commission* Fehler

Commission Fehler traten in beiden Experimenten auf. Während in Experiment I, in dem die Fehldiagnose vergleichsweise einfach zu entdecken war, nur 21 % der Probanden einen *commission*

Fehler begingen, waren es in Experiment II, in dem die Fehldiagnose nur durch eine aufwändige Prüfprozedur zu falsifizieren war, 74 % der Probanden, die der Diagnose folgten. Offenbar treten *commission* Fehler umso eher auf, je schwieriger der Fehler des Assistenzsystems zu erkennen ist. Die Fehldiagnose im Training (Experiment I) hatte dabei ebenso wenig Einfluss auf das Auftreten von *commission* Fehlern wie die Erfahrung von Systemausfällen (Experiment II). Allerdings zeigte sich in beiden Experimenten, dass sich Probanden, die einen *commission* Fehler begingen, bereits im Vorfeld durch deutlich höhere *complacency*-Ausprägungen, hier gegenüber den korrekten neun Diagnosen vor der Fehldiagnose, auszeichnen. Somit stellen *commission* Fehler eine mögliche Folge von *complacency* dar. Darüber hinaus unterschieden sich Probanden, die einen solchen Fehler begingen, in beiden Experimenten auch in weiteren Verhaltensmaßen systematisch von den übrigen Probanden: Sie riefen insgesamt weniger Informationen in der Fehlerdiagnosephase ab und investierten weniger Zeit in die Diagnoseprüfung bzw. initiierten schneller eine Fehlerbehebung. Letztgenannter Befund zeigt, dass – solange das Assistenzsystem zuverlässig arbeitet – mit einer vergleichsweise geringen Überprüfung desselben auch Nutzeneffekte im Sinne einer Zeitersparnis verbunden sind.

Jedoch machen die Ergebnisse auch deutlich, dass *complacency* bzw. eine reduzierte Überprüfungsintensität nicht die einzige Ursache für *commission* Fehler darstellt. Eine Analyse des Überprüfungsverhaltens isoliert für den jeweiligen Fehler, bei dem vom Assistenzsystem eine Fehldiagnose vorgeschlagen wurde, ergab Folgendes: In Experiment I riefen alle fünf Probanden, die der Fehldiagnose folgten, vorher falsifizierende Informationen ab, während in Experiment II die meisten Probanden den *commission* Fehler aufgrund variierender *complacency*-Ausprägungen begingen. Doch auch hier hatte immer noch ein Fünftel der Probanden eindeutig widersprechende Informationen abgerufen und folgte dennoch der Fehldiagnose. Dies belegt zum einen empirisch, was bislang nur theoretisch vermutet wurde (Mosier et al., 2001): *Commission* Fehler können die Folge einer unzureichenden Verifikation im Sinne von *complacency* sein, aber auch trotz des Abrufs widersprechender Information und damit „informiert“ auftreten. Zum anderen zeigt sich, dass bei einfach zu entdeckenden Fehlern nur noch ein vergleichsweise kleiner Teil der Probanden überhaupt einen *commission* Fehler begeht, dann aber „informiert“. Bei schwierigeren Fehlern begeht ein ebenso vergleichsweise kleiner Anteil der Probanden einen *commission* Fehler „informiert“ und zeigt damit letztlich eine Extremform des *automation bias*. Der Großteil der Probanden scheut aber offenbar den mit einer vollständigen Überprüfung verbundenen Aufwand und folgt der Fehldiagnose daher „bewusst uninformiert“.

Ein weiterer Aspekt, der bei der Genese von *commission* Fehlern eine Rolle zu spielen scheint, liegt im Selbstvertrauen der Probanden in die eigene Fähigkeit zur manuellen Ausübung der au-

tomatisierten Funktion: Offenbar folgen Operateure falschen Diagnosen eines Assistenzsystems unter anderem deshalb, weil sie – im Vergleich zu Probanden, die sich über die falsche Diagnose hinwegsetzen – ihre eigene Fehlerdiagnose- und Fehlermanagementkompetenz für geringer erachten. Mit Blick auf die Frage, ob diese subjektive Einschätzung auch objektiv ein geringeres Leistungsniveau widerspiegelt, lassen die Ergebnisse keinen eindeutigen Schluss zu: Einerseits zeichneten sich die Probanden, die einen *commission* Fehler begingen, zwar durch schlechtere Leistungen bei der Reaktionszeitaufgabe aus (in Experiment I nur in den Fehlerphasen, in Experiment II sowohl in Fehlerphasen wie auch in fehlerfreien Phasen). Andererseits zeigten sie beim Systemausfall gegen Ende der Testphase keine schlechtere Leistung bei der manuellen Fehlerdetektion und -diagnose im Vergleich zu Probanden, die der Fehldiagnose nicht folgten.

Zusammenfassend kann folgendes mit Blick auf *complacency* bezüglich der Diagnosefunktion und *commission* Fehler festgehalten werden: *Complacency* gegenüber der Diagnosefunktion kann durch die Erfahrung von Fehldiagnosen im Training reduziert, wenn auch nicht gänzlich vermieden werden. Eine mögliche, aber nicht zwingende Folge von *complacency* gegenüber der Diagnosefunktion stellen *commission* Fehler dar. Diese können jedoch auch ohne *complacency* auf Seiten des Operateurs auftreten: Auch eine vollständige Überprüfung von Fehldiagnosen schützt nicht davor, diesen dennoch zu folgen. An dieser Stelle mag dem Selbstvertrauen der Operateure in die eigene Diagnosekompetenz eine entscheidende Rolle zukommen: Ist dieses gering ausgeprägt, entscheidet sich der Operateur möglicherweise trotz widersprechender Informationen dafür, der fehlerhaften Direktive zu folgen.

4b) Omission Fehler

Omission Fehler wurden ausschließlich in Experiment II untersucht. 80 % der Probanden, die nur darüber informiert worden waren, dass Fehler des Assistenzsystems auftreten können, begingen einen solchen Fehler beim ersten Systemausfall in der Testphase. Dagegen konnte durch die Erfahrung von Systemausfällen im Training der Anteil von *omission* Fehlern auf 18 % gesenkt werden. Bei einem etwas später folgenden, zweiten Systemausfall, zeigte sich erneut die unmittelbare Wirkung von Fehlererfahrungen: Hier verschwand der Unterschied zwischen den Gruppen, d.h. dass der Anteil von *omission* Fehlern auch für die Informationsgruppe auf das Niveau der Erfahrungsgruppe sank.

Vermutet worden war, dass Probanden die einen *omission* Fehler begehen, sich auch im Vorfeld schon durch ein weniger intensives Überwachungsverhalten auszeichnen. Diese Annahme konnte nicht bestätigt werden. Zwar konnte das Überwachungsverhalten gegenüber der Alarmfunktion durch die Erfahrung von Fehlern derselben im Training verbessert werden; Probanden, die einen *omission* Fehler begingen, riefen aber während der vermeintlich fehlerfreien Phasen we-

der weniger Information insgesamt, noch einen geringeren Anteil relevanter Informationen ab. Darüber hinaus zeichnet ein Vergleich der Ergebnisse beider Experimente in Bezug auf die Überwachung der Alarmfunktion folgendes Bild: In beiden Experimenten riefen die Probanden auch in den vermeintlich fehlerfreien Phasen Informationen über den Systemzustand ab. Darauf, dass diese Abrufe zum Großteil dazu erfolgten, um Gewissheit über die Fehlerfreiheit des Systems zu erlangen, verweist der Anteil relevanter Informationsabrufe an den Gesamtabrufen von 64 % (Experiment I) bzw. 86 % (Experiment II). Allerdings wurden in Experiment I, mit durchschnittlich 7.14 Abrufen pro Minute, weniger Informationen abgerufen als in Experiment II (14.43 Informationsgruppe; 19.06 Erfahrungsgruppe). Eine mögliche Ursache hierfür könnte in den unterschiedlichen Instruktionen für Experiment I und II liegen. In beiden Experimenten wurden die Probanden darüber informiert, dass Fehler des Assistenzsystems vorkommen können. In Experiment I wurde ihnen mitgeteilt, dass daher die Diagnosen und das vorgeschlagene Fehlermanagement überprüft werden sollten. Dagegen wurde in Experiment II darauf hingewiesen, dass folglich die Parameter weiterhin überwacht werden sollten. In Betracht der Tatsache, dass in Experiment II nicht nur mehr Informationen in fehlerfreien Phasen abgerufen wurden, sondern auch in Fehlerdiagnosephasen tendenziell weniger Informationen abgerufen wurden als in Experiment I ($M = 7.28$ vs. $M = 4.45$), liegt die Annahme nahe, dass sich darin ein Instruktionseffekt andeutet. Da der Ablauf der Testphase in Experiment I und II mit Blick auf die Reihenfolge und Zeitpunkte der auftretenden Fehler keine hinreichende Vergleichbarkeit gewährleistet, ist eine statistische Prüfung jedoch anhand der vorliegenden Daten nicht möglich. Zudem legen bisherige Studien nahe, dass der Instruktion bezüglich der Automationsüberprüfung keine bedeutende Rolle zukommt. Selbst die Instruktion, dass verifiziert werden „muss“, erwies sich nicht als erfolgreich, ebenso wie während des Experiments eingeblendete „*verify*“-Hinweise es nicht vermochten, *automation misuse*-Effekte zu reduzieren (Skitka, Mosier, Burdick und Rosenblatt, 2000; Mosier et al., 2001).

Zusammenfassend lässt sich Folgendes mit Blick auf *complacency* bezüglich der Alarmfunktion und *omission* Fehler festhalten: *Complacency* gegenüber der Alarmfunktion, hier indirekt über das Überwachungsverhalten in vermeintlich fehlerfreien Phasen operationalisiert, kann durch die Erfahrung von Systemausfällen reduziert werden. Dieser Effekt scheint vermittelt über eine verminderte wahrgenommene Zuverlässigkeit der Alarmfunktion des Systems. Zudem kann durch eine solche Fehlererfahrung im Training die Anzahl späterer *omission* Fehler deutlich gesenkt werden. Allerdings weisen sich Probanden, die einen *omission* Fehler begehen, weder durch eine geringere Zuverlässigkeitseinschätzung der Alarmfunktion noch durch eine reduzierte Überwachungsintensität aus. Während sich Probanden, die *commission* Fehler begehen, in einer Vielzahl anderer Verhaltensmaße konsistent von den übrigen Probanden unterscheiden, scheinen Proban-

den, die einen *omission* Fehler begehen, keine systematischen Verhaltensabweichungen aufzuweisen. Auch zeigt sich kein Zusammenhang zwischen dem Begehen von *omission* und *commission* Fehlern. Das deutet daraufhin, dass bei *omission* Fehlern – im Vergleich zu *commission* Fehlern – in stärkerem Maße situative Aspekte sowie möglicherweise auch Personenmerkmale eine Rolle spielen könnten. Dies zu untersuchen wird jedoch Aufgabe künftiger Studien sein. Nicht auszuschließen ist, dass sich in den in Experiment II beobachteten *omission* Fehlern genau das zeigt, was Moray (2003) anhand seines Gedankenexperiments zu verdeutlichen suchte: Vielleicht waren die Probanden, die die Fehlfunktion nicht innerhalb des kritischen Zeitfensters von 20 bzw. 35 s entdeckten, genau in dieser Zeitspanne mit anderen Aufgaben befasst, wie etwa dem Ablesen und Eintragen des CO₂-Wertes oder der Verbindungsprüfung, und übersahen deshalb die Fehlfunktionen.

Offen bleibt die Frage, ob der beobachtete spezifische Effekt von Automationsfehlern auch auf andere Fehlertypen übertragbar ist. Interessant wäre in diesem Kontext eine Untersuchung der Wirkung von falschen Alarmen des Assistenzsystems. Betroffen wäre dabei – wie bei den in Experiment II untersuchten Auslassungsfehlern – die Alarmfunktion. Vor dem Hintergrund, dass dieselbe automatisierte Funktion Fehler aufweist, stellt sich die Frage, ob falsche Alarme überhaupt zu anderen Effekten als Auslassungsfehlern führen. Andererseits legen Befunde aus dem Kontext binärer Alarmsysteme sehr wohl nahe, dass falsche Alarme andere Effekte zeitigen als Auslassungsfehler (z.B. Dixon et al., 2007): Möglicherweise zeigt sich dies auch bei einer Übertragung auf den Kontext von Assistenzsystemen. So könnte vermutet werden, dass die Erfahrung falscher Alarme zu einer zeitlich verzögerten Reaktion auf Fehlermeldungen führt und somit Ausdruck in einer geringeren *compliance* gegenüber der Alarmfunktion im Sinne von Meyer (2004) findet. Ob sich falsche Alarme jedoch zugleich auch auf das Überwachungsverhalten in fehlerfreien Phasen (*reliance* sensu Meyer, 2004) und auf die *compliance* gegenüber der Diagnosefunktion auswirken, wäre empirisch zu klären.

4c) Fertigkeitsverlust

Erwartet worden war, dass *complacency* nicht nur das Auftreten von *omission* und *commission* Fehlern, sondern auch Fertigkeitsverluste begünstigt. Bereits in Experiment I zeigte sich jedoch stattdessen für alle Probanden ein deutlicher Trainingseffekt: Unabhängig von der experimentellen Bedingungsmanipulation verbesserten die Probanden ihre manuelle Fehlerdiagnosekompetenz. Zwischen dem manuellen Training zu Beginn der Untersuchung und den beiden Systemausfällen am Ende der Untersuchung, die ebenfalls eine manuelle Fehlerdiagnose und -behebung erforderten, reduzierten die Probanden ihre Fehlerdiagnose- bzw. Fehleridentifikationszeit um mehr als die Hälfte. Dieser Befund wurde in Experiment II repliziert. Trotz des relativ hohen

Trainingsaufwandes (ca. 4 h) profitierten die Probanden offenbar immer noch vom weiteren Training in der Testphase – wenngleich die Fehlerdiagnose und -behebung hier nicht mehr manuell erfolgen musste, sondern bis zum Systemausfall durch das Assistenzsystem unterstützt wurde. Möglicherweise war aber auch die Testphase zu kurz, um Fertigkeitsverluste sichtbar werden zu lassen (vgl. Sauer et al., 2000). Für künftige Studien kann daraus nur gelernt werden, dass deutlich umfangreichere Trainings- sowie automationsunterstützte Phasen zu realisieren sind, wenn es gilt, Fertigkeitsverluste im Umgang mit vergleichbaren Assistenzsystemen zu untersuchen.

Eine weitere Begrenzung der vorliegenden Arbeit mag darin gesehen werden, dass die gewählte Versuchsumgebung zwar im Vergleich zu anderen laborexperimentellen Aufgaben komplex, in Relation zu typischen Arbeitsumgebungen von Operateuren jedoch immer noch einfach gestaltet war. Zudem konnten in der vorliegenden Arbeit nur kurzfristige, schnell eintretende Folgen einer Automationsunterstützung untersucht werden. Anzunehmen ist, dass bei länger andauernder – ggf. jahrelanger – Erfahrung mit hoch reliablen Systemen, Probleme eines übersteigerten Vertrauens noch stärker zutage treten. Jedoch bleibt dies in Ermangelung empirischer Belege im Bereich der Mutmaßungen. Vor diesem Hintergrund ist, trotz des Bemühens um eine hohe ökologische Validität, eine Übertragung der Befunde auf reale Arbeitssituationen nur eingeschränkt möglich.

Der Forderung nach einer möglichst hohen ökologischen Validität steht jedoch die Notwendigkeit entgegen, zunächst die Grundlagen für eine wissenschaftliche Auseinandersetzung zu schaffen. Ein häufig angeprangertes Problem im Bereich der Mensch-Automation-Interaktion ist die Vermengung letztlich unterscheidbarer Phänomene (z.B. Dekker & Hollnagel, 2004). Dies mag vor dem Hintergrund, dass automationspsychologische Fragestellungen erst seit wenigen Jahren intensiv wissenschaftlich untersucht werden, verständlich sein. Jedoch ist das Bemühen um konzeptuelle Klarheit und schlüssige operationale Definitionen eine Grundvoraussetzung, um mittelfristig brauchbare wissenschaftliche Erkenntnisse für die konkrete Gestaltung zukünftiger automatisierter Mensch-Maschine-Systeme ableiten zu können. Die vorliegende Arbeit versuchte hierzu einen Beitrag zu leisten.

LITERATURVERZEICHNIS

- Aeronautica Civil of the Republic of Columbia (1996). *AA965 Cali accident report. Prepared for the WWW by Peter Ladkin, University of Bielefeld, Germany.* Verfügbar unter:
<http://sunnyday.mit.edu/accidents/calirep.html> [26.09.2007]
- Ajzen, I. & Fishbein, M. (1980). *Understanding attitudes and predicting social behavior*. Upper Saddle River, NJ: Prentice Hall.
- Bagheri, N. & Jamieson, G. A. (2004a). Considering subjective trust and monitoring behavior in assessing automation-induced “complacency.” In *Proceedings of the Human Performance, Situation Awareness and Automation Technology Conference*. Daytona Beach, FL: Embry-Riddle Aeronautical University.
- Bagheri, N. & Jamieson, G. A. (2004b). The impact of context-related reliability on automation failure detection and scanning behaviour. In *Proceedings of the 2004 IEEE International Conference on Systems, Man and Cybernetics* (S. 212-217). Piscataway, NJ: IEEE
- Bahner, J. E. & Manzey, D. (2004). Complacency – Begriffsklärung, Stand der Forschung und Implikationen für die Verlässlichkeit der Mensch-Maschine Interaktion. In M. Grandt (Hrsg.), *Verlässlichkeit der Mensch-Maschine Interaktion* (DGLR-Bericht 2004-03) (S. 35-48). Bonn: Deutsche Gesellschaft für Luft- und Raumfahrt.
- Bailey, N. R. & Scerbo, M. W. (2007). Automation-induced complacency for monitoring highly reliable systems: the role of task complexity, system experience, and operator trust. *Theoretical Issues in Ergonomics Science*, 8, 321-348.
- Bainbridge, L. (1983). Ironies of automation. *Automatica*, 19, 775-770.
- Barber, B. (1983). *The logic and limits of trust*. New Brunswick, NJ: Rutgers University Press.
- Beck, H. P., Dzindolet, M. T. & Pierce, L. G. (2002). Applying a decision-making model to understand misuse, disuse and appropriate automation use. In E. Salas (Hrsg.), *Advances in Human Factors and Cognitive Engineering: Automation* (Vol. 2, S. 37-78). Amsterdam: JAI.
- Beck, H. P., Dzindolet, M. T. & Pierce, L. G. (2007). Automation usage decisions: controlling intent and appraisal errors in a target detection task. *Human Factors*, 49, 429-437.
- Billings, C. E., Lauber, J. K., Funkhouser, H., Lyman, G. & Huff, E. M. (1976). *NASA aviation safety reporting system* (Tech Memo. No. TM-X-3445). Moffet Field, CA: NASA Ames Research Center.

- Breznitz, S. (1983). *Cry-wolf: The psychology of false alarms*. Mahwah, NJ: Erlbaum.
- Comstock, J. R. & Arnegard, R. J. (1992). *The multi-attribute task battery for human operator workload and strategic behavior research* (Tech. Memo. No. 104174). Hampton, VA: NASA Langley Research Center.
- Degani, A. (2003). *Taming bal: Designing interfaces beyond 2001*. New York: Palgrave MacMillian.
- Dekker, S. & Hollnagel, E. (2004). Human factors and folk models. *Cognition, Technology & Work*, 6(2), 79-86.
- Deutsch, M. (1960). The effect of motivational orientation upon trust and suspicion. *Human Relations*, 13, 123-139.
- De Vries, P., Midden, C. & Bouwhuis, D. (2003). The effects of errors on system trust, self-confidence, and the allocation of control in route planning. *International Journal of Human-Computer Studies*, 58, 719-735.
- Dijkstra, J. J., Liebrand, W. B. G. & Timminga, E. (1998). Persuasiveness of expert systems. *Behaviour and Information Technology*, 17, 155-163.
- Dixon, S. R. & Wickens, C. D. (2006). Automation reliability in unmanned aerial vehicle control: A reliance-compliance model of automation dependence in high workload. *Human Factors*, 48, 474-486.
- Dixon, S. R., Wickens, C. D. & McCarley, J. (2007). On the independence of compliance and reliance: Are automation false alarms worse than misses? *Human Factors*, 49, 564-572.
- Domeinski, J., Wagner, R., Schöbel, M. & Manzey, D. (2007). Human redundancy in automation monitoring. Effects of social loafing and social compensation. In *Proceedings of the Human Factors and Ergonomics Society 51st Annual Meeting* (S. 587-591). Santa Monica, CA: Human Factors and Ergonomics Society.
- Dzindolet, M. T., Peterson, S. A., Pomranky, R. A., Pierce, L. G. & Beck, H. P. (2003). The role of trust in automation reliance. *International Journal of Human-Computer Studies*, 58, 697-718.
- Dzindolet, M. T., Pierce, L. G., Beck, H. P. & Dawe, L. A. (1999). Misuse and disuse of automated aids. In *Proceedings of the Human Factors Society 43rd Annual Meeting* (S. 339-343). Santa Monica, CA: Human Factors and Ergonomics Society.
- Dzindolet, M. T., Pierce, L. G., Beck, H. P. & Dawe, L. A. (2002). The perceived utility of human and automated aids in a visual detection task. *Human Factors*, 44, 79-94.

- Elepfandt, M., Bahner, J. E. & Manzey, D. (2007). Automation Bias und Complacency: Der Einfluss von Systemausfällen im Training auf die Überwachung einer Automation. In M. Rötting, G. Wozny, A. Klostermann & J. Huss (Hrsg.), *Fortschritt-Berichte VDI Reihe 22 Nr. 25: Prospektive Gestaltung von Mensch-Technik-Interaktion* (S. 323-326). Düsseldorf: VDI
- Endsley, M. R. (1995). Toward a theory of situation awareness in dynamic Systems. *Human Factors* 37, 32-64.
- Endsley, M. R., Bolté, B. & Jones, D. B. (2003). *Designing for situation awareness. An approach to user-centered design*. London: Taylor & Francis.
- Feuerberg, B. V., Bahner, J. E. & Manzey, D. (2005). Interindividuelle Unterschiede im Umgang mit Automation – Entwicklung eines Fragebogens zur Erfassung des Complacency-Potenzials. In L. Urbas & C. Steffens (Hrsg.), *Fortschritt-Berichte VDI Reihe 22 Nr. 22: Zustanderkennung und Systemgestaltung, 6. Berliner Werkstatt Mensch-Maschine-Systeme* (S. 199-202). Düsseldorf: VDI.
- Frey, B. & Schulz-Hardt, S. (1997). Eine Theorie der gelernten Sorglosigkeit. In H. Mandl (Hrsg.), *Bericht über den 40. Kongress der Deutschen Gesellschaft für Psychologie* (S. 604-611). Göttingen: Hogrefe.
- Funk, K. & Lyall, B. (2000). *A comparative analysis of flightdecks with varying levels of automation* (Federal Aviation Administration, Grant 93-G-039, Phase 1 final report). Washington, DC: FAA.
- Funk, K., Lyall, B., Wilson, J., Vint, R., Niemczyk, M., Suroteguh, C. & Owen, G. (1999). Flight deck automation issues. *International Journal of Aviation Psychology*, 9, 109-123.
- Gempler, K. S. & Wickens, C. D. (1998). *Display of predictor reliability on a cockpit display traffic information* (Technical Report ARL-98-6/ROCKWELL-98-1). Savoy, IL: University of Illinois, Aviation Research Lab.
- Hart, S. G. & Staveland, L. E. (1988). Development of NASA-TLX (Task Load Index): Results of empirical and theoretical research. In P. A. Hancock & N. Meshkati (Hrsg.), *Human Mental Workload* (S. 139-183). North-Holland: Elsevier.
- Hauß, Y. & Timpe, K.-P. (2000). Automatisierung und Unterstützung in Mensch-Maschine-Systemen. In K.-P. Timpe, T. Jürgensohn & H. Kolrep (Hrsg.), *Mensch-Maschine-Systemtechnik: Konzepte, Modellierung, Gestaltung, Evaluation* (S. 41-62). Düsseldorf: Symposium.

- Hockey, G. R. J. & Maule, A. J. (1995). Unscheduled manual interventions in automated process control. *Ergonomics*, 38, 2504-2524.
- Hockey, G. R. J., Wastell, D. G. & Sauer, J. (1998). Effects of sleep deprivation and user interface on complex performance: a multilevel analysis of compensatory control. *Human Factors*, 40, 233-253.
- Hüper, A.-D. (2005). *Complacency als Verhalten: Einfluss von Reliabilitätserfahrung und Reliabilitätsinformation während des Trainings auf die Überwachung einer Automation*. Unveröffentlichte Diplomarbeit, Technische Universität Berlin, Deutschland.
- Hüper, A.-D., Bahner, J. E. & Manzey, D. (2005). Complacency-Effekte im Umgang mit Assistenzsystemen: Der Einfluss von Reliabilitätserfahrung und -information. In L. Urbas & C. Steffens (Hrsg.), *Fortschritt-Berichte VDI Reihe 22 Nr. 22: Zustanderkennung und Systemgestaltung, 6. Berliner Werkstatt Mensch-Maschine-Systeme* (S. 219-222). Düsseldorf: VDI.
- Jian, J. Y., Bisantz, A. M. & Drury, C. G. (2000). Foundations for an empirically determined scale of trust in automated systems. *International Journal of Cognitive Ergonomics*, 4, 53-71.
- Johns, J. L. (1996). A concept analysis of trust. *Journal of Advanced Nursing*, 24, 76-83.
- Kaber, D. B. & Endsley, M. R. (1997). Out-of-the-loop performance problems and the use of intermediate levels of automation for improved control system functioning and safety. *Process Safety Progress*, 16(3), 126-131.
- Kelly, C., Boardman, M., Goillau, P. & Jeannot, E. (2003). *Guidelines for trust in future ATM systems: A literature review* (Reference No. 030317-01). Brussels: EUROCONTROL, EATMP Info-centre.
- Kern, T. (1998). *Flight discipline*. New York, NY: McGraw-Hill.
- Kramer, R. M. (1999). Trust and distrust in organizations: Emerging perspectives, enduring questions. *Annual Review of Psychology*, 50, 569-598.
- Lee, J. D. & Moray, N. (1992). Trust, control strategies and allocation of function in human-machine systems. *Ergonomics*, 35, 1243-1270.
- Lee, J. D. & Moray, N. (1994). Trust, self-confidence, and operators' adaptation to automation. *International Journal of Human-Computer Studies*, 40, 153-184.
- Lee, J. D. & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46, 50-80.

- Lerch, F. J., Prietula, M. J. & Kulik, C. T. (1997). The turing effect: the nature of trust in expert system advice. In P. J. Feltovich & K. M. Ford (Hrsg.), *Expertise in context: Human and machine* (S. 417-448). Cambridge, MA: MIT Press.
- Lewandowsky, S., Mundy, M. & Tan, G. (2000). The dynamics of trust: comparing humans to automation. *Journal of Experimental Psychology: Applied*, 6, 104-123.
- Lorenz, B., Di Nocera, F., Röttger, S. & Parasuraman, R. (2002). Automated fault-management in a simulated spaceflight micro-world. *Aviation, Space, and Environmental Medicine*, 73, 886-897.
- Lüdtke, A. & Möbus, C. (2002). Prognose von Bedienungsfehlern durch Simulation der Entstehung gelernter Sorglosigkeit bei der Pilot-Cockpit Interaktion. In M. Grandt & K.-P. Gärtner (Hrsg.), *Situation awareness in der Fahrzeug- und Prozessführung. DGLR-Bericht 2002-04* (S. 163-180). Bonn: Deutsche Gesellschaft für Luft- und Raumfahrt e.V.
- Luhmann, N. (2000). *Vertrauen: Ein Mechanismus der Reduktion sozialer Komplexität* (4. Aufl.). Stuttgart: UTB.
- Madhavan, P. & Wiegman, D. A. (2007a). Similarities and differences between human-human and human-automation trust: An integrative review. *Theoretical Issues in Ergonomics Science*, 8(4), 277-301.
- Madhavan, P. & Wiegman, D. A. (2007b). Effects of information source, pedigree, and reliability on operator interaction with decision support systems. *Human Factors*, 49, 773-785
- Madhavan, P., Wiegmann, D. A. & Lacson, F. C. (2006). Automation failures on tasks easily performed by operators undermine trust in automated aids. *Human Factors*, 48, 241-256.
- Manzey, D. (in Druck). Systemgestaltung und Automatisierung. In P. Badke-Schaub, G. Hofinger & K. P. Lauche (Hrsg.), *Human Factors: Psychologie der Sicherheit*. Heidelberg: Springer.
- Manzey, D. & Bahner, J. E. (2005). Vertrauen in Automation als Aspekt der Verlässlichkeit von Mensch-Maschine-Systemen. In K. Karrer, B. Gauss & C. Steffens (Hrsg.), *Beiträge zur Mensch-Maschine-Systemtechnik aus Forschung und Praxis – Festschrift für Klaus-Peter Timpe* (S. 93-109). Düsseldorf: Symposium.
- Manzey, D. & Otte, L. (2007). The impact of automation reliability on monitoring behaviour: How adaptive are operators in interacting with automated systems? Vortrag beim HFES Europe Chapter Meeting, Braunschweig, 24-26 October 2007.

- Masalonis, A. J., Duley, J. A., Galster, S. M., Castano, D. J., Metzger, U. & Parasuraman, R. (1998). Air traffic controller trust in a conflict probe during free flight. In *Proceedings of the Human Factors and Ergonomics Society 42nd Annual Meeting* (S. 1601). Santa Monica, CA: Human Factors and Ergonomics Society.
- Mayer, R. C., Davis, J. H. & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20, 709-734.
- Meyer, J. (2001). Effects of warning validity and proximity on responses to warnings. *Human Factors*, 43, 563-572.
- Meyer, J. (2004). Conceptual issues in the study of dynamic hazard warnings. *Human Factors*, 46, 196-204.
- Moorman, C., Deshpande, R. & Zaltman, G. (1993). Factors affecting trust in market-research relationships. *Journal of Marketing*, 57, 81-101.
- Moray, N. (2003). Monitoring, complacency, scepticism and eutactic behavior. *International Journal of Industrial Ergonomics*, 31, 175-178.
- Moray, N. & Inagaki, T. (2000). Attention and complacency. *Theoretical Issues in Ergonomics Science*, 1, 354-365.
- Moray, N., Inagaki, T. & Itoh, M. (2000). Adaptive automation, trust and self-confidence in fault management of time critical tasks. *Journal of Experimental Psychology: Applied*, 6, 44-58.
- Mosier, K. L., Palmer, E. A. & Degani, A. (1992). Electronic checklists: Implications for decision making. In *Proceedings of the Human Factors Society 36th Annual Meeting* (S. 7-11). Santa Monica, CA: Human Factors and Ergonomics Society.
- Mosier, K. L. & Skitka, L. J. (1996). Human decision-makers and automated decision aids: Made for each other? In R. Parasuraman & M. Mouloua (Hrsg.), *Automation and Human Performance: Theory and Applications* (S. 201-220). Mahwah, NJ: Lawrence Erlbaum Associates.
- Mosier, K. L., Skitka, L. J., Dunbar, M. & McDonnell, L. (2001). Aircrews and automation bias: The advantages of teamwork? *International Journal of Aviation Psychology*, 11, 1-14.
- Mosier, K. L., Skitka, L. J., Heers, S. & Burdick, M. (1998). Automation bias: Decision making and performance in high-tech cockpits. *International Journal of Aviation Psychology*, 8, 47-63.
- Muir, B. (1994). Trust in automation part I: Theoretical issues in the study and human intervention in a process control simulation. *Ergonomics*, 37, 1905-1923.

- Muir, B. & Moray, N. (1996). Trust in automation. Part II. Experimental studies of trust and human intervention in a process control simulation. *Ergonomics*, 39, 429-460.
- O'Hare, D. & Roscoe, S. (1990). *Flightdeck performance: The human factor*. Ames, IA: Iowa State University Press.
- Parasuraman, R., Molloy, R. & Singh, I. L. (1993). Performance consequences of automation-induced "complacency". *International Journal of Aviation Psychology*, 3, 1-23.
- Parasuraman, R., Mouloua, M. & Molloy, R. (1996). Effects of adaptive task allocation on monitoring of automated systems. *Human Factors*, 38, 665-679.
- Parasuraman, R. & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39, 230-253.
- Parasuraman, R., Sheridan, T. B. & Wickens, C. D. (2000). A model for types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics - Part A: Systems and Humans*, 30, 286-297.
- Prinzel, L. J. (2002). *The relationship of self-efficacy and complacency in pilot-automation interaction* (Tech. Memo. No. TM-2002-211925). Hampton, VA: NASA Langley Research Center.
- Prinzel, L. J., De Vries, H., Freeman, F. G. & Mikulka, P. (2001). *Examination of automation-induced complacency and individual difference variates* (Tech. Memo. No. TM-2001-211413). Hampton, VA: NASA Langley Research Center.
- Rempel, J. K., Holmes, J. G. & Zanna, M. P. (1985). Trust in close relationships. *Journal of Personality and Social Psychology*, 49, 95-112.
- Riley, V. (1996). Operator reliance on automation: Theory and data. In R. Parasuraman & M. Mouloua (Hrsg.), *Automation theory and applications. Human factors in transportation* (S. 19-35). Mahwah, NJ: Lawrence Erlbaum Associates.
- Rotter, J. B. (1967). A new scale for the measurement of interpersonal trust. *Journal of Personality*, 35, 651-665.
- Sarter, N. B., Woods, D. D. & Billings, C. E. (1997). Automation surprises. In G. Salvendy (Hrsg.), *Handbook of human factors & ergonomics* (2. Aufl., S. 1926-1943). New York, NY: Wiley.
- Sauer, J., Hockey, G. R. J. & Wastell, D. G. (2000). Effects of training on short- and long-term skill retention in a complex multiple-task environment. *Ergonomics*, 43, 2043-2064.

- Sheridan, T. B. (1987). Supervisory control. In G. Salvendy (Hrsg.), *Handbook of human factors* (S. 1243-1268). New York: Wiley.
- Sheridan, T. B. (1988). Trustworthiness of command and control systems. In *Proceedings of the Conference on Analysis, Design and Evaluation of Man-Machine Systems, 1988* (S. 427-431). Oxford: Pergamon.
- Sheridan, T. B. & Parasuraman, R. (2006). Human-automation interaction. In R. S. Nickerson (Hrsg.), *Reviews of human factors and ergonomics: Volume 1* (S. 89-129). Santa Monica, CA: Human Factors and Ergonomics Society.
- Sheridan, T. B. & Verplank, W. L. (1978). *Human and computer control of undersea teleoperators*. Cambridge, MA: Man-Machine Systems Laboratory, Department of Mechanical Engineering, MIT.
- Singh, I. L., Molloy, R. & Parasuraman, R. (1993a). Automation-induced "complacency": development of the complacency potential rating scale. *International Journal of Aviation Psychology*, 3, 111-122.
- Singh, I. L., Molloy, R. & Parasuraman, R. (1993b). Individual differences in monitoring failures of automation. *Journal of General Psychology*, 120(3), 357-373.
- Singh, I. L., Molloy, R. & Parasuraman, R. (1997). Automation induced monitoring inefficiency: role of display location. *International Journal of Human-Computer Studies*, 46, 17-30.
- Skitka, L. J., Mosier, K. L. & Burdick, M. (1999). Does automation bias decision-making? *International Journal of Human-Computer Studies*, 51, 991-1006.
- Skitka, L. J., Mosier, K. L. & Burdick, M. (2000). Accountability and automation bias. *International Journal of Human-Computer Studies*, 52, 701-717.
- Skitka, L. J., Mosier, K. L., Burdick, M. & Rosenblatt, B. (2000). Automation bias and errors: Are crews better than individuals? *International Journal of Aviation Psychology*, 10, 85-97.
- Stokes, A. & Kite, K. (1994). *Flight stress: Stress, fatigue and performance in aviation*. Adlershot, UK: Avebury Aviation.
- Thackray, R. I. & Touchstone, R. M. (1989). Detection efficiency on an air traffic task with and without computer aiding. *Aviation, Space, and Environmental Medicine*, 60, 744-748.
- Tversky, A. & Kahneman, D. (1984). Judgment under uncertainty: Heuristics and biases. *Science*, 185, 1124-1131.

- Waern, Y. & Ramberg, R. (1996). People's perception of human and computer advice. *Computers in Human Behavior*, 12, 17-27.
- Wandke, H. (2005). Assistance in human-machine interaction: a conceptual framework and a proposal for a taxonomy. *Theoretical Issues in Ergonomics Science*, 6(2), 129-155.
- Wickens, C. D. & Hollands, J. G. (2000). *Engineering psychology and human performance* (3. Aufl.). Englewood Cliffs, NJ: Prentice-Hall.
- Wiegman, D. A., Rich, A. & Zhang, H. (2001). Automated diagnostic aids: The effects of aid reliability on users' trust and reliance. *Theoretical Issues in Ergonomics Science*, 2, 352-367.
- Wiener, E. L. (1981). Complacency: Is the term useful for air safety? In *Proceedings of the 26th Corporate Aviation Safety Seminar* (S. 116-125). Denver, CO: Flight Safety Foundation, Inc.
- Wiener, E. L. (1985). Cockpit automation: In need of a philosophy. In *Proceedings of the 1985 Behavioral Engineering Conference* (S. 369-375). Warrendale, PA: Society of Automotive Engineers.
- Williams, K. D. & Karau, S. J. (1991). Social loafing and social compensation – the effects of expectations of coworker performance. *Journal of Personality and Social Psychology*, 61, 570-581.
- Zajonc, R. B. (1965). Social facilitation. *Science*, 149, 269-274.

ANHANG

Anhang A: Fehlfunktionen von CAMS: Fehlertypen, -diagnose und -management

Anhang B: Experimentelle Gruppeneinteilung: Mittelwertsvergleiche basierend auf Daten vor der Manipulation (Experiment I)

Anhang C: Handout zur Fehlerdiagnose und -behebung (Experiment I)

Anhang D: Präsentation CAMS-Training (Experiment I)

Anhang E: Präsentation AutoCAMS-Training (Experiment I)

Anhang F: Übersicht über die mit Blick auf die verschiedenen Fehlerdiagnosen relevanten und irrelevanten Informationsquellen.

Anhang G: Übersicht über relevante und irrelevante Informationsquellen in vermeintlich fehlerfreien Phasen

Anhang H: Experimentelle Gruppeneinteilung: Mittelwertsvergleiche basierend auf Daten vor der Manipulation (Experiment II)

Anhang I: Benutzeroberfläche AutoCAMS (Experiment II)

Anhang J: Checkliste für die Fehlerdetektion, -diagnose und -behebung (Experiment II)

Anhang K: Handout zur Fehlerdiagnose und -behebung (Experiment II)

Anhang L: Präsentation CAMS-Training (Experiment II)

Anhang M: Präsentation AutoCAMS-Training (Experiment II)

Fehler	Sauerstoffsystem		Stickstoffsystem	
	Fehlerdiagnose	Fehlermanagement	Fehlerdiagnose	Fehlermanagement
Leck im Ventil	▪ O₂-Abnahme in der Kabine ▪ Gasmenge, die in das Ventil strömt > Gasmenge, die das Ventil verlässt ▪ Gasmenge, die das Ventil verlässt, kann auch kleiner als „stand.“ in „Standards Gasströme“ sein	▪ Stärke des O₂-Gasstroms von „stand.“ auf „hoch“ setzen ▪ Reparaturauftrag absenden	▪ Druck-Abnahme in der Kabine ▪ Gasmenge, die in das Ventil strömt > Gasmenge, die das Ventil verlässt ▪ Gasmenge, die das Ventil verlässt, kann auch kleiner als „stand.“ in „Standards Gasströme“ sein	▪ Stärke des N₂-Gasstroms von „stand.“ auf „hoch“ setzen ▪ Reparaturauftrag absenden
	▪ O₂-Abnahme in der Kabine ▪ Gemessener Gasstrom < „stand.“ in „Standards Gasströme“ ▪ Gasmenge, die in das Ventil strömt = Gasmenge, die das Ventil verlässt	Nach erfolgter Reparatur: ▪ Zurücksetzen des Gasstroms auf „stand.“		
Offen steckendes Ventil	▪ O₂-Anstieg in der Kabine ▪ Druck steigt zum oberer Normalwert (Ablassventil verhindert weiteren Anstieg) ▪ Ausschluss eines O₂ Sensordefektes durch Testprozedur : 1. Steuerung O₂ auf Manuell „aus“ 2. nach Rückkehr der O₂-Konzentration in Normalbereich wieder „Auto“ → wenn Konzentration wieder steigt, offen steckendes Ventil → wenn Konzentration sinkt, Sensordefekt	▪ Testprozedur ▪ Manuelle Steuerung von O₂ bis Fehler behoben wurde ▪ Reparaturauftrag absenden Nach erfolgter Reparatur: ▪ Steuerung wieder auf „Auto“	▪ Druck-Anstieg BIS ZUR Obergrenze des Normalbereiches in der Kabine (durch Eingreifen des Entlüftungssystems steigt Druck nicht über Normalbereich hinaus) ▪ O₂-Kurve zeigt ein Muster aus steilen Einbrüchen und langsamen Anstiegen	▪ Manuelle Steuerung des Drucks bis Fehler behoben wurde ▪ Reparaturauftrag absenden Nach erfolgter Reparatur: ▪ Steuerung wieder auf „Auto“
	▪ O₂-Anstieg ODER -Abnahme in der Kabine ▪ O₂-Anstieg: → Testprozedur (Ausschluss eines offen steckenden O₂ Ventils, siehe oben) ▪ O₂-Abnahme: → Gemessener Gasstrom = 0	▪ Ggf. Testprozedur ▪ Manuelle Steuerung von O₂ bis Fehler behoben wurde ▪ Reparaturauftrag absenden Nach erfolgter Reparatur: ▪ Steuerung wieder auf „Auto“		
Sensor-defekt	▪ O₂-Anstieg ODER -Abnahme in der Kabine ▪ O₂-Anstieg: → Testprozedur (Ausschluss eines offen steckenden O₂ Ventils, siehe oben) ▪ O₂-Abnahme: → Gemessener Gasstrom = 0	▪ Ggf. Testprozedur ▪ Manuelle Steuerung von O₂ bis Fehler behoben wurde ▪ Reparaturauftrag absenden Nach erfolgter Reparatur: ▪ Steuerung wieder auf „Auto“	▪ Druck-Anstieg ODER -Abnahme in der Kabine ▪ Druck-Anstieg: → Druck steigt ÜBER Normalbereich hinaus (Ablassventil bleibt geschlossen) ▪ Druck-Abnahme: → Gemessener Gasstrom = 0	▪ Manuelle Steuerung des Drucks bis Fehler behoben wurde ▪ Reparaturauftrag absenden Nach erfolgter Reparatur: ▪ Steuerung wieder auf „Auto“

Experimentelle Gruppeneinteilung: Mittelwertsvergleiche basierend auf Daten vor der Manipulation (Experiment I)

	Informationsgruppe <i>M (SD)</i>	Erfahrungsgruppe <i>M (SD)</i>	<i>t</i>	<i>df</i>	<i>P</i>
Mittlere manuelle Fehler- diagnosezeit (über alle 10 Fehlfunktionen des CAMS-Trainings)	103.98 (37.60)	107.65 (45.96)	-.21	22	.83
Vertrauen in CAMS	3.5 (1.38)	3.5 (1.09)	.00	22	1.0
Selbstvertrauen bezüglich der Bedienung von CAMS	3.25 (0.75)	3.75 (0.75)	-1.63	22	.12

Handout zur Fehlerdiagnose und -behebung (Experiment I)

Fehler von CAMS

Fehlertyp	Fehlerdiagnose	Fehlermanagement
O2 Leck im Ventil	▪ Gasmenge, die in das Ventil strömt ist größer als die Gasmenge, die das Ventil verlässt ▪ Beobachtung der Gasströme für Diagnose notwendig ▪ Unterschiedliche Größe der Gasströme die den Tank verlassen und der die gemessen werden	▪ Stärke des Gasstroms von „stand“ (standard) auf „hoch“ setzen ▪ Reparaturauftrag absenden ▪ Nach erfolgter Reparatur Zurücksetzen der Stärke des Gasstroms
N2 Leck im Ventil	▪ Gasmenge, die in das Ventil strömt ist größer als die Gasmenge, die das Ventil verlässt ▪ Beobachtung der Gasströme für Diagnose notwendig ▪ Unterschiedliche Größe der Gasströme die den Tank verlassen und der die gemessen werden	▪ Stärke des Gasstroms von „stand“ (standard) auf „hoch“ setzen ▪ Reparaturauftrag absenden ▪ Nach erfolgter Reparatur Zurücksetzen der Stärke des Gasstroms
O2 Blockierung des Ventils	▪ Führt zu einer Reduktion des Gasstroms ▪ Öffnen des „Standards Gasströme“ ▪ Bei fehlerfreiem Arbeiten von CAMS müssen die Gasströme mindestens so hoch wie die dort dargestellten Werte sein	▪ Stärke des Gasstroms von „stand“ (standard) auf „hoch“ setzen ▪ Reparaturauftrag absenden ▪ Nach erfolgter Reparatur Zurücksetzen der Stärke des Gasstroms
N2 Blockierung des Ventils	▪ Führt zu einer Reduktion des Gasstroms ▪ Öffnen des „Standards Gasströme“ ▪ Bei fehlerfreiem Arbeiten von CAMS müssen die Gasströme mindestens so hoch wie die dort dargestellten Werte sein	▪ Stärke des Gasstroms von „stand“ (standard) auf „hoch“ setzen ▪ Reparaturauftrag absenden ▪ Nach erfolgter Reparatur Zurücksetzen der Stärke des Gasstroms
O2 offen steckendes Ventil	▪ Ventil lässt sich nicht mehr schließen und führt so zu einem permanenten Zufluss von O2 in die Kabine ▪ Ausschluss eines Defektes der Sauerstoffsensoren - Sauerstoffstrom unterbrechen - nach Rückkehr der Sauerstoffkonzentration in den Normalbereich Regelung wieder an die Automatik übergeben - wenn Konzentration wieder steigt, dann liegt ein offen steckendes Ventil vor - wenn Konzentration sinkt, dann liegt ein Sensordefekt vor	▪ Nach Ausschluss eines defekten Sensors Reparaturauftrag absenden ▪ Manuelle Regelung des Parameters bis Fehler behoben wurde ▪ Nach erfolgter Reparatur Einschalten der automatischen Kontrolle
N2 offen steckendes Ventil	▪ Ventil lässt sich nicht mehr schließen und führt so zu einem permanenten Zufluss von N2 in die Kabine ▪ Durch Eingreifen des Entlüftungssystems bleibt Druck an der Obergrenze des Normalbereiches ▪ O2-Kurve zeigt ein Muster aus steilen Einbrüchen und langsamen Anstiegen	▪ Manuelle Steuerung des Parameters bis Fehler behoben wurde ▪ Absenden des Reparaturauftrages ▪ Nach erfolgter Reparatur Einschalten der automatischen Kontrolle

O2 Sensordefekt	▪ Überschreitung der Grenzen des Normalbereiches wird von der Automatik nicht erkannt ▪ Nötiges Ein-/ bzw. Ausschalten der Gastzufuhr unterbleibt ▪ Dementsprechend kann ein defekter Sensor sowohl für zu hohe als auch für zu niedrige Werte des O2-Parameters verantwortlich sein ▪ Bei steigender O2-Konzentration - Ausschluss eines offen steckenden Ventils durch Testprozedur ▪ Bei sinkender O2-Konzentration - Wert auf der Gasstromanzeige muss Null sein	▪ Nach Ausschluss eines offen steckenden Ventils Reparaturauftrag absenden ▪ Manuelle Regelung des Parameters bis Fehler behoben wurde ▪ Nach erfolgter Reparatur Einschalten der automatischen Kontrolle
N2 Sensordefekt	▪ Überschreitung der Grenzen des Normalbereiches wird von der Automatik nicht erkannt ▪ Nötiges Ein-/ bzw. Ausschalten der Gastzufuhr unterbleibt ▪ Dementsprechend kann ein defekter Sensor sowohl für zu hohe als auch für zu niedrige Werte des N2-Parameters verantwortlich sein ▪ Bei steigender N2-Konzentration - Das Ablassventil bleibt geschlossen und der Druck steigt permanent ▪ Bei sinkender N2-Konzentration - Wert auf der Gasstromanzeige muss Null sein	▪ Manuelle Steuerung des Parameters bis Fehler behoben wurde ▪ Absenden des Reparaturauftrages ▪ Nach erfolgter Reparatur Einschalten der automatischen Kontrolle

Präsentation CAMS-Training (Experiment I)



Versuchsdurchführung
Arbeiten mit der Mikrowelt CAMS

Erstellt für die Diplomarbeit von
Anke-Dorothea Hüper

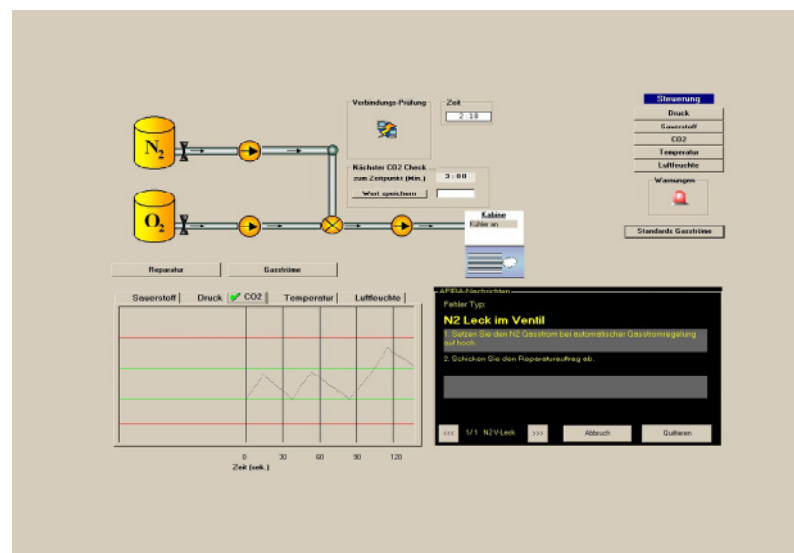
Agenda	
I.	Einführung CAMS
II.	Die fünf Subsysteme
	Sauerstoff
	Druck
	Kohlenstoffdioxid
	Temperatur
	Luftfeuchtigkeit
III.	Der Überwachungs- und Steuerungsbildschirm
	Bildschirm
	Verlaufsanzeige
	Steuerungsmenüs
III.	Aufgaben
IV.	Störungen und Fehldiagnosen



Was ist CAMS?

Cabin Air Management System

Lebenserhaltungssystem einer Raumstation



Einführung CAMS



Was ist CAMS?

- CAMS wird von einer/m Ingenieur/in in der Bodenstation einer Raumfahrtagentur überwacht
- CAMS sorgt für ausreichenden Druck, ausreichende Temperatur und genug atembare Luft in der Raumstation
- CAMS besteht aus fünf Subsystemen, die die folgenden Parameter überwachen und automatisch regeln:
 - Sauerstoff (O₂)
 - Druck
 - Kohlenstoffdioxid
 - Temperatur
 - Luftfeuchtigkeit

Einführung CAMS



Was müssen Sie tun?

- Überwachung von CAMS
- Fehlermanagement
- Ressourcen sparen
- Erfassen des Kohlenstoffdioxidstandes
- Bestätigen der Verbindung mit der Raumstation

Subsysteme allgemein



Vorstellung der Subsysteme

- Jedes Subsystem ist für einen Parameter zuständig
- Für jeden Parameter ist ein Normal- und ein Sollbereich definiert
 - Normalbereich durch grüne Begrenzung gekennzeichnet
 - Sollbereich durch rote Begrenzung gekennzeichnet
- Bei Werten außerhalb des Sollbereiches liegt eine Fehlfunktion von CAMS vor
- Werte außerhalb des Sollbereiches gefährden das Leben der Besatzung der Raumstation

Einführung in CAMS Anke-Dorothea Hüper 7

Subsysteme Sauerstoff



Sauerstoff

- Kontinuierliche Sauerstoffversorgung der Raumstation
- Anteil von Sauerstoff in der Raumstation circa 20 %
- Sollbereich: 19 – 20,5 %
- Normalbereich: 19,5 – 20 %
- Wenn der Sauerstoffanteil der Kabine unter 19,5 % fällt, wird automatisch das Ventil des Sauerstofftanks geöffnet
- Wenn der Sauerstoffanteil 20 % erreicht, wird die Sauerstoffzufuhr unterbrochen

Einführung in CAMS Anke-Dorothea Hüper 8

Subsysteme

Druck



Druck

- Konstanter Druck in der Raumstation erforderlich
- Normaler atmosphärischer Druck auf Meereshöhe circa 1000 mbar
- Regelung an Bord der Raumstation durch Stickstoffzufuhr (N₂) in die Kabine
- Sollbereich: 970 – 1040 mbar
- Normalbereich: 990 – 1020 mbar
- Wenn der Kabinendruck unter 990 mbar fällt, wird automatisch das Ventil des Stickstofftanks geöffnet
- Wenn der Kabinendruck 1020 mbar erreicht, wird das Entlüftungsventil geöffnet

Einführung in CAMS

Anke-Dorothea Hüper

9

Subsysteme

Kohlenstoffdioxid



Kohlenstoffdioxid

- Kohlenstoffdioxid entsteht durch die Atmung der Besatzungsmitglieder
- Muss aus der Kabinenluft entfernt werden
- Aufnahme durch Filter
- Sollbereich: 0,1 – 1,5 %
- Normalbereich: 0,5 – 1 %
- der Filter wird aktiv, sobald die Kohlenstoffdioxidkonzentration 1 % erreicht
- der Filter schaltet sich aus, sobald die Kohlenstoffdioxidkonzentration unter 0,5 % sinkt

Einführung in CAMS

Anke-Dorothea Hüper

10

Subsysteme

Temperatur und Luftfeuchtigkeit

TU

Temperatur

- Sollbereich: 18,5 – 23 °C
- Normalbereich: 20 – 22 °C
- Regelung durch Luftkühler und Heizung

Luftfeuchtigkeit

- Sollbereich: 36 – 44 %
- Normalbereich: 38 – 42 %
- Regelung durch Luftentfeuchter

Einführung in CAMS

Anke-Dorothea Hüper

11

Der Überwachungs- und Steuerungsbildschirm

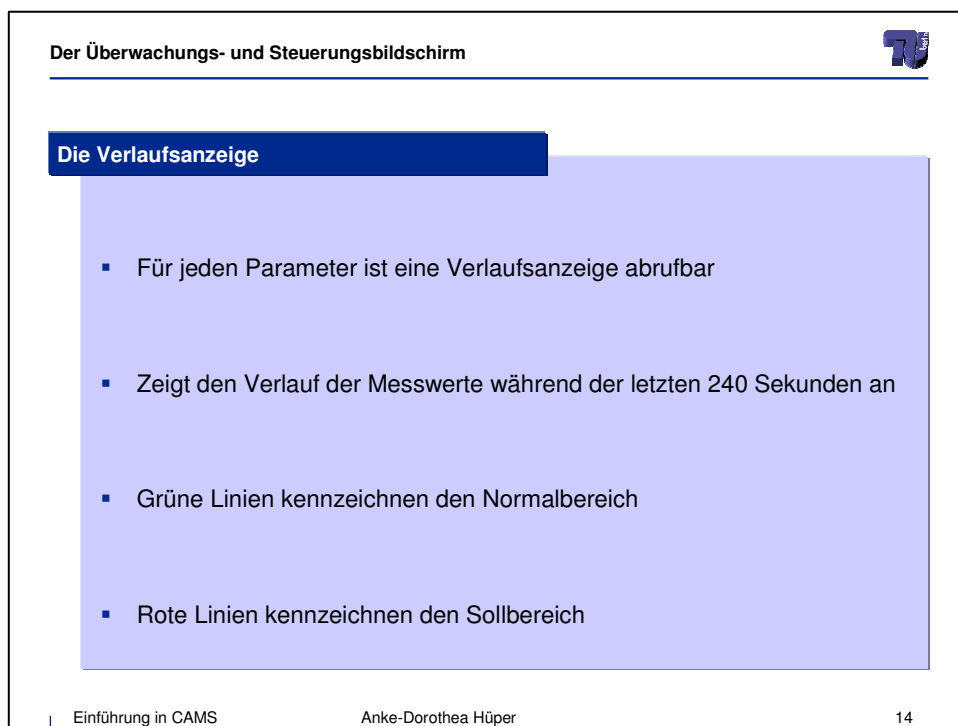
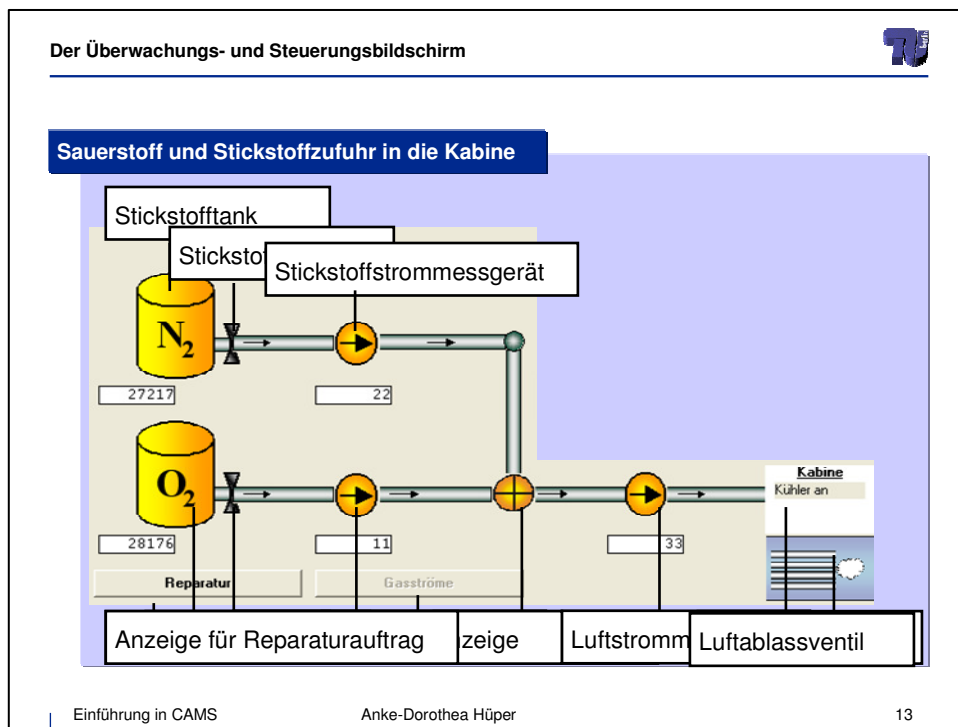
TU

The screenshot displays a complex monitoring and control interface. On the left, there are two gas cylinders labeled N₂ and O₂ with associated flow lines and valves. A central control menu is highlighted with a blue box and labeled 'Steuerungsmenüs der einzelnen Parameter'. This menu includes options for 'Reparatur', 'Gasströme', 'Sauerstoff', 'Druck', 'CO2', 'Temperatur' (with a green checkmark), and 'Luftfeuchte'. To the right, there are several status indicators: 'Verbindungs-Prüfung' with a 'Zeit' field showing '0:14', 'Nächster CO2 Check ... zum Zeitpunkt (Min.)' showing '1:00', and a 'Wert speichern' button. Below these, there's a 'Kabine' icon. Further right, a 'Steuerung' panel lists 'Druck', 'Sauerstoff', 'CO2', 'Temperatur', and 'Luftfeuchte'. Below this, a 'Warnungen' panel shows a red alarm icon. At the bottom right, there's a 'Standard Gasströme' panel. The bottom of the interface features a 'Zeit (sek.)' field showing '0'.

Einführung in CAMS

Anke-Dorothea Hüper

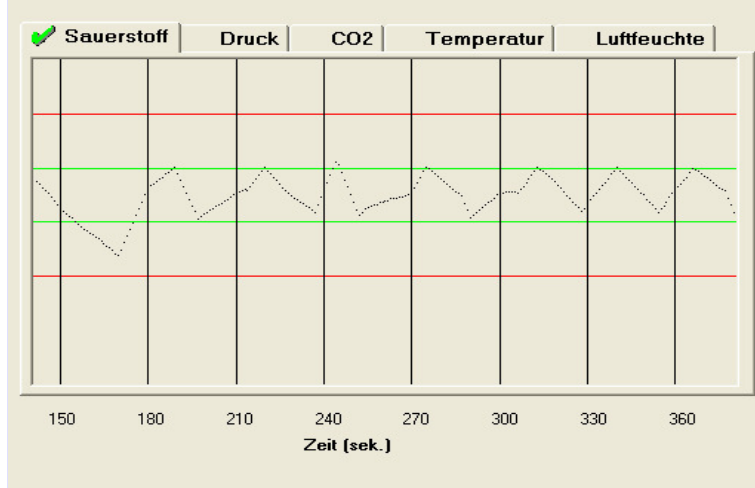
12



Der Überwachungs- und Steuerungsbildschirm



Die Verlaufsanzeige



Einführung in CAMS

Anke-Dorothea Hüper

15

Der Überwachungs- und Steuerungsbildschirm



Die Steuerungsmenüs

Sauerstoff **Steuerung**

19.0 20.5

19.5 20.0

19.6%

Steuerung

☒ Auto

☐ Manuell

☐ ein

☐ aus

Gasstrom

☐ hoch ☐ mittel

☒ stand. ☐ null

Standards Gasströme

	Stickstoff	Sauerstoff
stand.	15	6
mittel	25	9
hoch	45	18

Menüs der einzelnen Subsysteme

Soll- und Normalbereich

Aktueller Messwert

Warnleuchte (Masteralarm)

Gasstromstärke

Allgemeine Gasstromwerte

Einführung in CAMS

Anke-Dorothea Hüper

16

Aufgaben



Ihre Aufgaben sind...

- Halten der Parameter im Sollbereich
 - Überwachung der Parameter und wenn nötig steuernd eingreifen
- Fehlermanagement
 - Sofortiges Einschreiten bei Fehlfunktion von CAMS durch:
 1. Ursache der Fehlfunktion herausfinden
 2. Manuelle Regelung des Parameters
 3. Absenden eines Reparaturauftrages an die Besatzung der Raumstation
 4. Nach erfolgtem Fehlermanagement (nach einer Minute) wieder Überwachung von CAMS

Aufgaben



Ihre Aufgaben sind...

- Ressourcen sparen
 - Verschwenderische Steuerungsstrategien haben eine rapide Abnahme der O2- und N2-Vorräte zur Folge
- Bestätigen der Verbindung mit der Raumstation
 - Sobald Verbindungszeichen auftaucht, muss es mit der linken Maustaste bestätigt werden



Aufgaben



Ihre Aufgaben sind...

- Erfassen des Kohlenstoffdioxidwertes in regelmäßigen Abständen
 1. Steuerungsmenü von CO2 öffnen und Wert ablesen
 2. Wert in das dafür vorgesehene Feld jeweils zur vollen Minute eintragen
 3. „Wert speichern“ anklicken

The screenshot shows a control interface with two main sections. The top section is titled 'Zeit' and contains a digital display showing '0:14'. The bottom section is titled 'Nächster CO2 Check ... zum Zeitpunkt (Min.)' and contains a digital display showing '1:00'. Below this display is a button labeled 'Wert speichern' and an empty input field.

Störungen und Fehldiagnosen



Einführung Training

Bitte starten Sie „Einführung“ und überwachen Sie das System.

Bearbeiten Sie bitte auch die Aufgaben „Ablesen des CO2-Wertes“ und „Bestätigen der Verbindung mit der Raumstation“.



Fehlerdiagnose & -management

- Bei Anzeichen von Fehlfunktionen des Systems (rote Warnleuchte) sind eine Reihe von Diagnoseschritten notwendig
- Reparatur erfolgt durch die Crew an Bord, sie benötigt allerdings einen Reparaturauftrag
- Alle möglichen Reparaturaufträge sind im Reparaturmenü aufgeführt
- voreilige Handlungen müssen aufgrund des großen Ressourcenverlustes vermieden werden
- Absendung des Reparaturauftrages nur bei eindeutiger Diagnose
- Sobald keine Fehlfunktion des Systems mehr vorliegt, wird Warnleuchte grün

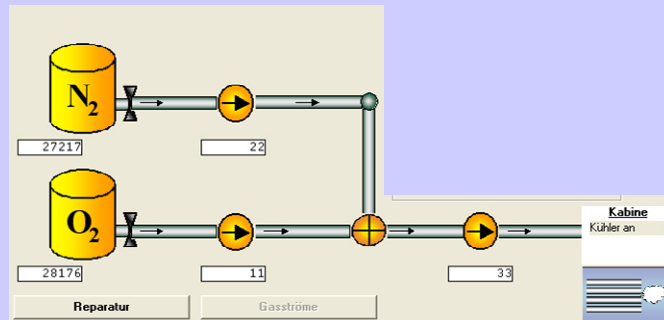


Fehlfunktionen

- Alle möglichen Fehler lassen sich in generell vier Typen einteilen, betroffen ist nur das Sauerstoff (O₂) - oder das Stickstoffsystem (N₂)
 1. Leck im Ventil
 - Verlust von O₂ oder N₂, verringerte Zufuhr in die Kabine, Ressourcenverlust
 2. Blockierung eines Ventils
 - Keine vollständige Öffnung, verringerte Zufuhr in die Kabine
 3. Offen steckendes Ventil
 - Ventil lässt sich nicht mehr schließen, zu großer Gasstrom in die Kabine, Ressourcenverlust
 4. Defekte Sensoren
 - Abweichung des Parameters von den Grenzwerten wird nicht mehr erkannt, keine automatische Regelung mehr

Diagnose Leck im Ventil

- Gasmenge die in das Ventil strömt ist größer als Gasmenge, die das Ventil verlässt



- Beobachtung der Gasströme für Diagnose notwendig

Management Leck im Ventil

1. Stärke des Gasstrom muss im betreffenden Steuerungsmenü von „stand.“ (standard) auf „hoch“ gesetzt werden

trotz des Gasverlustes auf dem Weg von den Tanks in die Kabine gelangt ausreichend Sauerstoff/Stickstoff in die Kabine

2. Absenden des Reparaturauftrages an die Besatzung
3. nach erfolgter Reparatur Zurücksetzen der Stärke des Gasstroms

Störungen und Fehldiagnosen



Leck im O2-Ventil Training

Bitte starten Sie „Leck O2“ und beheben Sie den auftretenden Fehler.

Einführung in CAMS

Anke-Dorothea Hüper

25

Störungen und Fehldiagnosen



Leck im N2-Ventil Training

Bitte starten Sie „Leck N2“ und beheben Sie den auftretenden Fehler.

Einführung in CAMS

Anke-Dorothea Hüper

26



Diagnose Blockierung eines Ventils

- Blockierung eines Ventils führt zu einer Reduktion des Gasstroms in der betreffenden Leitung
- Je nach Stärke des Gasstroms gibt es unterschiedliche Minimalwerte die mindestens auf den Gasstromanzeigen erscheinen müssen, wenn CAMS fehlerfrei arbeitet
- können durch Anklicken der Schaltfläche „Standards Gasströme“ angezeigt werden

Standards Gasströme		
	Stickstoff	Sauerstoff
stand.	15	6
mittel	25	9
hoch	45	18



Management von Blockierungen

Identisch mit dem Management von Lecks

Störungen und Fehldiagnosen



Blockierung O2-Ventil Training

Bitte starten Sie „Blockierung O2“ und beheben Sie den auftretenden Fehler.

Einführung in CAMS

Anke-Dorothea Hüper

29

Störungen und Fehldiagnosen



Blockierung N2-Ventil Training

Bitte starten Sie „Blockierung N2“ und beheben Sie den auftretenden Fehler.

Einführung in CAMS

Anke-Dorothea Hüper

30

Störungen und Fehldiagnosen


Training Blockierung oder Leck

Bitte starten Sie „Übung 1“ und beheben Sie den auftretenden Fehler.

Einführung in CAMS
Anke-Dorothea Hüper
31

Störungen und Fehldiagnosen


Diagnose offen steckendes Ventil O2

- Ventil lässt sich nicht mehr schließen
- Führt zu einem permanenten Zufluss von O2 in die Kabine und der daraus resultierenden steigenden Konzentration in der Kabine
- Gasstromanzeige zeigt einen permanenten Durchfluss von Sauerstoff
- Anhaltender Sauerstoffstrom erhöht auch den Kabinendruck bis zum oberen Normalwert
- Ablassventil wird geöffnet, um überschüssigen Druck abzubauen

Einführung in CAMS
Anke-Dorothea Hüper
32

**Management offen steckendem Ventil O2**

1. Ausschluss eines Defektes der Sauerstoffsensoren
 - Sauerstoffstrom im Steuerungsmenü muss unterbrochen werden
 - Nach Rückkehr der Sauerstoffkonzentration in den Normalbereich Regelung der O2-Konzentration wieder an die Automatik zurückgeben
 - Wenn O2-Konzentration sofort wieder steigt, dann besteht die Fehlfunktion im offen steckenden Ventil
 - Wenn O2-Konzentration sinkt, dann liegt ein defekter Sensor vor
2. Manuelle Regelung des Parameters bis Fehler behoben wurde
3. Abschicken des Reparaturauftrages
4. Einschalten der automatischen Kontrolle nach erfolgter Reparatur

**Offen steckendes Ventil O2 Training**

Bitte starten Sie „Offen steckendes Ventil O2“ und beheben Sie den auftretenden Fehler.

Störungen und Fehldiagnosen



Diagnose offen steckendes Ventil N2

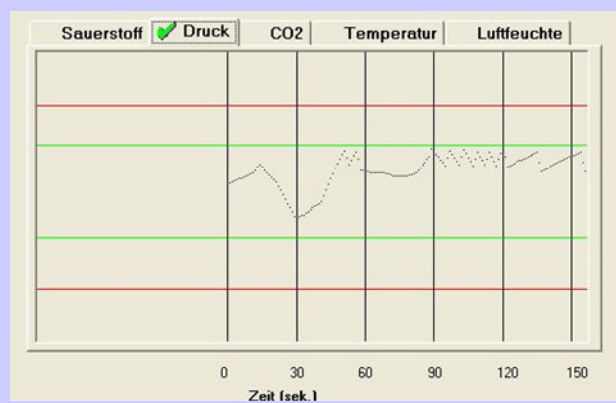
- Ventil lässt sich nicht mehr schließen
- Führt zu einem permanenten Zufluss von N2 in die Kabine und der daraus resultierenden steigenden Konzentration in der Kabine
- Gasstromanzeige zeigt einen permanenten Durchfluss von N2
- Durch Eingreifen des Entlüftungssystems bleibt Druck an der Obergrenze des Normalbereiches
- Sauerstoffkonzentration in der Kabine nimmt rapide ab und kann nur langsam wieder erhöht werden
- Verlaufsanzeige von Sauerstoff zeigt dadurch ein Muster aus steilen Einbrüchen und langsamen Anstiegen

Störungen und Fehldiagnosen



Diagnose offen steckendes Ventil N2

- Druckverlauf

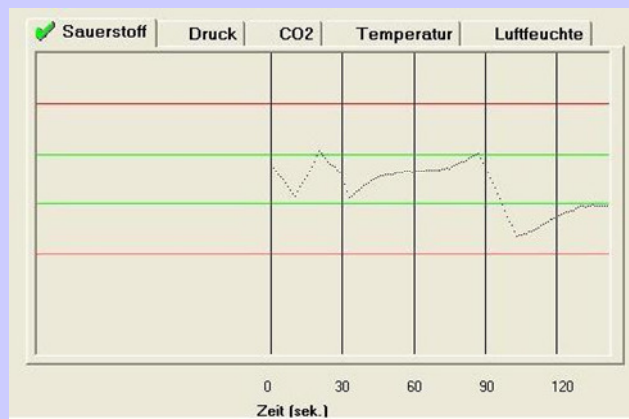


Störungen und Fehldiagnosen



Diagnose offen steckendes Ventil N2

- Sauerstoffverlauf



Störungen und Fehldiagnosen



Management offen steckendes Ventil N2

- Manuelle Steuerung des Drucks
- Abschicken des Reparaturauftrages
- Einschalten der automatischen Kontrolle nach erfolgter Reparatur

**Offen steckendes Ventil N2 Training**

Bitte starten Sie „Offen steckendes Ventil N2“ und beheben Sie den auftretenden Fehler.

**Diagnose defekter Sensor O2**

- Überschreitung der Grenzen des Normalbereiches wird von der Automatik nicht erkannt
- Nötiges Ein-/ bzw. Ausschalten der Gaszufuhr unterbleibt
- Dementsprechend kann ein defekter Sensor sowohl für zu hohe als auch für zu niedrige Werte des O2-Parameters verantwortlich sein
- Bei steigender O2-Konzentration
 - Ausschluss eines offen steckendem Ventil durch Testprozedur
- Bei sinkender O2-Konzentration
 - Wert auf der Gasstromanzeige muss Null sein



Management von defekten Sensoren O2

- Manuelle Steuerung des Parameters bis Fehler behoben wurde
- Abschicken des Reparaturauftrages
- Einschalten der automatischen Kontrolle nach erfolgter Reparatur



Defekter Sensor O2 Training

Bitte starten Sie „Defekter Sensor O2“ und beheben Sie den auftretenden Fehler.

Störungen und Fehldiagnosen



Diagnose defekter Sensor N2

- Überschreitung der Grenzen des Normalbereiches wird von der Automatik nicht erkannt
- Nötiges Ein-/ bzw. Ausschalten der Gaszufuhr unterbleibt
- Dementsprechend kann ein defekter Sensor sowohl für zu hohe als auch für zu niedrige Werte des N2-Parameters verantwortlich sein
- Bei steigender N2-Konzentration
 - Das Ablassventil bleibt geschlossen und der Druck steigt permanent
- Bei sinkender N2-Konzentration
 - Wert auf der Gasstromanzeige muss Null sein

Störungen und Fehldiagnosen



Management von defekten Sensoren N2

- Manuelle Steuerung des Parameters bis Fehler behoben wurde
- Abschicken des Reparaturauftrages
- Einschalten der automatischen Kontrolle nach erfolgter Reparatur

Störungen und Fehldiagnosen



Defekter Sensor N2 Training

Bitte starten Sie „Defekter Sensor N2“ und beheben Sie den auftretenden Fehler.

Einführung in CAMS

Anke-Dorothea Hüper

45

Störungen und Fehldiagnosen



Training

Bitte starten Sie „Übung 2“ und beheben Sie den auftretenden Fehler.

Einführung in CAMS

Anke-Dorothea Hüper

46

Störungen und Fehldiagnosen



Training

Bitte starten Sie „Übung 3“ und beheben Sie den auftretenden Fehler.

Einführung in CAMS

Anke-Dorothea Hüper

47



Fragen?

Einführung in CAMS

Anke-Dorothea Hüper

48



Versuchsdurchführung

Auto-CAMS

Erstellt für die Diplomarbeit von
Anke-Dorothea Hüper

Einführung CAMS 

Was ist CAMS?

**Cabin Air Management
System**

Lebenserhaltungssystem einer Raumstation

Einführung in CAMS Anke-Dorothea Hüper 2

Einführung CAMS



Was ist CAMS?

- CAMS wird von einer/m Ingenieur/in in der Bodenstation einer Raumfahrtagentur überwacht
- CAMS sorgt für ausreichenden Druck, ausreichende Temperatur und genug atembare Luft in der Raumstation
- CAMS besteht aus fünf Subsystemen, die die folgenden Parameter überwachen und automatisch regeln:
 - Sauerstoff (O₂)
 - Druck
 - Kohlenstoffdioxid
 - Temperatur
 - Luftfeuchtigkeit

Aufgaben



Ihre Aufgaben sind...

- Halten der Parameter im Sollbereich
 - Überwachung der Parameter und wenn nötig steuernd eingreifen
- Fehlermanagement
 - Sofortiges Einschreiten bei Fehlfunktion von CAMS durch:
 1. Manuelle Regelung des Parameters
 2. Ursache der Fehlfunktion herausfinden
 3. Absenden eines Reparaturauftrages an die Besatzung der Raumstation
 4. Nach erfolgtem Fehlermanagement (nach einer Minute) wieder Überwachung von CAMS

Aufgaben


Ihre Aufgaben sind...

- Ressourcen sparen
 - Verschwenderische Steuerungsstrategien haben eine rapide Abnahme der O2- und N2-Vorräte zur Folge
- Bestätigen der Verbindung mit der Raumstation
 - Sobald Verbindungszeichen auftaucht, muss es mit der linken Maustaste bestätigt werden



Einführung in CAMS
Anke-Dorothea Hüper
5

Aufgaben


Ihre Aufgaben sind...

- Erfassen des Kohlenstoffdioxidwertes in regelmäßigen Abständen
 1. Steuerungsmenü von CO2 öffnen und Wert ablesen
 2. Wert in das dafür vorgesehene Feld jeweils zur vollen Minute eintragen
 3. „Wert speichern“ anklicken



Einführung in CAMS
Anke-Dorothea Hüper
6



Fehlfunktionen

- Alle möglichen Fehler lassen sich in generell vier Typen einteilen, betroffen ist nur das Sauerstoff (O₂) - oder das Stickstoffsystem (N₂)
 1. Leck im Ventil
 - Verlust von O₂ oder N₂, verringerte Zufuhr in die Kabine, Ressourcenverlust
 2. Blockierung eines Ventils
 - Keine vollständige Öffnung, verringerte Zufuhr in die Kabine
 3. Offen steckendes Ventil
 - Ventil lässt sich nicht mehr schließen, zu großer Gasstrom in die Kabine, Ressourcenverlust
 4. Defekte Sensoren
 - Abweichung des Parameters von den Grenzwerten wird nicht mehr erkannt, keine automatische Regelung mehr



Training

Bitte starten Sie „Übung 4“ und beheben Sie den auftretenden Fehler.

Das Assistenzsystem von Auto-CAMS



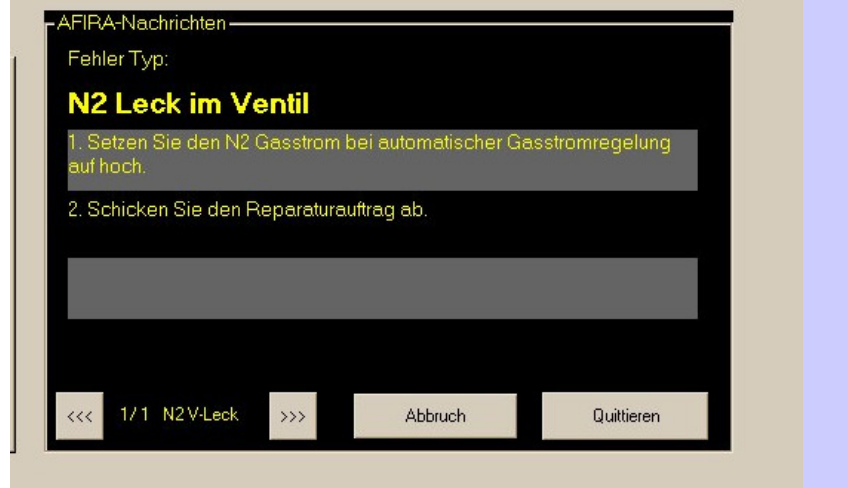
Was ist das Assistenzsystem?

- Wird AFIRA genannt
- Assistenzsystem unterstützt die Personen, die mit CAMS arbeiten
- Assistenzsystem schaltet sich automatisch ein, sobald Fehlfunktionen im System erkennbar sind (zeitgleich zum Aufleuchten des Masteralarms)
- Assistenzsystem unterstützt bei der Fehlerdiagnose und schlägt Fehlermanagement vor
- Da die Möglichkeit besteht, dass Fehler des Assistenzsystems vorkommen, sollten die Diagnose und das vorgeschlagene Fehlermanagement überprüft werden

Das Assistenzsystem



Was ist AFIRA?



Das Assistenzsystem


Was ist AFIRA?

AFIRA-Nachrichten

Fehler erfolgreich behoben

Setzen Sie das System auf die Standardeinstellungen zurück.

<<< 1/1 02 V-Leck >>>

Abbruch

Quitieren

Einführung in CAMS
Anke-Dorothea Hüper
11



Fragen?

Einführung in CAMS
Anke-Dorothea Hüper
12

Übersicht über die mit Blick auf die verschiedenen Fehlerdiagnosen relevanten und irrelevanten Informationsquellen

Informationsquelle	Sauerstoffsubsystem				Stickstoffsubsystem			
	Leck im Ventil	Blockierung des Ventils	Offenstehendes Ventil	Sensordefekt	Leck im Ventil	Blockierung des Ventils	Offenstehendes Ventil	Sensordefekt
Verlaufsanzeige O ₂	X	X	X	X			X	
Verlaufsanzeige N ₂					X	X	X	X
Verlaufsanzeige CO ₂								
Verlaufsanzeige Temperatur								
Verlaufsanzeige Luftfeuchtigkeit								
Gasströme	X	X		X (nur sinkend)	X	X		X (nur steigend)
Standards Gasströme		X				X		
Steuerungsmenü O ₂	X	X	X	X			X	
Steuerungsmenü N ₂					X	X	X	X
Steuerungsmenü CO ₂								
Steuerungsmenü Temperatur								
Steuerungsmenü Luftfeuchtigkeit								
O ₂ manuell auf „aus“			X	X (nur steigend)				
O ₂ manuell auf „Auto“			X	X (nur steigend)				

Anmerkung: mit „X“ gekennzeichnet sind jeweils relevante Informationsquellen, alle übrigen Informationsquellen sind mit Blick auf die Diagnoseprüfung irrelevant.

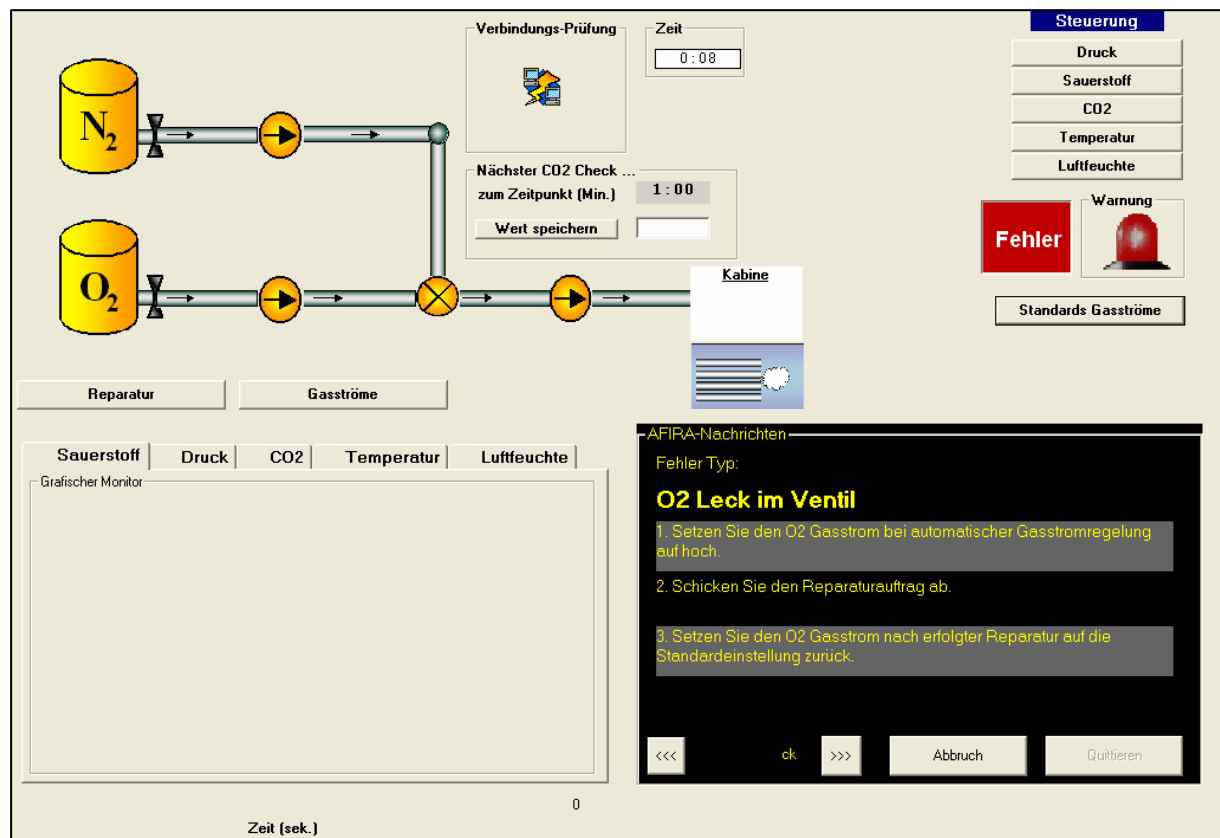
Übersicht über relevante und irrelevante Informationsquellen in vermeintlich fehlerfreien Phasen

Informationsquelle	Relevant für Fehlerdetektion	Irrelevant für Fehlerdetektion
Verlaufsanzeige O ₂	X	
Verlaufsanzeige N ₂	X	
Verlaufsanzeige CO ₂		X
Verlaufsanzeige Temperatur		X
Verlaufsanzeige Luftfeuchtigkeit		X
Gasströme	X	
Standards Gasströme	X	
Steuerungsmenü O ₂	X	
Steuerungsmenü N ₂	X	
Steuerungsmenü CO ₂		X
Steuerungsmenü Temperatur		X
Steuerungsmenü Luftfeuchtigkeit		X
O ₂ manuell auf „aus“	X	
O ₂ manuell auf „Auto“	X	

Experimentelle Gruppeneinteilung: Mittelwertsvergleiche basierend auf Daten vor der Manipulation (Experiment II)

	Informationsgruppe <i>M (SD)</i>	Erfahrungsgruppe <i>M (SD)</i>	<i>t</i>	<i>df</i>	<i>P</i>
Mittlere manuelle Fehlerdetektionszeit (über alle 10 Fehlfunktionen)	27.06 (9.90)	26.27 (8.54)	.203	21	.84
Mittlere manuelle Fehlerdetektionszeit (über N ₂ -Lecks)	21.50 (18.59)	12.00 (10.03)	1.50	21	.15
Mittlere manuelle Fehlerdetektionszeit (über O ₂ -Blockierungen)	26.50 (27.85)	25.50 (15.04)	.11	21	.92
Mittlere manuelle Fehleridentifikationszeit (über alle 10 Fehlfunktionen)	68.43 (26.16)	63.16 (47.90)	.33	21	.74
Mittlere manuelle Fehleridentifikationszeit (über N ₂ -Lecks)	81.00 (102.35)	45.27 (42.03)	1.08	21	.29
Mittlere manuelle Fehleridentifikationszeit (über O ₂ -Blockierungen)	52.50 (43.02)	53.45 (62.30)	-.04	21	.97
Informationssuchverhalten in den fehlerfreien Phasen (Informationsabrufe gesamt)	17.58 (6.07)	20.07 (5.10)	-1.06	21	.30
Selbsteinschätzung Fehlerdetektionsleistung	3.67 (0.89)	3.64 (0.81)	.09	21	.93
Selbsteinschätzung Fehlerdiagnoseleistung	3.08 (0.79)	3.36 (1.21)	-.66	21	.51
Selbsteinschätzung Fehlermanagementleistung	3.50 (0.52)	3.18 (0.75)	1.19	21	.25
Wahrgenommene Zuverlässigkeit N ₂ -Subsystem	2.75 (0.97)	2.64 (0.92)	.29	21	.78
Wahrgenommene Zuverlässigkeit O ₂ -Subsystem	2.58 (1.17)	2.27 (0.91)	.71	21	.49

Benutzeroberfläche AutoCAMS (Experiment II)



CAMS Checkliste – Übungen Fehler

Fehler	Fehlerrisiko ein	Gesetzlich/ Steuerung ändern	Reparaturantrag ab senden	Gesetzlich/ Steuerung zurück setzen	Fehlerrisiko aus
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐

2/10

CAMS Checklisten

VP	_____
Datum	_____
Uhrzeit	_____

Fehler	Fehlerrisiko ein	Gesetzlich / Steuerung ändern	Reparaturantrag abgeben	Gesetzlich / Steuerung zurücksetzen	Fehlerrisiko aus c
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐

L1/O1-7

Fehler	Fehlernummer ein	Gesamtsumme / Steuerung ändern	Reparaturtag abenden	Gesamtsumme / Steuerung zurücksetzen	Fehlernummer aus

3/10

Fehler	Fehlernummer ein	Gesetzlich Steuerung ändern	Reparaturtag abgeben	Gesetzlich Steuerung zurücksetzen	Fehlernummer aus
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐

6/10

[illegible]

[illegible][illegible]

CAMS Checkliste – Training 2



Fehler	Fehlerrisiko ein	Gesetzlich / Steuerung ändern	Reparaturantrag abgeben	Gesetzlich / Steuerung zurücksetzen	Fehlerrisiko aus
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐

9/10

CAMS Checkliste – Training 3



Fehler	Fehlerrückmeldung	Gesamt-/Steuerung ändern	Reparaturantrag abschicken	Gesamt-/Steuerung zurücksetzen	Fehlerrückmeldung ausgeben
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐
	☐	☐	☐	☐	☐

10/10

Handout zur Fehlerdiagnose und -behebung (Experiment II)



Fehler O2-System		
Fehlertyp	Fehlerdiagnose	Fehlermanagement
O2 Leck im Ventil	- O2-Abnahme in der Kabine - Gasmenge, die in das Ventil strömt > Gasmenge, die das Ventil verlässt - Gasmenge, die das Ventil verlässt, kann auch kleiner als „stand.“ in „Standards Gasströme“ sein	- Fehlermodus ein - Stärke des O2-Gasstroms von „stand.“ auf „hoch“ setzen - Reparaturauftrag abenden
O2 Blockierung des Ventils	- O2-Abnahme in der Kabine - Gemessener Gasstrom < „stand.“ in „Standards Gasströme“ - Gasmenge, die in das Ventil strömt > Gasmenge, die das Ventil verlässt	Nach erfolgreicher Reparatur: - Zurück setzen des Gasstroms auf „stand.“ - Fehlermodus aus
O2 offen stecken des Ventils	- O2-Anstieg in der Kabine (Ventil lässt sich nicht mehr schließen) - Druck steigt zum oberen Normalwert (Ablasseventil verhindert weiteren Anstieg) - Ausschluss eines O2 Sensordefektes durch Testprozedur: - Steuerung O2 auf Manuell „aus“ - nach Rückkehr der O2-Konzentration in Normalbereich wieder „Auto“ → wenn Konzentration wieder steigt: offen steckendes Ventil → wenn Konzentration sinkt: Sensordefekt	- Fehlermodus ein - Testprozedur - Manuelle Steuerung von O2 bis Fehler behoben wurde - Reparaturauftrag abenden Nach erfolgreicher Reparatur: - Steuerung wieder auf „Auto“ - Fehlermodus aus
O2 Sensordefekt	- O2-Anstieg ODER -Abnahme in der Kabine - O2-Anstieg: Testprozedur (Ausschluss eines offen steckenden O2 Ventils, siehe oben) - O2-Abnahme: → Gemessener Gasstrom = 0	- Fehlermodus ein - Ggf. Testprozedur - Manuelle Steuerung von O2 bis Fehler behoben wurde - Reparaturauftrag abenden Nach erfolgreicher Reparatur: - Steuerung wieder auf „Auto“ - Fehlermodus aus

Fehler N2-System		
Fehlertyp	Fehlerdiagnose	Fehlermanagement
N2 Leck im Ventil	- Druck-Abnahme in der Kabine - Gasmenge, die in das Ventil strömt > Gasmenge, die das Ventil verlässt - Gasmenge, die das Ventil verlässt, kann auch kleiner als „stand.“ in „Standards Gasströme“ sein	- Fehlermodus ein - Stärke des N2-Gasstroms von „stand.“ auf „hoch“ setzen - Reparaturauftrag abenden
N2 Blockierung des Ventils	- Druck-Abnahme in der Kabine - Gemessener Gasstrom < „stand.“ in „Standards Gasströme“ - Gasmenge, die in das Ventil strömt > Gasmenge, die das Ventil verlässt	Nach erfolgreicher Reparatur: - Zurück setzen des Gasstroms auf „stand.“ - Fehlermodus aus
N2 offen stecken des Ventils	- Druck-Anstieg BIS ZUR Obergrenze des Normalbereiches in der Kabine (Ventil lässt sich nicht mehr schließen; durch Eingreifen des Entlüftungssystems steigt Druck nicht über Normalbereich hinaus) - O2-Kurve zeigt ein Muster aus steilen Entwürfen und langsamen Anstiegen	- Fehlermodus ein - Manuelle Steuerung des Drucks bis Fehler behoben wurde - Reparaturauftrag abenden Nach erfolgreicher Reparatur: - Steuerung wieder auf „Auto“ - Fehlermodus aus
N2 Sensordefekt	- Druck-Anstieg ODER -Abnahme in der Kabine - Druck-Anstieg: → Druck steigt ÜBER Normalbereich hinaus (Ablasseventil bleibt geschlossen) - Druck-Abnahme: → Gemessener Gasstrom = 0	- Fehlermodus ein - Manuelle Steuerung des Drucks bis Fehler behoben wurde - Reparaturauftrag abenden Nach erfolgreicher Reparatur: - Steuerung wieder auf „Auto“ - Fehlermodus aus

Präsentation CAMS-Training (Experiment II)



Versuchsdurchführung

Arbeiten mit der Mikrowelt CAMS

Erstellt für die Diplomarbeit von
Monika Bepfandt

Agenda	
I.	Einführung CAMS
II.	Die fünf Subsysteme
	Sauerstoff
	Druck
	Kohlenstoffdioxid
	Temperatur
	Luftfeuchtigkeit
III.	Der Überwachungs- und Steuerungsbildschirm
	Bildschirm
	Verlaufsanzeige
	Steuerungsmenüs
III.	Aufgaben
IV.	Fehlfunktionen von CAMS

Einführung CAMS



Was ist CAMS?

Cabin Air Management System

Lebenserhaltungssystem einer Raumstation

Einführung in CAMS

Monika Elepfandt

3

Einführung CAMS



Was ist CAMS?

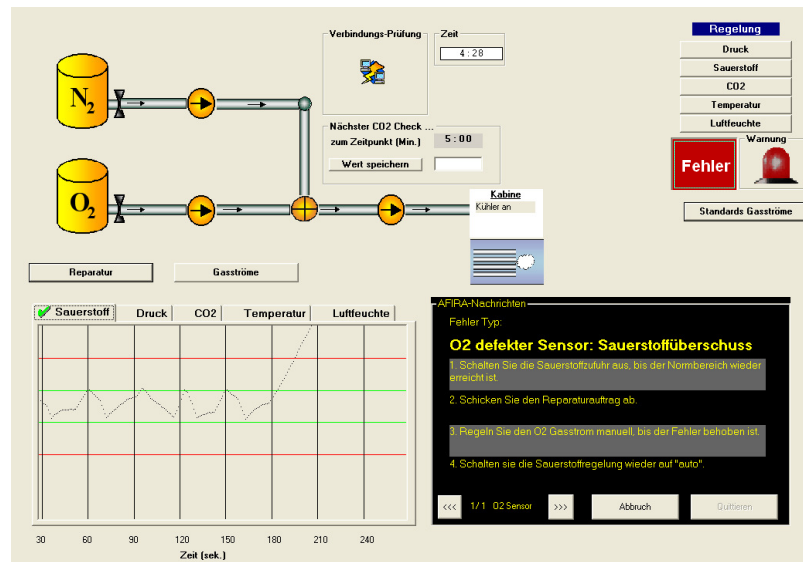
- CAMS wird von einer/m Ingenieur/in in der Bodenstation einer Raumfahrtagentur überwacht
- CAMS sorgt für ausreichenden Druck, ausreichende Temperatur und genug atembare Luft in der Raumstation
- CAMS besteht aus fünf Subsystemen, die die folgenden Parameter überwachen und automatisch steuern:
 - Sauerstoff
 - Druck
 - Kohlenstoffdioxid
 - Temperatur
 - Luftfeuchtigkeit

Einführung in CAMS

Monika Elepfandt

4

Einführung CAMS



Einführung in CAMS

Monika Elepfandt

5

Einführung CAMS



Was müsst ihr tun?

- Überwachung von CAMS
- Fehlermanagement
- Ressourcen sparen
- Bestätigen der Verbindung mit der Raumstation
- Erfassen des Kohlenstoffdioxidstandes

Einführung in CAMS

Monika Elepfandt

6

Subsysteme allgemein



Vorstellung der Subsysteme

- Jedes Subsystem ist für einen Parameter zuständig
- Für jeden Parameter ist ein Normal- und ein Sollbereich definiert
 - **Sollbereich** durch rote Begrenzung gekennzeichnet
 - **Normalbereich** durch grüne Begrenzung gekennzeichnet
- Bei Werten außerhalb des **Normalbereiches** liegt eine Fehlfunktion von CAMS vor
- Werte außerhalb des **Sollbereiches** gefährden das Leben der Besatzung der Raumstation

Einführung in CAMSMonika Elepfandt7

Subsysteme Sauerstoff



Sauerstoff

- Kontinuierliche Sauerstoffversorgung der Raumstation erforderlich
- Anteil von Sauerstoff in der Raumstation circa 20 %
- **Sollbereich:** 19 – 20,5 %
- **Normalbereich:** 19,5 – 20 %
- Wenn der Sauerstoffanteil der Kabine unter 19,5 % fällt, wird automatisch das Ventil des Sauerstofftanks geöffnet
- Wenn der Sauerstoffanteil 20 % erreicht, wird die Sauerstoffzufuhr unterbrochen

Einführung in CAMSMonika Elepfandt8

Subsysteme

Druck



Druck

- Konstanter Druck in der Raumstation erforderlich
- Normaler atmosphärischer Druck auf Meereshöhe circa 1000 mbar
- Wird an Bord der Raumstation durch Stickstoffzufuhr in die Kabine gesteuert
- **Sollbereich:** 970 – 1040 mbar
- **Normalbereich:** 990 – 1020 mbar
- Wenn der Kabinendruck unter 990 mbar fällt, wird automatisch das Ventil des Stickstofftanks geöffnet
- Wenn der Kabinendruck 1020 mbar erreicht, wird das Druckablassventil geöffnet

Einführung in CAMS

Monika Elepfandt

9

Subsysteme

Kohlenstoffdioxid



Kohlenstoffdioxid

- Kohlenstoffdioxid entsteht durch die Atmung der Besatzungsmitglieder
- Muss aus der Kabinenluft entfernt werden
- Aufnahme durch Filter
- **Sollbereich:** 0,1 – 1,5 %
- **Normalbereich:** 0,5 – 1 %
- Der Filter wird aktiv, sobald die Kohlenstoffdioxidkonzentration 1 % erreicht
- Der Filter schaltet sich aus, sobald die Kohlenstoffdioxidkonzentration unter 0,5 % sinkt

Einführung in CAMS

Monika Elepfandt

10

Subsysteme

Temperatur und Luftfeuchtigkeit



Temperatur

- **Sollbereich:** 18,5 – 23 °C
- **Normalbereich:** 20 – 22 °C
- Steuerung durch Luftkühler und Heizung

Luftfeuchtigkeit

- **Sollbereich:** 36 – 44 %
- **Normalbereich:** 38 – 42 %
- Steuerung durch Luftentfeuchter

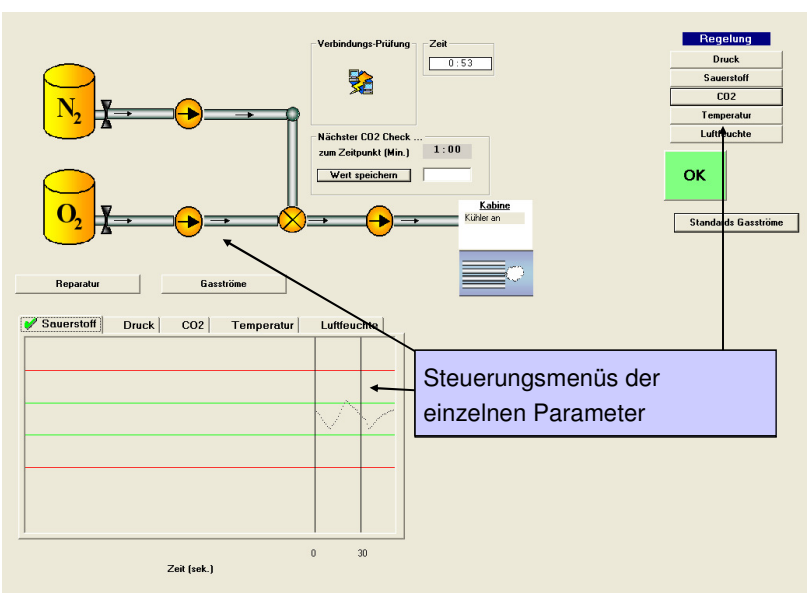
Einführung in CAMS

Monika Elepfandt

11

Der Überwachungs- und Steuerungsbildschirm





Verbindungs-Prüfung
 Zeit: 0:53
 Nächster CO2 Check ...
 zum Zeitpunkt (Min.) 1:00
 Weit speichern

Kabine
 Kühlt an

Regelung
 Druck
 Sauerstoff
 CO2
 Temperatur
 Luftfeuchte
 OK
 Standards Gasströme

Reparatur Gasströme
 Sauerstoff Druck CO2 Temperatur Luftfeuchte

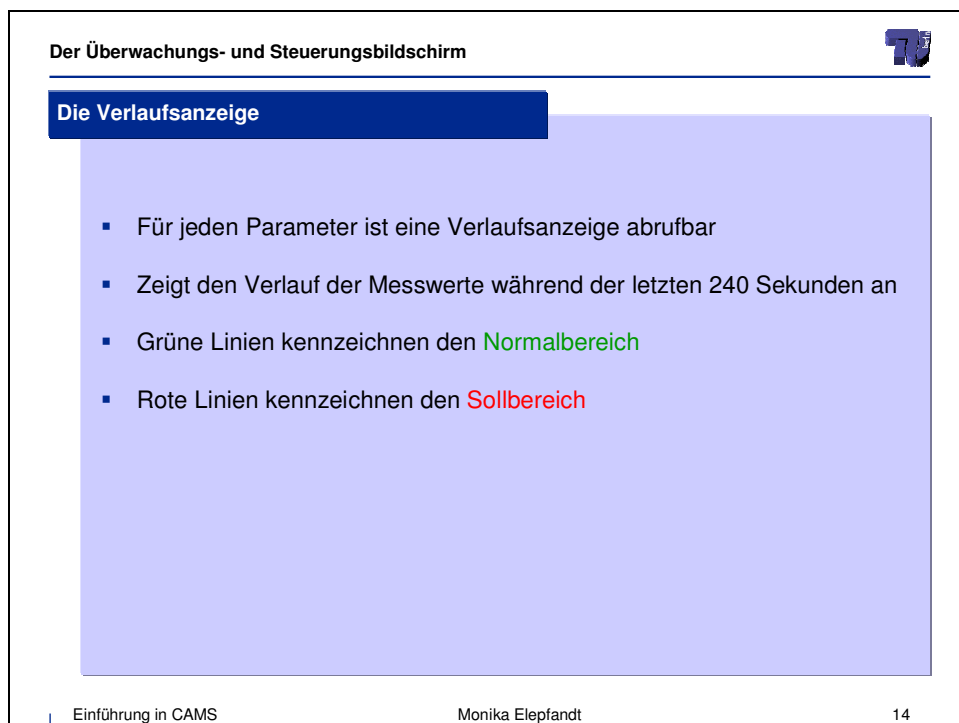
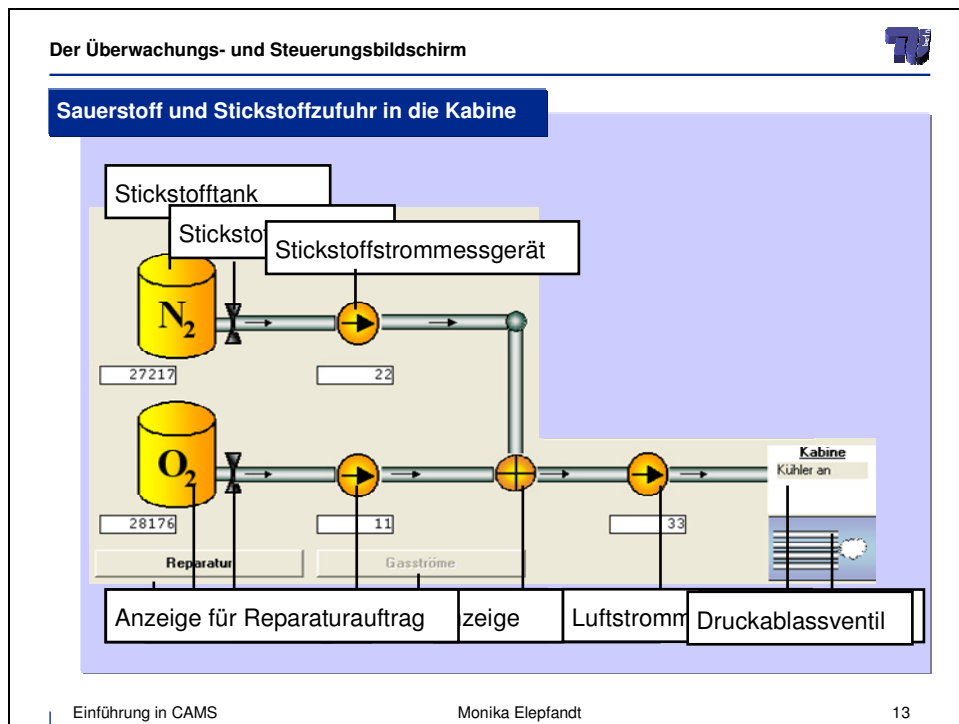
 Zeit (sek.) 0 30

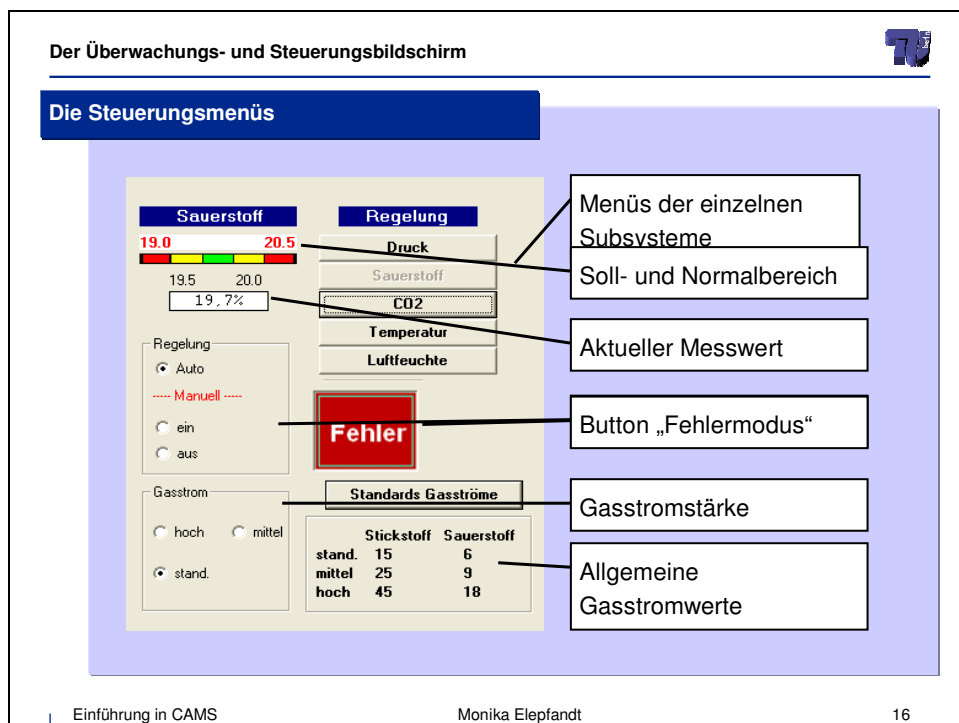
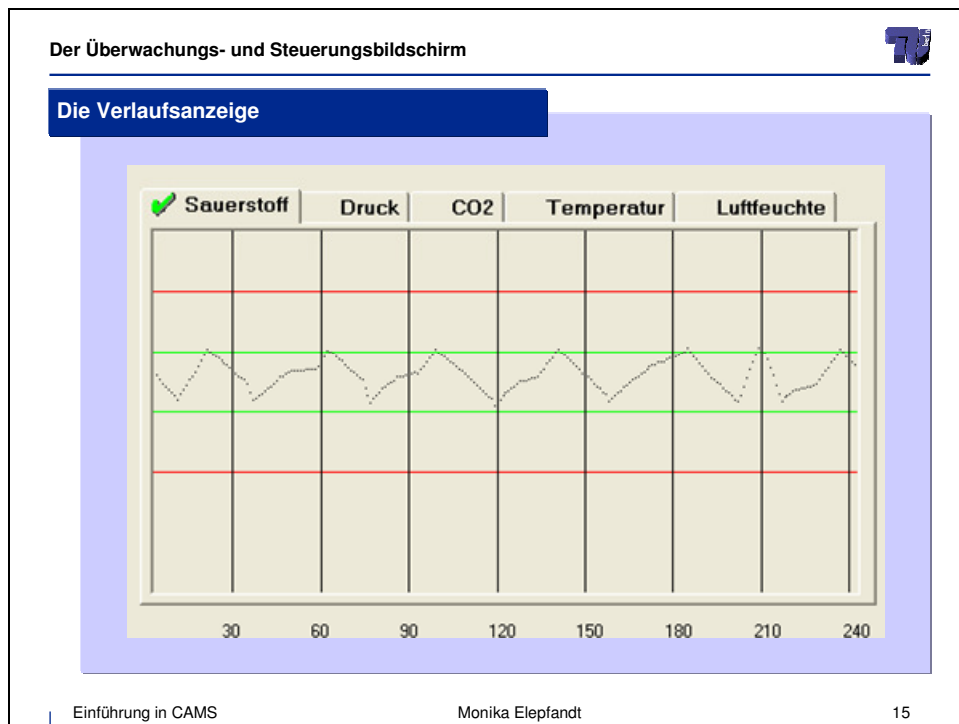
Steuerungsmenüs der einzelnen Parameter

Einführung in CAMS

Monika Elepfandt

12





Aufgaben


Eure Aufgaben sind...

- Halten der Parameter im **Sollbereich**
 - Überwachung der Parameter und wenn nötig steuernd eingreifen
- Fehlerdiagnose & -management
 - Sofortiges Einschreiten bei Fehlfunktion von CAMS durch:
 1. Fehlermodus einschalten
 2. Ursache der Fehlfunktion herausfinden
 3. Absenden eines Reparaturauftrages an die Besatzung der Raumstation
 4. Manuelle Steuerung des Parameters
 5. Fehlermodus ausschalten

Einführung in CAMS
Monika Elepfandt
17

Aufgaben


Eure Aufgaben sind...

- Ressourcen sparen
 - Verschwenderische Steuerungsstrategien haben eine rapide Abnahme der O2- und N2-Vorräte zur Folge
- Bestätigen der Verbindung mit der Raumstation
 - Sobald Verbindungszeichen auftaucht, muss es mit der linken Maustaste bestätigt werden



Einführung in CAMS
Monika Elepfandt
18

Aufgaben


Eure Aufgaben sind...

- Erfassen des CO₂-Wertes in regelmäßigen Abständen
 1. Steuerungsmenü von CO₂ öffnen und Wert ablesen
 2. Wert in das dafür vorgesehene Feld jeweils zur vollen Minute eintragen
 3. „Wert speichern“ anklicken



Einführung in CAMS
Monika Elepfandt
19

Aufgaben


Eure Aufgaben sind...

- Ressourcen sparen
 - Verschwenderische Steuerungsstrategien haben eine rapide Abnahme der O₂- und N₂-Vorräte zur Folge
- Bestätigen der Verbindung mit der Raumstation
 - Sobald Verbindungszeichen auftaucht, muss es mit der linken Maustaste bestätigt werden



Einführung in CAMS
Monika Elepfandt
20

Aufgaben



Einführung Training

Bitte startet die „Einführung“ und überwacht das System.

Bearbeitet bitte auch die Aufgaben „Ablesen des CO2-Wertes“ und „Bestätigen der Verbindung mit der Raumstation“.

Fehlfunktionen



Fehlerdiagnose & -management

- Bei eindeutiger Fehlfunktion des Systems Fehlermodus einschalten
- Fehlfunktion diagnostizieren
- Reparatur erfolgt durch die Besatzung an Bord, sie benötigt allerdings einen Reparaturauftrag (Dauer der Reparatur: 1 Minute)
- Alle möglichen Reparaturaufträge sind im Reparaturmenü aufgeführt
- Sobald keine Fehlfunktion des Systems mehr vorliegt, Fehlermodus ausschalten



Voreilige Handlungen müssen aufgrund des großen Ressourcenverlustes vermieden werden

- Nur bei eindeutiger Fehlfunktion Fehlermodus einschalten
- Nur bei eindeutiger Diagnose einen Reparaturauftrag absenden

Fehlfunktionen



Fehlfunktionen allgemein

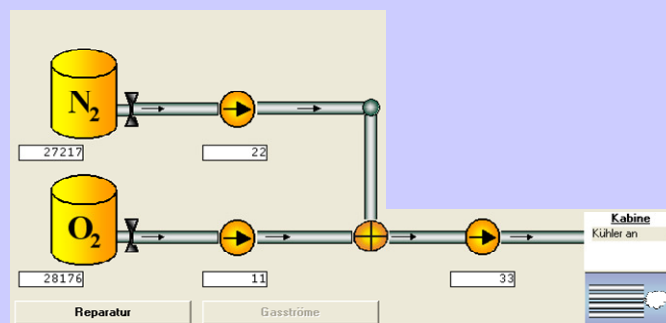
- Alle möglichen Fehler lassen sich in generell vier Typen einteilen, betroffen sind nur das Sauerstoff (O₂) - oder das Stickstoffsystem (N₂)
 1. Leck im Ventil
 - Verlust von O₂ oder N₂, verringerte Zufuhr in die Kabine, Ressourcenverlust
 2. Blockierung eines Ventils
 - Keine vollständige Öffnung, verringerte Zufuhr in die Kabine
 3. Offen steckendes Ventil
 - Ventil lässt sich nicht mehr schließen, zu großer Gasstrom in die Kabine, Ressourcenverlust
 4. Defekte Sensoren
 - Abweichung des Parameters von den Grenzwerten wird nicht mehr erkannt, keine automatische Steuerung mehr

Fehlfunktionen



Diagnose Leck im Ventil

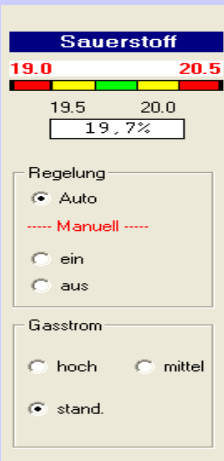
- Gasmenge, die in das Ventil strömt, ist größer als Gasmenge, die das Ventil verlässt



- Beobachtung der Gasströme für Diagnose notwendig

Fehlfunktionen


Management Leck im Ventil



1. Fehlermodus einschalten Fehler
2. Stärke des Gasstrom muss im betreffenden Steuerungsmenü von „stand.“ (standard) auf „hoch“ gesetzt werden
trotz des Gasverlustes auf dem Weg von den Tanks in die Kabine gelangt ausreichend Sauerstoff/Stickstoff in die Kabine
3. Absenden des Reparaturauftrags an die Besatzung
4. Nach erfolgreicher Reparatur Stärke des Gasstroms zurücksetzen und Fehlermodus ausschalten OK

Einführung in CAMS
Monika Elepfandt
25

Fehlfunktionen


Leck im O2-Ventil Training

Bitte startet „Leck O2“ und behebt den auftretenden Fehler.

Einführung in CAMS
Monika Elepfandt
26

Fehlfunktionen

Leck im N2-Ventil Training

Bitte startet „Leck N2“ und behebt den auftretenden Fehler.

Einführung in CAMS Monika Elepfandt 27

Fehlfunktionen

Diagnose Blockierung eines Ventils

- Blockierung eines Ventils führt zu einer Reduktion des Gasstroms in der betreffenden Leitung
- Je nach Stärke des Gasstroms gibt es unterschiedliche Minimalwerte, die mindestens auf den Gasstromanzeigen erscheinen müssen, wenn CAMS fehlerfrei arbeitet
- können durch Anklicken der Schaltfläche „Standards Gasströme“ angezeigt werden
- Gasstrom, der durch das Ventil strömt ist kleiner als bei „Standards Gasströme“

	Stickstoff	Sauerstoff
stand.	15	6
mittel	25	9
hoch	45	18

Einführung in CAMS Monika Elepfandt 28

Fehlfunktionen



Management von Blockierungen

Identisch mit dem Management von Lecks

Einführung in CAMS

Monika Elepfandt

29

Fehlfunktionen



Blockierung O2-Ventil Training

Bitte startet „Blockierung O2“ und behebt den auftretenden Fehler.

Einführung in CAMS

Monika Elepfandt

30

Fehlfunktionen



Blockierung N2-Ventil Training

Bitte startet „Blockierung N2“ und behebt den auftretenden Fehler.

Einführung in CAMS

Monika Elepfandt

31

Fehlfunktionen



Training Leck oder Blockierung

Bitte startet „Übung 1“ und behebt den auftretenden Fehler.

Einführung in CAMS

Monika Elepfandt

32

Fehlfunktionen



Diagnose offen steckendes Ventil O2

- Ventil lässt sich nicht mehr schließen
- Führt zu einem permanenten Zufluss von O2 in die Kabine und einer daraus resultierenden steigenden Konzentration in der Kabine
- Gasstromanzeige zeigt einen permanenten Durchfluss von O2
- Anhaltender O2-Strom erhöht den Kabinendruck bis zum oberen Normalwert
- Ablassventil wird geöffnet, um überschüssigen Druck abzubauen

Fehlfunktionen



Diagnose & Management offen steckendes Ventil O2

1. Fehlermodus einschalten
2. Ausschluss eines Sensordefektes O2
 - O2-Strom im Steuerungsmenü ausschalten
 - Nach Rückkehr der O2-Konzentration in den **Normalbereich** Steuerung des O2-Stroms wieder an die Automatik zurückgeben
 - Wenn O2-Konzentration sofort wieder steigt, liegt ein offen steckendes Ventil vor
 - Wenn O2-Konzentration sinkt, liegt ein Sensordefekt vor
3. Manuelle Steuerung von O2
4. Abschicken des Reparaturauftrages
5. Nach erfolgreicher Reparatur automatische Steuerung wieder einschalten und Fehlermodus ausschalten

Fehlfunktionen


Offen steckendes Ventil O2 Training

Bitte startet „Offen steckendes Ventil O2“ und behebt den auftretenden Fehler.

Einführung in CAMS
Monika Elepfandt
35

Fehlfunktionen


Diagnose offen steckendes Ventil N2

- Ventil lässt sich nicht mehr schließen
- Führt zu einem permanenten Zufluss von N2 in die Kabine und der daraus resultierenden steigenden Konzentration in der Kabine
- Gasstromanzeige zeigt einen permanenten Durchfluss von N2
- Durch Eingreifen des Entlüftungssystems bleibt Druck an der Obergrenze des Normalbereiches
- O2-Konzentration in der Kabine nimmt rapide ab und erhöht sich nur langsam wieder
- Verlaufsanzeige von O2 zeigt dadurch ein Muster aus steilen Einbrüchen und langsamen Anstiegen

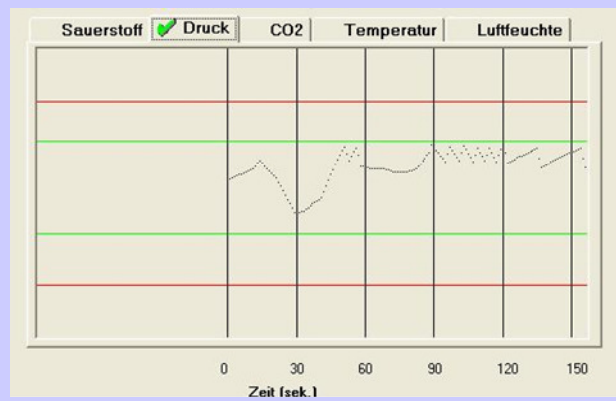
Einführung in CAMS
Monika Elepfandt
36

Fehlfunktionen



Diagnose offen steckendes Ventil N2

Druckverlauf



Einführung in CAMS

Monika Elepfandt

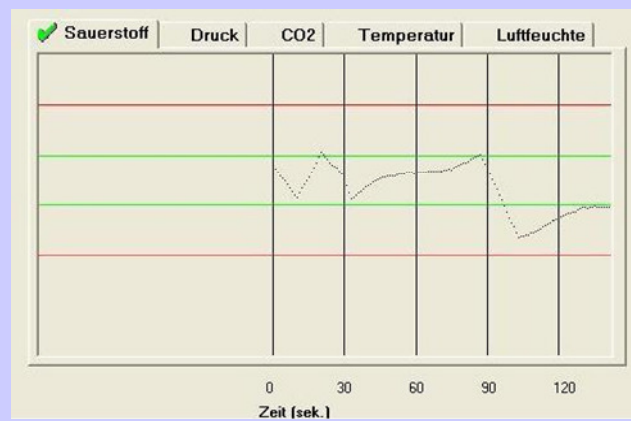
37

Fehlfunktionen



Diagnose offen steckendes Ventil N2

O2-Verlauf



Einführung in CAMS

Monika Elepfandt

38

Fehlfunktionen



Management offen steckendes Ventil N2

1. Fehlermodus einschalten
2. Manuelle Steuerung des Drucks
3. Abschicken des Reparaturauftrages
4. Nach erfolgreicher Reparatur automatische Steuerung wieder einschalten und Fehlermodus ausschalten

Einführung in CAMS

Monika Elepfandt

39

Fehlfunktionen



Offen steckendes Ventil N2 Training

Bitte startet „Offen steckendes Ventil N2“ und behebt den auftretenden Fehler.

Einführung in CAMS

Monika Elepfandt

40

Fehlfunktionen



Diagnose defekter Sensor O2

- Überschreitung der Grenzen des **Normalbereiches** wird von der Automatik nicht erkannt
- Nötiges Ein-/ bzw. Ausschalten der Gaszufuhr unterbleibt
- Dementsprechend kann ein defekter Sensor sowohl für zu hohe als auch für zu niedrige Werte des O2-Parameters verantwortlich sein
- Bei steigender O2-Konzentration
 - Ausschluss eines offen steckenden Ventils durch Testprozedur
- Bei sinkender O2-Konzentration
 - Wert auf der Gasstromanzeige muss Null sein

Fehlfunktionen



Management von defekten Sensoren O2

1. Fehlermodus einschalten
2. Manuelle Steuerung von O2
3. Abschicken des Reparaturauftrages
4. Nach erfolgreicher Reparatur automatische Steuerung wieder einschalten und Fehlermodus ausschalten

Fehlfunktionen


Defekter Sensor O2 Training

Bitte startet „Defekter Sensor O2“ und behebt den auftretenden Fehler.

Einführung in CAMS
Monika Elepfandt
43

Fehlfunktionen


Diagnose defekter Sensor N2

- Überschreitung der Grenzen des **Normalbereiches** wird von der Automatik nicht erkannt
- Nötiges Ein-/ bzw. Ausschalten der Gaszufuhr unterbleibt
- Dementsprechend kann ein defekter Sensor sowohl für zu hohe als auch für zu niedrige Werte des N2-Parameters verantwortlich sein
- Bei steigender N2-Konzentration
 - Das Ablasventil bleibt geschlossen und der Druck steigt über den Normalbereich hinaus
- Bei sinkender N2-Konzentration
 - Wert auf der Gasstromanzeige muss Null sein

Einführung in CAMS
Monika Elepfandt
44

Fehlfunktionen



Management von defekten Sensoren N2

1. Fehlermodus einschalten
2. Manuelle Steuerung des Drucks
3. Abschicken des Reparaturauftrages
4. Nach erfolgreicher Reparatur automatische Steuerung wieder einschalten und Fehlermodus ausschalten

Einführung in CAMS

Monika Elepfandt

45

Fehlfunktionen



Defekter Sensor N2 Training

Bitte startet „Defekter Sensor N2“ und behebt den auftretenden Fehler.

Einführung in CAMS

Monika Elepfandt

46

Fehlfunktionen



Training offen steckendes Ventil oder Sensordefekt

Bitte startet „Übung 2“ und behebt den auftretenden Fehler.

Einführung in CAMS

Monika Elepfandt

47

Fehlfunktionen



Training

Bitte startet „Übung 3“ und behebt den auftretenden Fehler.

Einführung in CAMS

Monika Elepfandt

48



Fragen?

Einführung in CAMS

Monika Elepfandt

49



Versuchsdurchführung

Auto-CAMS

Erstellt für die Diplomarbeit von
Monika Elepfandt

Agenda

I.	Einführung CAMS
II.	Aufgaben
III.	Fehlfunktionen
IV.	Das Assistenzsystem AFIRA

Einführung CAMS

Was ist CAMS?

Cabin Air Management System

Lebenserhaltungssystem einer Raumstation

Einführung in CAMS

Monika Elepfandt

3

Einführung CAMS

Was ist CAMS?

- CAMS wird von einer/m Ingenieur/in in der Bodenstation einer Raumfahrtagentur überwacht
- CAMS sorgt für ausreichenden Druck, ausreichende Temperatur und genug atembare Luft in der Raumstation
- CAMS besteht aus fünf Subsystemen, die die folgenden Parameter überwachen und automatisch steuern:
 - Sauerstoff
 - Druck
 - Kohlenstoffdioxid
 - Temperatur
 - Luftfeuchtigkeit

Einführung in CAMS

Monika Elepfandt

4

Aufgaben


Eure Aufgaben sind...

- Halten der Parameter im **Sollbereich**
 - Überwachung der Parameter und wenn nötig steuernd eingreifen
- Fehlerdiagnose & -management
 - Sofortiges Einschreiten bei Fehlfunktion von CAMS durch:
 1. Fehlermodus einschalten
 2. Ursache der Fehlfunktion herausfinden
 3. Absenden eines Reparaturauftrages an die Besatzung der Raumstation
 4. Manuelle Steuerung des Parameters
 5. Fehlermodus ausschalten

Einführung in CAMS
Monika Elepfandt
5

Aufgaben


Eure Aufgaben sind...

- Ressourcen sparen
 - Verschwenderische Steuerungsstrategien haben eine rapide Abnahme der O2- und N2-Vorräte zur Folge
- Bestätigen der Verbindung mit der Raumstation
 - Sobald Verbindungszeichen auftaucht, muss es mit der linken Maustaste bestätigt werden



Einführung in CAMS
Monika Elepfandt
6



Aufgaben

Eure Aufgaben sind...

- Erfassen des CO₂-Wertes in regelmäßigen Abständen
 1. Steuerungsmenü von CO₂ öffnen und Wert ablesen
 2. Wert in das dafür vorgesehene Feld jeweils zur vollen Minute eintragen
 3. „Wert speichern“ anklicken



Einführung in CAMS
Monika Elepfandt
7



Fehlfunktionen

Fehlfunktionen allgemein

- Alle möglichen Fehler lassen sich in generell vier Typen einteilen, betroffen sind nur das Sauerstoff (O₂) - oder das Stickstoffsystm (N₂)
 1. Leck im Ventil
Verlust von O₂ oder N₂, verringerte Zufuhr in die Kabine, Ressourcenverlust
 2. Blockierung eines Ventils
Keine vollständige Öffnung, verringerte Zufuhr in die Kabine
 3. Offen steckendes Ventil
Ventil lässt sich nicht mehr schließen, zu großer Gasstrom in die Kabine, Ressourcenverlust
 4. Defekte Sensoren
Abweichung des Parameters von den Grenzwerten wird nicht mehr erkannt, keine automatische Steuerung mehr

Einführung in CAMS
Monika Elepfandt
8

Fehlfunktionen


Training

Bitte startet die „Übung 4“ und behebt den auftretenden Fehler.

Einführung in CAMS
Monika Elepfandt
9

Das Assistenzsystem AFIRA


Was ist das Assistenzsystem AFIRA?

AFIRA

- unterstützt die Personen, die mit CAMS arbeiten
- schaltet sich automatisch ein, sobald Fehlfunktionen im System erkennbar sind
- unterstützt bei der Fehlerdiagnose und schlägt ein Fehlermanagement vor
- da die Möglichkeit besteht, dass Fehler des Assistenzsystems vorkommen, sollten die Parameter weiterhin überwacht werden.

Einführung in CAMS
Monika Elepfandt
10

Das Assistenzsystem AFIRA

The screenshot displays the AFIRA control interface. On the left, a schematic diagram shows two gas cylinders, N_2 and O_2 , connected to a network of pipes and valves. The O_2 line includes a valve labeled 'Kabine K hler an'. To the right of the diagram, a 'Verbindungs-Pr fung' (Connection Test) section shows a 'Zeit' (Time) of 0:49 and a 'N chster CO2 Check ... zum Zeitpunkt (Min.)' (Next CO2 Check ... at time (Min.)) of 1:00, with a 'Wert speichern' (Save value) button. Further right, a 'Regelung' (Control) panel lists parameters: Druck (Pressure), Sauerstoff (Oxygen), CO2, Temperatur (Temperature), and Luftfeuchte (Humidity). Below this is a 'Fehler' (Error) indicator with a red alarm bell icon and a 'Standards Gasstr me' (Standard Gas Flows) button. At the bottom left, a 'Reparatur' (Repair) and 'Gasstr me' (Gas Flows) section includes a 'Sauerstoff' (Oxygen) tab and a 'Grafischer Monitor' (Graphic Monitor) showing a 'Zeit (sek.)' (Time in seconds) scale from 0 to 30. On the bottom right, a 'AFIRA-Nachrichten' (AFIRA Messages) window displays the error: 'Fehler Typ: O2 defekter Sensor' (Error Type: O2 defective sensor). It provides four steps: 1. Schalten Sie die automatische O2 Gasstromkontrolle aus. (Switch off the automatic O2 gas flow control.) 2. Schicken Sie den Reparaturauftrag ab. (Submit the repair order.) 3. Regeln Sie den O2 Gasstrom manuell, bis der Fehler behoben ist. (Regulate the O2 gas flow manually until the error is resolved.) 4. Schalten Sie die Sauerstoffregelung wieder auf 'auto'. (Switch the oxygen control back to 'auto'). Navigation buttons include '<<< 1/1 O2 Sensor >>>', 'Abbruch' (Cancel), and 'Quittieren' (Acknowledge).

Einf hrung in CAMS

Monika Elepandt

11

Das Assistenzsystem AFIRA

This screenshot shows the AFIRA control interface after the error has been resolved. The 'Verbindungs-Pr fung' section now shows a 'Zeit' of 3:27 and a 'N chster CO2 Check ... zum Zeitpunkt (Min.)' of 4:00. The 'Regelung' panel remains the same. The 'Fehler' indicator is now green and shows 'OK'. The 'AFIRA-Nachrichten' window displays the success message: 'Fehler erfolgreich behoben' (Error successfully resolved) and 'Setzen Sie das System auf die Standardeinstellungen zur ck.' (Set the system back to the default settings.). The navigation buttons are the same as in the previous screenshot.

Einf hrung in CAMS

Monika Elepandt

12



Fragen?

Einführung in CAMS

Monika Elepfandt

13