



Thèse

2020

Open Access

This version of the publication is provided by the author(s) and made available in accordance with the copyright holder(s).

Konfigurationsales kausales Modellieren : Ein theoretisches Fundament und Verfahren für Kausalanalysen mit crisp-set Konfigurationen

Falk, Christoph

How to cite

FALK, Christoph. Konfigurationsales kausales Modellieren : Ein theoretisches Fundament und Verfahren für Kausalanalysen mit crisp-set Konfigurationen. Doctoral Thesis, 2020. doi: 10.13097/archive-ouverte/unige:133145

This publication URL: <https://archive-ouverte.unige.ch/unige:133145>

Publication DOI: [10.13097/archive-ouverte/unige:133145](https://doi.org/10.13097/archive-ouverte/unige:133145)

Konfigurationsales kausales Modellieren

**Ein theoretisches Fundament und Verfahren für Kausalanalysen
mit crisp-set Konfigurationen**

Inauguraldissertation
an der Fakultät für Geisteswissenschaften
der Universität Genf

vorgelegt von
Christoph Falk
von Alterswil

Leiter der Arbeit:
Prof. Dr. Michael Baumgartner, Universität Bergen

2019

Konfigurationsales kausales Modellieren

**Ein theoretisches Fundament und Verfahren für Kausalanalysen
mit crisp-set Konfigurationen**

Inauguraldissertation
an der Fakultät für Geisteswissenschaften
der Universität Genf

vorgelegt von

Christoph Falk

von Alterswil

Promotionskommission:

Prof. Dr. Marcel Weber, Universität Genf (Vorsitzender)

Prof. Dr. Michael Baumgartner, Universität Bergen (Leiter der Arbeit)

Prof. Dr. Claus Beisbart, Universität Bern

Prof. Dr. Jens Harbecke, Universität Witten/Herdecke

Danksagung

Ich bedanke mich bei all den Menschen, die in irgendeiner Form zur Entstehung der vorliegenden Arbeit beigetragen haben. Mein grösster Dank gilt dabei meinem Mentor, Prof. Dr. Michael Baumgartner. Er hat mich in das Gebiet der Kausalität und des kausalen Schliessens eingeführt und mich während meines Forschungsprojekts begleitet. Seine Unterstützung hat die Entstehung dieser Arbeit entscheidend mitgeprägt. Den unzähligen Diskussionen mit ihm und seinen stets sehr wertvollen Beiträgen und Rückmeldungen ist es zu verdanken, dass meine Ideen inhaltlich und strukturell jene Gestalt angenommen haben, die ich in der vorliegenden Arbeit präsentiere.

Weiter gilt es, ein grosses Dankeschön an meine Schwester, Stephanie Falk, zu richten. Trotz ihrer stets knapp bemessenen freien Zeit hat sie zahlreiche Stunden geopfert, um Kapitel für Kapitel gegenzulesen. Von ihren ausgeprägten redaktionellen Fähigkeiten konnte ich stark profitieren.

Schliesslich möchte ich mich auch bei meiner Arbeitgeberin, der Innosuisse, bedanken. Meine Arbeitskolleginnen und -kollegen haben mir das nötige Verständnis und die Flexibilität entgegengebracht, um diese Arbeit neben meiner beruflichen Tätigkeit verfassen zu können.

Inhaltsverzeichnis

Danksagung	i
1. Einleitung	1
1.1. Qualitative Comparative Analysis (QCA)	6
1.2. Coincidence Analysis (CNA)	9
1.3. Ziel und Aufbau	10
2. Der kausaltheoretische Rahmen	13
2.1. Grundlagen	13
2.1.1. Metaphysische Einbettung	14
2.1.2. Relata der Kausalrelation	14
2.1.3. Komplexität von Kausalstrukturen	17
2.1.4. Graphennotation	19
2.1.5. Faktormengen	21
2.2. Suffizienz, Notwendigkeit und Redundanzfreiheit	24
2.3. INUS-Theorie	28
2.4. RDN-Bikonditional	32
2.5. Strukturelle Redundanz	39
2.6. Minimale Theorie	44
2.7. Ambiguitäten	48
2.8. Vollständigkeits- und Erweiterungsklausel	49
2.8.1. Vollständigkeitsklausel	50
2.8.2. Erweiterungsklausel	51
2.9. Verursachungsbegriff	53
2.9.1. Kausale Relevanz	54
2.9.2. Ursachen einer Wirkung	55
2.10. Eigenschaften des Verursachungsbegriffs	57
2.10.1. Die kausalen Prinzipien	58
2.10.2. Kausale Wohlgeformtheit und minimaler Grad an kausaler Komplexität	62

2.11. Komplexe Minimale Theorien und komplexe Kausalstrukturen	66
2.11.1. Common-Cause-Strukturen	69
2.11.2. Kausalketten	70
2.11.3. Feedbackstrukturen	90
2.12. In Kürze	94
3. Empirische Grundlage	95
3.1. Von den Rohdaten zur Konfigurationstabelle	95
3.2. Skizzierung des analytischen Moments	102
3.3. Kausaler Fehlschluss	109
3.3.1. Entstehung eines kausalen Fehlschlusses	112
3.3.2. Störursachen	117
3.4. Homogenität	120
3.4.1. Homogenisierung der potenziellen Störursachen	120
3.4.2. Die Rolle von Homogenität für das analytische Moment	132
3.5. Maximale Diversität	135
3.6. In Kürze	139
4. Konfigurationsalgebra	140
4.1. Boolesche Algebren	140
4.1.1. Die Konfigurationsalgebra als mathematisches Modell einer zweiele- mentigen Booleschen Algebra	141
4.1.2. Instanzfunktion	144
4.2. Die Konfigurationsalgebra als algebraisches Fundament konfiguraler Me- thoden	156
4.2.1. Instanzfunktionen als Darstellungen des Instantiierungsverhaltens von Faktoren	157
4.2.2. Die Booleschen Gesetze als Schlussregeln des konfiguralen kausa- len Modellierens	160
4.2.3. Minimale Theorien als spezielle Darstellungsform einer Instanzfunktion	163
4.3. In Kürze	172
5. Das Aufdeckungsverfahren csCNA+	174
5.1. Das analytische Moment	175
5.2. Bestimmung aller RDN-Bikonditionale	181
5.2.1. Systematisierte Suche nach den RDN-Formen der inneren Instanzfunk- tionen	185
5.2.2. csQCA und das Quine-McCluskey-Verfahren	187
5.2.3. csCNA und das Baumgartner-Verfahren	204

5.2.4.	csQCA-Algorithmus und csCNA-Algorithmus im direkten Vergleich	216
5.3.	Bildung von Minimalen Theorien	218
5.3.1.	Bildung aller Konjunktionen von RDN-Bikonditionalen von Literalen paarweise verschiedener Faktoren	221
5.3.2.	Prüfung der Äquivalenz zur Konjunktion aller RDN-Bikonditionale	222
5.3.3.	Ausarbeitung aller strukturell minimalen Konjunktionen von RDN-Bikonditionalen	223
5.4.	Bildung von Kausalmodellen	225
5.4.1.	Ermittlung komplexer Minimaler Theorien	226
5.4.2.	Bestätigung indirekter Zusammenhänge	228
5.5.	csCNA+ in Kürze	234
6.	Fazit und Ausblick	236
6.1.	Fazit	236
6.2.	Ausblick	242
A.	Anhang	245
A.1.	Notation	246
A.2.	Beweise	247
A.3.	R-Replikationsskript	251
	Sachregister	258
	Literaturverzeichnis	261

1. Einleitung

Ein wesentlicher Teil der Bemühungen vieler Wissenschaften, die sich mit der Erklärung von Phänomenen oder der Prognose von Ereignissen beschäftigen, besteht in der Suche nach Kausalzusammenhängen. Im Idealfall erfolgt die Aufdeckung eines Kausalzusammenhangs zwischen einer Ursache *A* und einer Wirkung *B* dabei durch ein reproduzierbares Experiment, das die Kontrolle aller störenden Einflüsse ermöglicht und durch gezielte Manipulationen der untersuchten Faktoren sicherstellt, dass das Auftreten von *B* auf *A* zurückgeführt werden kann. In vielen Disziplinen ist man von der Möglichkeit einer solchen kontrollierten Herangehensweise jedoch weit entfernt, da der Untersuchungsgegenstand oft komplexe Zusammenhänge beinhaltet und man sich mit vielen Einflüssen konfrontiert sieht, die nicht alle berücksichtigt werden können, unbekannt sind oder sich nicht kontrollieren beziehungsweise manipulieren lassen.¹ Davon betroffen sind, um nur ein Beispiel zu nennen, Disziplinen wie die Human- und Sozialwissenschaften, deren Forschungsinteresse darauf gerichtet ist, die kausalen Verknüpfungen zwischen mehreren Faktoren wie zum Beispiel Alter, Ausbildung sowie Wertvorstellungen auf der einen Seite und der Parteipräferenz auf der anderen Seite zu finden (vgl. u. a. Bortz & Döring (2009), S. 11 f.). Zum einen lässt sich der Einfluss dieser Faktoren, im Fall des Alters besonders erkennbar, oftmals nicht innerhalb einer experimentellen Versuchsanordnung mit gezielten Manipulationen analysieren und die Untersuchung der Kausalzusammenhänge beruht zumeist auf der Beobachtung von natürlichen Variationen in den Ausprägungen der Faktoren (vgl. u. a. Bortz & Schuster (2010), S. 338). Zum anderen kann nicht ausgeschlossen werden, dass es viele weitere unkontrollierte Faktoren gibt, die eine Parteipräferenz unter gewissen Umständen auch beeinflussen können, wie etwa bestimmte Lebenslagen oder die politischen Erfolge aktueller Regierungsparteien.

Aus methodologischer Sicht werden zur Aufdeckung solcher komplexen Kausalzusammenhänge vor allem zwei Klassen von Methoden eingesetzt: die qualitativen und die quantitativen. Erstere zeichnen sich durch die Betrachtung weniger Fälle aus, die mit Einbezug von bereits bestehendem Kausalwissen und anderen wissenschaftlichen Theorien beispielsweise in Fallstudien miteinander verglichen werden. Diese Vergleiche zielen darauf ab, Gemeinsamkeiten

¹Einen Überblick über verschiedene Formen von Experimenten und ihre Merkmale liefern u. a. Shadish et al. (2001).

und Unterschiede herauszuarbeiten und zu beschreiben. Die daraus gewonnenen Resultate liefern in der Regel gründliche Kenntnisse über Einzelfälle, deren Geltungsbereich sich jedoch auf diese beschränkt. Im Gegensatz dazu basieren die quantitativen Methoden üblicherweise auf breit angelegten Studien mit Berücksichtigung möglichst vieler Fälle zur Folgerung von auf statistischen Verfahren basierenden Verallgemeinerungen (vgl. u. a. Ragin (2000), S. 22).² Zur Aufdeckung komplexer Kausalzusammenhänge werden dabei vor allem multiple Regressionsanalysen genutzt. Diese schätzen unter Benutzung eines multivariaten Regressionsmodells der Form $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_n X_n + \epsilon$ die Stärke des Einflusses eines jeden Ursachenfaktors (unabhängige Variablen $X_1, X_2, \dots, X_n$) auf die Wirkung (abhängige Variable Y), wobei ϵ ein Fehlerterm zur Berücksichtigung weiterer, im Modell unbeachteter Einflüsse darstellt (vgl. u. a. Kenny (1979); Antonakis et al. (2010)). Dabei geht man unter anderem davon aus, dass sich der Einfluss jeder X -Variable auf die Y -Variable unabhängig vom gleichzeitigen Einfluss aller weiteren im Modell spezifizierten Variablen schätzen lässt (vgl. bspw. Urban & Mayerl (2008), S. 81).

Aus kausaltheoretischer Sicht wird mit einer solchen Annahme die mögliche Komplexität der analysierbaren Kausalstrukturen limitiert: Oftmals möchte man eher behaupten, dass der Effekt eines Ursachenfaktors auf die Wirkung nur dann zu Tage tritt, wenn der Ursachenfaktor in bestimmten Kombinationen mit anderen Faktoren anwesend ist (vgl. bspw. Ragin (2008), S. 181). Untersucht man etwa Ursachen für Arbeitslosigkeit, darf davon ausgegangen werden, dass sich diese vor allem dann merklich erhöht, wenn neben einer Weltwirtschaftskrise die inländischen Firmen auch einen hohen Exportanteil aufweisen und es sich bei den exportierten Gütern um Luxusprodukte und nicht um lebenswichtige Arzneimittel handelt. Je nach Bedingungen unterscheiden sich die Effekte und Faktoren können in verschiedenen gegebenen Kombinationen gleiche oder aber auch unterschiedliche Auswirkungen haben. Das multiple Regressionsmodell, das Teil der grundsätzlich weit entwickelten und mathematisch strengen quantitativen Methoden ist, kann eine solche kausale Komplexität nur in eingeschränktem Masse erfassen, indem es den untersuchten Phänomenen formgebende Eigenschaften zugrunde legt, die aus der hier angesprochenen Perspektive nur vereinfachte Kausalstrukturen zulassen.

Um mit der multiplen Regressionstechnik auch Kausalzusammenhänge zu erfassen, bei denen nur die Kombination von mehreren Faktoren die Wirkung hervorrufen, werden Produktterme von X -Variablen in das Regressionsmodell integriert (vgl. z. B. Friedrich (1982); Brambor et al. (2006)). Bezeichnet beispielsweise die Kombination $X_1 X_2$ eine Ursache von Y , wird der Pro-

²Eine Gegenüberstellung der beiden Klassen von Forschungsmethoden findet man beispielsweise bei Mahoney & Goertz (2006).

duktterm X_1X_2 , der sogenannte Interaktionsterm, in das Regressionsmodell integriert und eine Gleichung der folgenden Form gebildet:

$$Y = \beta_0 + \beta_1X_1 + \beta_2X_2 + \beta_3X_1X_2 + \epsilon \quad (1.1)$$

Stellt sich nun heraus, dass der Koeffizient $\beta_3 > 0$ des Interaktionsterms statistisch signifikant ist, kann ein Interaktionseffekt abgeleitet werden, nach dem die Stärke des Einflusses von X_1 auf Y von der Stärke des Einflusses von X_2 abhängig ist und umgekehrt. Zur Veranschaulichung werde hier der Interaktionseffekt für die hypothetischen Werte $X_2 = 0$ sowie $X_2 = 1$ betrachtet, wobei β_3 und X_1 konstant den Wert 1 annehmen und $\beta_1 = \beta_2 = 0$ seien. Man stellt fest, dass X_1 einen Effekt auf Y hat, wenn X_2 den Wert 1, nicht aber wenn X_2 den Wert 0 annimmt, womit der Einfluss von X_1 abhängig von der Ausprägung von X_2 ist. Solche Interaktionsterme können um beliebig viele Variablen erweitert werden und man ist damit in der Lage, Kombinationen von mehreren Faktoren auf ihre gegenseitige Abhängigkeit für die Erzeugung ihres Effektes auf Y zu prüfen. Wie im Folgenden erklärt, stösst man dabei jedoch sehr schnell an Rechengrenzen.

Komplexe Kausalstrukturen können neben Kombinationen von gegebenen Ursachenfaktoren, die gemeinsam eine Wirkung hervorrufen, auch alternative Ursachen für eine Wirkung aufweisen. Beispielsweise kann die Arbeitslosigkeit eines Landes alternativ zur Weltwirtschaftskrise auch durch einen Bürgerkrieg stark steigen: In mindestens einem Fall führt eine Weltwirtschaftskrise zu steigender Arbeitslosigkeit und in mindestens einem anderen Fall ist ein Bürgerkrieg die Ursache.³ Damit hätte man zwei alternative Ursachen für denselben Wirkungstyp. Es lässt sich auch eine Kausalstruktur mit drei darin enthaltenen Ursachenfaktoren X_1 , X_2 und X_3 bilden, die die Wirkung Y hervorrufen, wenn X_1 zusammen mit X_2 oder X_1 zusammen mit X_3 gegeben sind. Beschreibe X_1 beispielsweise eine gute Infrastruktur, X_2 hohe Investitionen in Bildung und Forschung, X_3 eine liberale Wirtschaftspolitik und Y die hohe Wettbewerbsfähigkeit. Dem entsprechenden hypothetischen Kausalzusammenhang folgend wäre also das Vorhandensein einer guten Infrastruktur (X_1) zusammen mit hohen Investitionen in Bildung und Forschung (X_2) oder das Vorhandensein einer guten Infrastruktur (X_1) verbunden mit einer liberalen Wirtschaftspolitik (X_3) für hohe Wettbewerbsfähigkeit (Y) verantwortlich. Die Zusammenhänge zwischen X_1X_2 und Y sowie X_1X_3 und Y werden wiederum mit zwei Regressionsmodellen der Form wie (1.1) untersucht. Unter der Annahme, dass die verwendeten Daten den Anforderungen einer Regressionsanalyse genügen, werden die Koeffizienten der Interaktionsterme statistisch signifikant sein und so ermöglichen, auf den Zusammenhang zwischen der jeweiligen Faktorkombination und Y zu schliessen. Ein Problem entsteht nun dann, wenn überprüft wird, ob die beiden Kombinationen zu $X_1X_2X_3$ zusammengeführt werden sollen oder als alternative Ursachen X_1X_2 und

³Neben den beiden genannten Fällen können auch Fälle auftreten, bei denen sowohl die Weltwirtschaftskrise als auch ein Bürgerkrieg gegeben sind.

X_1X_3 von Y agieren. Folgt man etwa Brambor et al. (2006), wird zur Prüfung der Interaktionen zwischen X_1 , X_2 und X_3 bezüglich Y das folgende Modell gebildet:

$$Y = \beta_0 + \beta_1X_1 + \beta_2X_2 + \beta_3X_3 + \beta_4X_1X_2 + \beta_5X_1X_3 + \beta_6X_2X_3 + \beta_7X_1X_2X_3 + \epsilon \quad (1.2)$$

Durch die mehrmalige Verwendung gleicher Variablen in verschiedenen Termen wird dieses Modell in Anbetracht der tatsächlichen kausalen Abhängigkeiten sehr starke Korrelationen aufweisen und eine Änderung des Y -Wertes kann mit Variationen von mehreren Termen einhergehen, sodass sich die Ausprägungen einzelner Terme durch (lineare) Kombinationen anderer beschreiben lassen. Damit kann der Einfluss der einzelnen Variablen und ihrer verschiedenen Kombinationen nicht mehr präzise genug bestimmt und folglich keine Aussage über den Effekt der Interaktion zwischen X_1 , X_2 und X_3 getroffen werden.⁴ Dieses unter dem Namen Multikollinearität bekannte Problem hat damit starke Auswirkungen auf die Interpretation der einzelnen Koeffizienten β_i , die für das Erschliessen der kausalen Komplexität von zentraler Bedeutung sind.⁵ Eine Regressionsanalyse mit Interaktionstermen ist also, wie beispielsweise auch Clark et al. (2006) zeigen, durchaus in der Lage, den Effekt verschiedener Ausprägungen eines Faktors auf eine Wirkung in Abhängigkeit von den Ausprägungen anderer Faktoren zu bestimmen und damit Kombinationen von verschiedenen Faktoren als Ursachen zu folgern. Wird jedoch die kausale Komplexität um eine weitere Komponente erhöht und werden auch Kausalzusammenhänge betrachtet, die aus mehreren solchen alternativen Kombinationen bestehen, kann eine Analyse dieser Form alleine aufgrund der Multikollinearität vor grosse Probleme gestellt werden. Diese verschärfen sich weiter, wenn zusätzlich beispielsweise kausale Verknüpfungen zwischen den X -Variablen bestehen, was etwa dann der Fall wäre, wenn mindestens eine X -Variable eine weitere Wirkung darstellen würde, deren Hervorbringung teilweise von der Ausprägung der restlichen Variablen des Modells abhängig ist.

Das Problem der Multikollinearität im obigen Beispiel der Wettbewerbsfähigkeit kann grösstenteils darauf zurückgeführt werden, dass Regressionsverfahren im Allgemeinen nicht solche Arten von kausalen Hypothesen untersuchen, die primär die Komplexität von Kausalstrukturen betreffen. Der Schwerpunkt von Regressionsanalysen liegt nicht darin, herauszufinden, welche Faktoren zusammen oder alternativ zueinander bestimmte Wirkungen hervorrufen, sondern vielmehr bei der Untersuchung von direktionalen Hypothesen, die oftmals als Vergleichssätze der

⁴Das Weglassen von einzelnen Termen zur Verhinderung des Problems bietet dabei keine Lösung, da dadurch Annahmen getroffen werden, die oftmals nicht gerechtfertigt werden können (vgl. u. a. Brambor et al. (2006), S. 66 ff.).

⁵Es soll hier noch bemerkt werden, dass auch bei einer nachweisbaren Signifikanz eines Koeffizienten niedrigerer Ordnung, wie beispielsweise β_3 aus (1.2), nicht auf einen generellen Effekt von X_3 auf Y geschlossen werden kann, sondern lediglich als Effekt auf Y für $X_2 = X_3 = 0$ interpretiert werden darf (vgl. bspw. Braumoeller (2004)). Diese Einschränkung führt oft zu fehlerhaften Interpretationen der Modelle und verlangt zusätzliche Analysen, zumal Bedingungen wie $X_i = 0$, wobei X_i zum Beispiel das Alter oder Gewicht von Personen beschreibt, teilweise auch keinen Sinn ergeben.

Form „je mehr von X , desto mehr von Y “ verfasst sind oder vorzugsweise in präziseren Varianten wie beispielsweise: „Verdoppelt sich der Wert von X , steigt Y um das Dreifache“. Aus kausaltheoretischer Sicht zielt die Untersuchung solcher Hypothesen vor allem auf Erkenntnisse betreffend die Stärke des kausalen Einflusses von Faktoren ab und entsprechend basieren Regressionstechniken auf Grundlagen und Vorgehensweisen, die sich vor allem für die Betrachtung dieser Eigenschaft von kausalen Abhängigkeiten eignen.

Die oben am Beispiel der Wettbewerbsfähigkeit kurz skizzierte kausale Komplexität, die in der vorliegenden Arbeit im Zentrum steht, kann durch konfigurationale Methoden erfasst werden.⁶ Diese Methoden basieren auf der Booleschen Algebra und untersuchen implikationale Hypothesen der Form „wenn X , dann Y “.⁷ Dabei werden Phänomene basierend auf verschiedenen Konfigurationen von Faktoren untersucht, von denen angenommen wird, dass sie je nach Zusammensetzung die im Fokus der Untersuchung stehende Wirkung in ihrer Entstehung beeinflussen. Bezüglich des obigen Beispiels stellt man sich also die Frage, welche der Faktoren X_1 , X_2 und X_3 zusammen an- oder abwesend sein müssen, damit eine hohe Wettbewerbsfähigkeit begünstigt wird und ob es mehrere alternative Kombinationen davon gibt, die einen ähnlichen Einfluss haben. Dieses Interesse ist eng verknüpft mit der Suche nach notwendigen sowie hinreichenden Bedingungen der Wirkung (vgl. z. B. Schneider & Wagemann (2012), S. 78 f.). Solche Forschungsvorhaben basieren typischerweise auf Daten mittlerer Fallzahl, wobei jeder Datenpunkt konkreten Fällen einer Konfiguration der Ursachenfaktoren sowie der Wirkung entspricht. Für die Untersuchung der hohen Wettbewerbsfähigkeit Y im Zusammenhang mit X_1 , X_2 und X_3 könnte eine solche Datengrundlage etwa die Form von Tabelle 1.1 haben.

Jede Zeile c_i der Tabelle entspricht einer Konfiguration der Faktoren, wobei eine Eins beziehungsweise eine Null die Anwesenheit beziehungsweise Abwesenheit des entsprechenden Faktors in der Spaltenüberschrift kennzeichnet. Aus Zeile c_1 folgt also, dass bei der Erhebung der Daten Länder identifiziert wurden, die eine gute Infrastruktur haben, viel in Bildung und Forschung investieren sowie eine liberale Wirtschaftspolitik und hohe Wettbewerbsfähigkeit aufweisen. Aus Zeile c_2 ergibt sich eine ähnliche Folgerung mit dem Unterschied, dass in entsprechenden Fällen keine liberale Wirtschaftspolitik gegeben ist. Genau in diesen Unterschieden zwischen den einzelnen Konfigurationen und sich allenfalls daraus ergebenden Regelmäßigkeiten im Auftrittsverhalten der Faktoren liegt der Kern zur Etablierung von komplexen Kausalbeziehungen. So lässt sich aus Tabelle 1.1 etwa konkret entnehmen, dass immer, wenn X_1 und X_2 anwesend sind, auch Y anwesend ist. Doch welche Bedingungen müssen solche Regelmäßigkeiten genau erfüllen, um kausal interpretiert werden zu können? Wie lassen sich komplexe

⁶Alternativ werden diese auch Boolesche Methoden (vgl. z. B. Baumgartner (2009)) oder mengentheoretische Methoden (vgl. z. B. Schneider & Wagemann (2012)) genannt.

⁷Eine Gegenüberstellung von Regressionsverfahren und konfiguralen Methoden findet man u. a. bei Thiem et al. (2016).

	X_1	X_2	X_3	Y
c_1	1	1	1	1
c_2	1	1	0	1
c_3	1	0	1	1
c_4	1	0	0	0
c_5	0	1	1	0
c_6	0	1	0	0
c_7	0	0	1	0
c_8	0	0	0	0

Tabelle 1.1.: Hypothetische Datengrundlage für die Untersuchung des kausalen Zusammenhangs zwischen guter Infrastruktur (X_1), hohen Investitionen in Bildung und Forschung (X_2), liberaler Wirtschaftspolitik (X_3) und der als Wirkung im Fokus stehenden hohen Wettbewerbsfähigkeit (Y).

Kausalbeziehungen über derartigen Regelmässigkeiten einfangen? Und wie lässt sich eine entsprechende Suche systematisieren?

Kausaltheoretische Ansätze, die Kausalzusammenhänge mittels auf der Suffizienz und Notwendigkeit basierenden Regelmässigkeiten beschreiben und analysieren, werden als sogenannte Regularitätstheorien der Kausalität bezeichnet. Genau solche auf den Philosophen, Ökonomen und Historiker David Hume (1711–1776) zurückgehende Regularitätstheorien (vgl. Hume (1739)) bilden das Fundament des konfiguralen kausalen Modellierens. Das zentrale Ziel dieser Arbeit ist, dieses Fundament zu präsentieren und zu diskutieren sowie ein technisches Instrumentarium bereitzustellen, das in der Lage ist, Kausalstrukturen beliebiger Komplexität aus Daten wie in Tabelle 1.1 aufzudecken. In der Literatur lassen sich dazu zwei vielversprechende moderne Ansätze finden, die sich mit gleichen oder zumindest sehr ähnlichen Fragestellungen befassen: die Qualitative Comparative Analysis (QCA) und die Coincidence Analysis (CNA).

1.1. Qualitative Comparative Analysis (QCA)

Die Qualitative Comparative Analysis, kurz QCA, ist ein methodologischer Ansatz aus den Sozialwissenschaften, der sich mit dem Vergleich von qualitativ unterscheidbaren Konfigurationen und den daraus ableitbaren komplexen Kausalzusammenhängen beschäftigt. Ins Leben gerufen wurde QCA von Charles Ragin, der in seinem Buch *The Comparative Method* von 1987 ein auf der Booleschen Algebra beruhendes Verfahren vorstellt, das einen Mittelweg zwischen qualitativen und quantitativen Methoden darstellen sollte (Ragin (1987), S. 84). Im Zuge der weiteren

Auseinandersetzung mit QCA wurde auch dahingehend argumentiert, dass sich dieses Verfahren vor allem für die Untersuchung von Daten mittlerer Fallzahlen anbieten und so helfen kann, die methodologische Lücke zwischen Einzelfallstudien und statistischen Analysen bei grossen Fallzahlen zu schliessen (vgl. u. a. Schneider & Wagemann (2007), S. 26). Mit fortschreitenden Entwicklungen und Verfeinerungen festigte sich bei Vertretern dieser Analysemethode schliesslich die Meinung, dass QCA nicht nur als Lückenfüller zwischen etablierten Methoden betrachtet werden soll. Die Anwendung von QCA kann vielmehr mit dem theoretischen Argument gerechtfertigt werden, dass man sich bei vielen Untersuchungen für mengentheoretische Zusammenhänge und damit auch für die kausale Komplexität eines Phänomens interessiert, die alternative und jeweils aus mehreren Faktoren bestehende Ursachen einer Wirkung berücksichtigt (vgl. u. a. Schneider & Wagemann (2012), S. 12 und S. 76 f.). Für diese Interessenlage, so die Meinung, eignet sich QCA besonders gut.

QCA wurde direkt an Realdaten entwickelt und hat den Anspruch, unmittelbar in der Praxis eingesetzt werden zu können. Dieser Anspruch führte von Beginn weg zu Kritik, die vor allem die Aufbereitung der Daten betrifft. Die Hauptkritik gilt der Transformation der gesammelten Datenpunkte in analysierbare Konfigurationstabellen, wozu die beobachteten Faktoren dichotomisiert werden müssen.⁸ Im Bezug auf obiges Beispiel muss damit für jedes in Betracht gezogene Land beispielsweise bestimmt werden, ob es eine liberale Wirtschaftspolitik aufweist oder nicht. Viele Ereignisse und Zustände lassen sich jedoch nicht ohne starke Verzerrung in diese Zweiwertigkeit aufteilen. Um die Güte einer Analyse bewahren zu können, ist oftmals eine graduelle Abstufung wie „sehr liberale Wirtschaftspolitik“ oder „eher liberale Wirtschaftspolitik“ anzuwenden. Als Antwort auf diese Kritik wurde fuzzy-set QCA (fsQCA) entwickelt, welche die Fuzzylogik und die damit verbundenen Fuzzy-Mengen in QCA integriert und eine Verallgemeinerung der ursprünglichen Version, nun als crisp-set QCA (csQCA) bekannt, darstellt (vgl. z. B. Ragin (2008), Kapitel 2). Im Laufe der Zeit wurden auch weitere Aspekte wie beispielsweise die zeitliche Komponente integriert und so neben csQCA und fsQCA unter anderem auch

⁸In der QCA-Literatur spricht man anstelle von „Faktoren“ von „Bedingungen“ und „Outcome“. Die Verwendung des Begriffs „Bedingung“ lässt sich auf die notwendigen und hinreichenden Bedingungen zurückführen, deren Auffinden für das Verfahren zentral ist. Der Begriff „Outcome“, der eine Umschreibung für das zu erklärende bzw. verstehende Phänomen ist, wird mangels passenden deutschen Synonyms mit gleichem Bedeutungsumfang direkt aus dem Englischen übernommen. Zudem soll mit dem Begriffspaar „Bedingung-Outcome“ eine klare Abgrenzung von QCA zu rein statistischen Methoden gewährleistet werden, die, wie oben bei der Regressionsanalyse kurz angesprochen, dazu üblicherweise den Begriff „Variable“ mit der Unterscheidung „unabhängig-abhängig“ oder auch „exogen-endogen“ verwenden (vgl. bspw. Schneider & Wagemann (2007), 31 f.; Urban & Mayerl (2008), S. 28). In der vorliegenden Arbeit wird jeweils von Faktoren gesprochen, die sich in Ursachen- bzw. Wirkungsfaktoren unterteilen lassen. Damit ist hier dasselbe gemeint wie mit den in QCA benutzten Bedingungen bzw. Outcome.

temporal-QCA, kurz tQCA, entwickelt (vgl. u. a. Caren & Panofsky (2005)).⁹ Insgesamt hat dies dazu geführt, dass QCA heute eher als Oberbegriff für eine Familie von verschiedenen Methoden verwendet wird.

QCA wird nicht nur als technisches Instrument zur Analyse von Konfigurationen verstanden, sondern im weiteren Sinne auch als Forschungsansatz (vgl. Schneider & Wagemann (2012), S. 11). Das heisst, QCA besteht nicht nur aus einem Algorithmus, der aus Konfigurationen, wie beispielsweise jenen aus Tabelle 1.1, kausale Zusammenhänge ableitet. Vielmehr beinhaltet QCA unter anderem auch die Erhebung von Daten, die Definition der zu untersuchenden Faktoren und die Festlegung von Kriterien durch die entschieden werden kann, ob ein bestimmtes Datenelement einem Faktor zugeordnet werden kann oder nicht. Der algorithmische Teil von QCA, der aus den aufbereiteten Konfigurationen Kausalbeziehungen ableitet, wird als „analytisches Moment“ bezeichnet. Ohne QCA auf das analytische Moment reduzieren zu wollen, kann dieses trotzdem als Anker jeder Analyse mittels QCA verstanden werden, dessen Einsatz nicht nur gute von schlechter QCA-Praxis unterscheidet, sondern unverzichtbar für QCA ist.

Der Kern von QCA bildet die Suche nach Zusammensetzungen von notwendigen und hinreichenden Bedingungen für das Auftreten einer Wirkung *W* und die Minimierung dieser Bedingungen. Für die Ermittlung solcher Regularitäten greift QCA auf Verfahren der Booleschen Algebra zurück, vor allem auf die sogenannte Quine-McCluskey-Optimierung, die überwiegend in der Digitaltechnik zum Einsatz kommt (vgl. u. a. Quine (1952); Quine (1959); Nelson et al. (1995), S. 211 ff.). Auf diese Weise hat QCA den Vorteil, direkt von der Entwicklung etablierter Verfahren profitieren zu können. Ein wesentlicher Nachteil liegt im wenig ausgereiften kausaltheoretischen Konzept, das QCA zugrundeliegt. Indem QCA mehr oder weniger direkt aus der Anwendung Boolescher Methoden auf Realdaten entwickelt wurde, fehlt eine fundierte Auseinandersetzung mit dem Verursachungsbegriff. Zum einen ist man bemüht, hinreichende und/oder notwendige Bedingungen für eine Wirkung zu finden, und andererseits spricht man von kausaler Komplexität solcher Regularitäten, deren Bestandteile alle kausal interpretiert werden. In welchem Zusammenhang jedoch die notwendigen und/oder hinreichenden Bedingungen und die Kausalbeziehung genau stehen lässt sich nicht klar erkennen, da in den praxisnahen Abhandlungen den kausaltheoretischen Grundlagen zu wenig Beachtung geschenkt wird. Hinzu kommt, dass die kausaltheoretische Debatte aufgezeigt hat, dass sich jener Verursachungsbegriff, auf den sich QCA hauptsächlich beruft, nicht für das konfigurationale kausale Modellieren eignet. Bereits der Begründer der entsprechenden Kausaltheorie selbst, John Leslie Mackie, weist diese sogenannte INUS-Theorie gleich nach deren Entwicklung zurück, weil sie hinreichende Bedingungen als Ursachen auszeichnet, die tatsächlich keine Ursachen sind (vgl. Mackie (1980), Ka-

⁹Weitere bekannte Versionen wären Two-Step-QCA, mvQCA, ESA und TESA. Einen kurzen Überblick zu den verschiedenen Varianten liefern Schneider & Wagemann (2012), Kapitel 8 und 10.

pitel 3). Diese in der vorliegenden Arbeit nochmals dargelegte Debatte hat auch Auswirkungen auf das Aufdeckungsverfahren. Ersetzt man den von QCA üblicherweise eingesetzten Verursachungsbegriff durch einen Begriff, der die Schwierigkeiten von ersterem überwindet, müssen via der Quine-McCluskey-Optimierung teils unhaltbare Annahmen eingeführt werden, um zu Bedingungen mit Eigenschaften zu gelangen, die für eine kausale Interpretation gewünscht sind. Ohne Einführung solcher Annahmen können auch bei idealer Datenlage nur Kausalstrukturen analysiert werden, die aus untereinander unabhängigen direkten Ursachen einer Wirkung bestehen. Damit können zwar beispielsweise mehrere alternative, jeweils aus Kombinationen von anwesenden Faktoren bestehende Ursachen einer Wirkung aufgedeckt werden, nicht aber Strukturen, die direkte Ursachen mehrerer Wirkungen beinhalten oder indirekte Kausalzusammenhänge aufweisen (vgl. u. a. Baumgartner (2013a); Baumgartner & Epplé (2014)).

1.2. Coincidence Analysis (CNA)

Im Gegensatz zur praxisnahen Entwicklung von QCA entsprang die Coincidence Analysis von Michael Baumgartner, kurz CNA genannt, direkt aus der Entwicklung einer Kausaltheorie mit dem Ziel, ein auf diesen theoretischen Grundlagen basierendes Analyseverfahren zur Aufdeckung von komplexen Kausalstrukturen herauszubilden (vgl. u. a. Baumgartner (2006)). Die Entwicklung von CNA folgt dabei auf die Ausarbeitung einer modernen Regularitätstheorie der Kausalität, die sogenannte Minimale Theorien zur Analyse nutzt (vgl. u. a. Graßhoff & May (2001)). Kernbestandteile der Minimalen Theorien sind wie bei QCA Zusammensetzungen von notwendigen und hinreichenden Bedingungen für Wirkungen. Für die Bildung solcher Minimalen Theorien gilt es jedoch Anforderungen zu erfüllen, die über jene der teils bei QCA genutzten Kriterien hinausgehen und gleichzeitig wesentlich sind, um die Zusammensetzungen von Bedingungen kausal interpretieren zu können (vgl. u. a. Baumgartner & Epplé (2014)). Baumgartner überträgt diese Anforderungen relativ direkt auf einen Algorithmus, welcher deren Erfüllung garantiert, und ist so in der Lage, basierend auf der Beobachtung zusammen auftretender Faktoren Kausalstrukturen mit Wirkungen aufzudecken, die sowohl komplexe sowie alternative direkte Ursachen besitzen als auch indirekt verursacht werden können (vgl. bspw. Baumgartner (2013a)).

Dieses noch relativ junge Analyseverfahren deckt sich im wesentlichen mit jenem zentralen Aspekt der rein thematisch umfangreicheren QCA, der sich mit dem sogenannten analytischen Moment beschäftigt. Als Datengrundlage dienen dabei wie bei QCA die Konfigurationen von Faktoren – in diesem Fall von QCA terminologisch abweichend Koinzidenzen genannt, woraus

sich der Name des Verfahrens ableitet –, die miteinander verglichen werden, um so die darin enthaltenen kausalen Abhängigkeiten auszuarbeiten.

Der Vorteil von CNA liegt darin, dass sie aus einer modernen und gut fundierten Kausaltheorie entspringt, die unter anderem das erklärte Ziel verfolgt, Kausalstrukturen beliebiger Komplexität einzufangen. Somit hat CNA von Grund auf den Anspruch und das damit einhergehende Potenzial, ohne Hinzunahme teils problematischer Annahmen auch Kausalstrukturen aufzudecken, die mehrere Wirkungen aufweisen (vgl. u. a. Baumgartner (2009)). Konkret bedeutet dies, dass zuzüglich zu Kausalstrukturen mit nur einer Wirkung auch Common-Cause-Strukturen und kausale Verkettungen analysiert werden können. CNA hat also das Ziel, relativ zur Datengrundlage in Form von koinzidierenden Faktoren, Kausalstrukturen mit dem höchsten Grad an kausaler Komplexität einzufangen. Es wird sich jedoch zeigen, dass auch jener Verursachungsbegriff, der dazu in der Vergangenheit von CNA eingesetzt wurde, um zusätzliche Eigenschaften ergänzt werden muss, damit beliebige komplexe Kausalstrukturen korrekt eingefangen werden. Des Weiteren hat CNA Nachteile von praktischer Natur: CNA wurde aus dem für die Analyse benutzten Verursachungsbegriff heraus entwickelt und wandelt diesen mehr oder weniger direkt in einen Algorithmus zur Aufdeckung von Kausalzusammenhängen um, womit die dabei nicht im Fokus stehende Recheneffizienz relativ gering ist. Schliesslich weist CNA eine Lücke betreffend die Aufdeckung von indirekten Kausalbeziehungen auf, deren Folgerungen nicht als formalisierter Schritt im Algorithmus integriert sind.

1.3. Ziel und Aufbau

Die vorliegende Arbeit verfolgt das Ziel, ein auf QCA und CNA basierendes Verfahren zu entwickeln, das in der Lage ist, jede beliebige Kausalstruktur aus Konfigurationen wie jene aus Tabelle 1.1 aufzudecken. Dieses Vorhaben ist an die Beantwortung der folgenden drei wesentlichen Fragen gebunden:

1. Was genau wird gesucht?
2. Welche Informationen stehen dafür zur Verfügung?
3. Wie lässt sich basierend auf diesen Informationen das Gesuchte algorithmisch finden?

Die erste dieser drei Fragen betrifft den Verursachungsbegriff, der bei konfiguralen Methoden zum Einsatz kommt. Um ein Verfahren zur Aufdeckung von komplexen Kausalstrukturen zu entwickeln, muss in einem ersten Schritt definiert werden, wie genau die Beziehung zwischen

Ursache und Wirkung zu verstehen ist. Des Weiteren muss geklärt werden, wie die empirischen Daten aussehen, die für die Aufdeckung von Kausalzusammenhängen mittels konfiguratorischer Methoden zur Verfügung stehen. Dieser Klärungsbedarf wird durch die zweite der obigen drei Fragen ausgedrückt. Die Antworten auf die ersten beiden Fragen sind dabei eng miteinander verbunden, da im Hinblick auf das primäre Ziel nicht nur ein Verursachungsbegriff gefordert wird, der bekannte Kausalbeziehungen als solche auszeichnet, sondern zusätzlich auch ermöglicht, unbekannte Kausalzusammenhänge basierend auf empirischen Daten zu erschliessen. Wie genau auf Basis der erhobenen Daten kausale Abhängigkeiten algorithmisch aufgedeckt werden können, entspricht der Antwort auf die dritte Frage.

Um sicherzustellen, dass das Verfahren in der Lage ist, beliebige komplexe Kausalstrukturen aus den entsprechenden Konfigurationen korrekt abzuleiten, wird es unter idealisierten Bedingungen entwickelt. Von der Datenerhebung bis hin zur Folgerung von Kausalbeziehungen gibt es bei der Untersuchung eines Phänomens eine Vielzahl an Fehlerquellen, die zu Fehlschlüssen führen können. Unter anderem können Messfehler auftreten oder unvollständige Daten vorliegen. Der Fokus der vorliegenden Arbeit ist auf das analytische Moment des gesamten Forschungsprozesses gerichtet und es soll von Beginn weg sichergestellt werden, dass das Analyseverfahren keine Fehlschlüsse produziert, wenn die Konfigurationen Träger aller relevanten und nicht durch praktische Mängel verzerrten Informationen sind. Dies wird erreicht, indem das Verfahren in einem idealisierten Entdeckungskontext entwickelt wird, dessen Rahmenbedingungen so definiert sind, dass Fehlschlüsse ausschliesslich auf Fehlfunktionen des technischen Analyseinstruments oder ein unangemessenes kausaltheoretisches Fundament zurückgeführt werden können. Im Zusammenhang mit QCA bedeutet dies konkret, dass sich die Anwendung vorerst auf ideale und dichotome Daten und damit auf die crisp-set Variante (csQCA) beschränkt. Damit wird der Grundsatz verfolgt, dass ein Verfahren erst dann in einem praxisnäheren Entdeckungskontext betrachtet werden soll, wenn dessen Korrektheit unter idealen Bedingungen garantiert ist. Ist Letzteres der Fall, können die idealisierenden Bedingungen Schritt für Schritt abgebaut und das Verfahren kann entsprechend angepasst beziehungsweise durch approximative Methoden ersetzt werden.

Die erste der obigen drei Fragen wird in Kapitel 2 beantwortet. Unter Berücksichtigung der kausalen Komplexität von Kausalstrukturen werden in diesem Kapitel die kausaltheoretischen Grundlagen ausgearbeitet, die bestimmen, welche Eigenschaften notwendige und hinreichende Bedingungen aufweisen müssen, damit man diese basierend auf Konfigurationsvergleichen kausal interpretieren kann. Dazu wird das von CNA genutzte Analyseinstrument der Minimalen Theorien eingeführt und deren Charakteristika diskutiert, die für eine Kausalanalyse wesentlich sind. Es wird sich zeigen, dass der in der Vergangenheit genutzte, auf Minimalen Theorien basierende Verursachungsbegriff weiterentwickelt werden muss, damit über Konfigurationen er-

mittelte Regularitäten Kausalzusammenhänge korrekt einfangen. Kapitel 3 beschäftigt sich anschließend mit der Datengrundlage in Form von Konfigurationstabellen wie jene aus Tabelle 1.1, die für die Aufdeckung von Kausalzusammenhängen zur Verfügung stehen. Solche Tabellen können je nach Entstehung und darin enthaltenen Informationen einen starken Einfluss auf die Aussagekraft der daraus ableitbaren Schlussfolgerungen haben. Auf den beiden Kapiteln aufbauend widmet sich Kapitel 4 dem mathematischen Fundament, das beim konfiguralen kausalen Modellieren bei der crisp-set Variante zum Einsatz kommt. Dieses Fundament hilft, eine Brücke zwischen der Kausaltheorie, der empirischen Grundlage und dem konfiguralen kausalen Modellieren zu schlagen und erlaubt eine präzise Ausarbeitung der wesentlichen Prinzipien, die von QCA und CNA bei crisp-set Konfigurationstabellen zum Einsatz kommen. Letzteres ist Teil des Kapitels 5, in dem schliesslich ein Aufdeckungsverfahren präsentiert wird, das sich an der crisp-set Variante von CNA orientiert. Neben der Berücksichtigung der Weiterentwicklung der Kausaltheorie aus Kapitel 2, die es im Aufdeckungsverfahren zu implementieren gilt, werden Schritte aus csQCA zur Steigerung der Performance übernommen und Verfahrensschritte für den Schluss auf indirekte Kausalzusammenhänge konkret formalisiert und in den Algorithmus integriert. Das daraus resultierende erweiterte CNA-Verfahren für crisp-set Konfigurationstabellen, kurz als csCNA+ bezeichnet, ist in der Lage, innerhalb eines idealisierten Entdeckungskontextes Kausalstrukturen beliebiger Komplexität aus Konfigurationen korrekt abzuleiten. Das abschliessende Kapitel 6 fasst nochmals die zentralen Punkte zusammen und wagt einen Ausblick dahingehend, wie die in der vorliegenden Arbeit durchgängig verwendeten idealisierenden Bedingungen abgebaut werden können und wo die Grenzen von csCNA+ liegen.

2. Der kausaltheoretische Rahmen

Die Entwicklung eines Verfahrens zur Aufdeckung von Kausalzusammenhängen setzt einen angemessenen Verursachungsbegriff voraus. Bevor man sich die Frage stellt, wie Kausalzusammenhänge basierend auf Konfigurationen erschlossen werden können, muss man sich im Klaren darüber sein, wie Kausalität im Kontext des konfiguralen kausalen Modellierens überhaupt zu verstehen ist. Das vorliegende Kapitel gibt mit der darin präsentierten Regularitätstheorie der Kausalität eine entsprechende Antwort.

2.1. Grundlagen

Es gibt verschiedene Theorien der Kausalität, die unterschiedliche Verursachungsbegriffe liefern. Zu den bekanntesten gehören die Regularitätstheorien (vgl. u. a. Baumgartner (2008b)), kontrafaktische (vgl. z. B. Lewis (1986)) und probabilistische (vgl. z. B. Suppes (1970)) Ansätze sowie Interventions- (vgl. z. B. Woodward (2003)) und Transferenztheorien (vgl. z. B. Dowe (1995)). Während beispielsweise letztere den Zusammenhang zwischen einer Ursache *A* und deren Wirkung *B* primär mittels Energie- oder Impulsübertragungen beschreiben, basieren probabilistische Ansätze auf der Idee, dass *A* die Eintrittswahrscheinlichkeit von *B* ändert. Interventionstheorien wiederum gehen von möglichen Manipulationen von *A* aus, die zu Veränderungen bei *B* führen. Je nach Ansatz wird der Verursachungsbegriff auch mit unterschiedlichen Eigenschaften ausgestattet. Zum Beispiel ist Kausalität gemäss einigen Theorien eine deterministische Relation, der zufolge die gleichen Ursachen immer die gleichen Wirkungen haben. Gemäss anderen Theorien trifft dies nicht zu. Oder die Relation ist laut einigen Theorien transitiv, wonach aus „*A* verursacht *B*“ und „*B* verursacht *C*“ immer „*A* verursacht *C*“ folgt, was gemäss anderen Theorien nicht gilt.¹ Trotz dieser und weiterer, teils unvereinbarer Unterschiede hat sich bis heute keine der Theorien zuungunsten aller anderen durchgesetzt. Das hat zur gegenwärtig an Popularität gewinnenden Position des kausalen Pluralismus geführt, gemäss der unterschiedliche Theorien der Kausalität unterschiedliche Begriffe der Kausalität liefern, die alle in ihren jeweili-

¹Mehr zur deterministischen bzw. transitiven Relation siehe S. 58 bzw. S. 72 ff.

gen Anwendungsfeldern ihre Rechtfertigung haben. Sie widersprechen einander nicht, sondern ergänzen sich, beleuchten unterschiedliche Facetten von kausalen Strukturen und bilden unterschiedliche vor-theoretische kausale Intuitionen ab (vgl. z. B. Psillos (2009)). Diese Position wird auch in der vorliegenden Arbeit vertreten. Ohne die Angemessenheit anderer Theorien der Kausalität in anderen Anwendungsfeldern infrage stellen zu wollen, wird eine Regularitätstheorie der Kausalität präsentiert, die einen adäquaten Verursachungsbegriff für das konfigurationale kausale Modellieren bereitstellt.

2.1.1. Metaphysische Einbettung

So wie sich verschiedene Theorien der Kausalität unter anderem bezüglich der Eigenschaften unterscheiden, die dem entsprechenden Verursachungsbegriff zugeschrieben werden, sind sie teils auch in unterschiedlichen metaphysischen Auffassungen eingebettet. Regularitätstheorien stehen diesbezüglich in der Tradition ihres Begründers David Hume und sind aus metaphysischer Sicht in einem Hume'schen Aktualismus (bezüglich möglicher Welten) und Anti-Nezessitarismus verortet. Dieser Auffassung zufolge superveniert Kausalität auf der aktuellen Verteilung von Sachverhalten, wobei die Verteilung ihrerseits eine nackte Tatsache (*brute fact*) ist, die nicht weiter erklärt werden kann. Das heisst, dass es nicht die Kausalzusammenhänge sind, die diese Verteilung bestimmen, vielmehr ist es im Verständnis eines Anti-Nezessitarismus die Verteilung, die die Kausalzusammenhänge bestimmt. Bei letzteren handelt es sich um passende und zweckmässige Strukturierungen von Informationen, die in der aktuellen Verteilung von Sachverhalten enthalten sind (vgl. Hume (1748), Abschnitt 3). Obwohl diese Sichtweise – wie jeder andere metaphysische Ansatz auch – kontrovers ist, soll diese Kontroverse an dieser Stelle aufgrund des analytischen Schwerpunktes der vorliegenden Arbeit ausgeklammert bleiben. Einer pragmatischen Haltung folgend wird ein Hume'scher Anti-Nezessitarismus als ausreichend gerechtfertigt betrachtet, wenn sich die darauf basierende Kausaltheorie als angemessenes Fundament für das konfigurationale kausale Modellieren erweist.

2.1.2. Relata der Kausalrelation

Es wird zwischen mindestens zwei Arten von Kausalrelationen unterschieden: solche auf Typen-Ebene und solche auf Token-Ebene (vgl. u. a. Hausman (2005)). Bei letzterer werden singuläre Ereignisse kausal miteinander verbunden, die sich wesentlich durch ihre Lokalisierung in Raum und Zeit auszeichnen. So ist etwa der Autounfall von heute morgen um 10.00 Uhr in der Nähe des Bahnhofs ein singuläres Ereignis und stellt die Ursache (auf Token-Ebene) jenes singulä-

ren Ereignisses dar, das die um 10.30 Uhr im nahegelegenen Spital durchgeführte Notoperation beschreibt. Bei der Kausalrelation auf Typen-Ebene werden Faktoren kausal miteinander verknüpft. Solche Faktoren repräsentieren natürliche Eigenschaften, die Klassen von singulären Ereignissen oder Zuständen definieren können. Dabei haben die einem Faktor zugeschriebenen singulären Ereignisse oder Zustände die Eigenschaften gemeinsam, die durch den Faktor repräsentiert werden. Der Faktor „pluralistisches Parteiensystem“ beispielsweise repräsentiert die Klasse der politischen Systeme, bei denen mehrere Parteien an der Regierung eines Staates beteiligt sind. Solche Faktoren sind nicht in Raum und Zeit lokalisiert, können jedoch von singulären Ereignissen in Raum und Zeit instantiiert werden (vgl. bspw. Baumgartner (2006), S. 57 f.; Baumgartner & Graßhoff (2004), S. 39). Das Mehrparteiensystem der Schweiz oder Deutschlands aus dem Jahr 2018 sind zum Beispiel Instanzen des Faktors „pluralistisches Parteiensystem“ und die Anwesenheit des Faktors „Pluralistisches Parteiensystem“ eine Ursache (auf Typen-Ebene) dafür, dass der Faktor „Erschwerung einer Mehrheitsbildung“ gegeben ist. Bezüglich der genauen Kriterien, die eine durch einen Faktor repräsentierte Eigenschaft als natürlich auszeichnen, herrscht nach wie vor eine gewisse Unklarheit, weil sie nur schwierig zu präzisieren sind. Während standardmässig unter anderem jene Eigenschaften als natürlich betrachtet werden, die sich durch die eben erwähnte Instantiierbarkeit in Raum und Zeit eines entsprechenden Faktors auszeichnen, werden Eigenschaften wie beispielsweise Nelson Goodmans „grue“ oder solche, die „gerrymandered“ sind, ausgeschlossen (vgl. u. a. Lewis (1983); Fodor (1997)).² Wie in der Literatur zum kausalen Modellieren üblich, wird auch in der vorliegenden Arbeit auf eine präzise Definition des Begriffs der natürlichen Eigenschaft verzichtet. Stattdessen wird versucht, sich betreffend die Ontologie der Kausalität so wenig wie möglich festzulegen und einfach davon ausgegangen, dass alle Faktoren, von denen die Rede ist, jeweils Eigenschaften repräsentieren, die natürlich sind.

Den beiden Ebenen entsprechend gibt es auch mindestens zwei Arten von Kausaltheorien: (a) solche, welche die Kausalzusammenhänge primär auf der Token-Ebene beschreiben, und (b) solche, die Kausalzusammenhänge primär auf Typen-Ebene definieren. Das bedeutet, dass bei Kausaltheorien der Art (a) zuerst Tokenkausalität definiert wird und Typenkausalität dann basierend auf Tokenkausalität. Bei Kausaltheorien der Art (b) verhält es sich genau umgekehrt. Die in der vorliegenden Arbeit präsentierte Regularitätstheorie der Kausalität ist von der zweiten Art (b). Sie beschreibt die Kausalrelation primär als eine Verknüpfung zwischen Faktoren und analysiert die Kausalrelation auf Token-Ebene über das Instantiierungsverhalten von Faktoren, die (auf Typen-Ebene) kausal miteinander verbunden sind. Faktoren werden durch mit kursiven Grossbuchstaben *A*, *B*, *C* etc. bezeichneten Variablen dargestellt, denen Werte zugewiesen werden können. Diese Variablen können zweiwertig sein und die Werte 0 und 1 oder Werte eines Intervalls $[0, 1]$ annehmen. Weiter können sie auch mehrwertig sein und Werte aus einer end-

²Zur Eigenschaft „grue“ vgl. bspw. Goodman (1979), S. 74 f.

lichen Menge $\{0, 1, 2, \dots, n\}$ annehmen. Mit der Wertezuordnung von 0 und 1 lässt sich die An- und Abwesenheit eines Faktors ausdrücken. Ist ein Faktor wie das „pluralistische Parteiensystem“ (P) im Fall von Deutschland instantiiert, nimmt P den Wert 1 an. Ist P nicht instantiiert, wie beispielsweise im Fall von Kuba, wird der Wert 0 zugewiesen. Die Zuordnung eines Wertes im Intervall $[0, 1]$ kommt dann zum Einsatz, wenn man eine Graduierung zum Ausdruck bringen und etwa mit der Zuweisung von 0.7 beziehungsweise 0.2 zu P sagen möchte, das ein Land eher ein pluralistisches Parteiensystem aufweist beziehungsweise ein anderes Land eher nicht. Mittels Wertezuordnungen aus einer Menge $\{0, 1, 2, \dots, n\}$ können mit Faktoren schliesslich mehrwertige Eigenschaften repräsentiert werden, wie beispielsweise der Faktor „Ampel“ (A), die auf grün (etwa durch Wertezuordnung 0), gelb (Zuordnung 1) und rot (Zuordnung 2) gestellt sein kann. Die vorliegende Arbeit konzentriert sich auf den zweiwertigen Fall mit einer Wertezuweisung von 0 und 1. Dies entspricht dem fundamentalen crisp-set Fall, an dem sich die Entwicklung einer Kausaltheorie sowie eines Aufdeckungsverfahrens innerhalb eines idealisierten Entdeckungskontextes besonders gut darstellen lässt.

Kausalzusammenhänge bestehen nicht nur zwischen anwesenden Faktoren. Ein Kausalzusammenhang kann auch die Abwesenheit eines Faktors beinhalten. Bezeichnet zum Beispiel der Faktor K die Konsumfreudigkeit, ist die Abwesenheit dieses Faktors K eine Ursache für Umsatzeinbussen im Einzelhandel (U). Notiert wird dieses „nicht K “ in der vorliegenden Arbeit mit dem in der Booleschen Algebra oft genutzten Querstrich oberhalb des Faktors, womit „nicht K “ dargestellt wird als „ $\bar{K}$ “.³ Dieses Beispiel macht deutlich, dass die auf Faktoren basierenden Kausalzusammenhänge, die in der vorliegenden Arbeit betrachtet werden, präziser ausgedrückt Werte annehmende Faktoren oder kurz *Faktorwerte*, in Beziehung zueinander setzen. Nicht der Faktor K ist eine Ursache, sondern die Abwesenheit des Faktors K beziehungsweise der den Wert 0 annehmende Faktor K , kurz $\bar{K}$ ist eine Ursache für Umsatzeinbussen. Genauso verursacht die Abwesenheit von K strenggenommen nicht den Faktor U , sondern genauer gesagt das Vorhandensein des Faktors U beziehungsweise den den Wert 1 annehmenden Faktor U . Mit der gängigen Ausdrucksweise, dass ein Faktor X einen Faktor Y verursacht, ist somit beim konfiguralen kausalen Modellieren jeweils implizit gemeint, dass der gegebene Faktor X eine Ursache für den gegebenen Faktor Y ist. Um Letzteres explizit zu machen, soll an dieser Stelle der in der Booleschen Algebra übliche Begriff des *Literals* eingeführt werden. Konkret wird im Kontext der hier vorliegenden Typenkausalität genau dann von einem *Faktorliteral* oder kurz *Literal* (eines Faktors X) gesprochen, wenn es sich um den Wert 0 oder 1 annehmenden Faktor X handelt, wobei weiter der den Wert 0 annehmende Faktor X dem negativen Faktorliteral $\bar{X}$ und der den Wert 1 annehmende Faktor dem positiven Faktorliteral X entspricht. Um die damit einhergehende Verwechslungsgefahr zwischen dem Faktor „ X “ und dem identisch notierten

³Gebräuchliche alternative Schreibweisen zur Darstellung von $\bar{K}$ sind u. a. (Kleinbuchstaben) „ k “, „ $\neg K$ “ oder „ $\sim K$ “ (vgl. u. a. Schneider & Wagemann (2012), S. 54).

positiven Literal „X“ zu verhindern, ist mit dem kursiven Grossbuchstaben „X“, sofern nicht explizit als Faktor bezeichnet, jeweils das positive Faktorliteral gemeint.

2.1.3. Komplexität von Kausalstrukturen

In der Einleitung wurde bereits zum Ausdruck gebracht, dass das konfigurationale kausale Modellieren hauptsächlich durch dessen Umgang mit kausaler Komplexität motiviert wird. Dementsprechend muss auch ein Verursachungsbegriff diese Komplexität erfassen können, um als angemessenes Fundament für das konfigurationale kausale Modellieren infrage zu kommen. Bevor also ein adäquater Verursachungsbegriff formuliert werden kann, muss geklärt werden, welche Formen von kausaler Komplexität es überhaupt gibt.

Einfache Beispiele wie jenes der Wettbewerbsfähigkeit (vgl. S. 3 ff.) haben gezeigt, dass kausale Abhängigkeiten grundsätzlich nicht aus einer Verknüpfung von nur zwei Literalen bestehen und man für eine umfassende Modellierung von Kausalzusammenhängen Literale mehrerer Faktoren benötigt. So kann die Zunahme der Arbeitslosigkeit (A) neben der Automatisierung von Arbeitsprozessen (M) etwa auch durch den Wegzug grosser Firmen (F) in Kombination mit dem Fehlen eines Sozialplanes für die Entlassenen ($\bar{S}$) verursacht werden. Während bei der Kombination von F und $\bar{S}$ von einer *komplexen Ursache* von A gesprochen wird, entspricht M in diesem Beispiel einer *Alternativursache* von A . Komplexe Ursachen werden durch einfaches Aneinanderreihen der entsprechenden Literale dargestellt und alternative Ursachen durch das Symbol „+“ verknüpft, womit die zwei Alternativursachen für die Zunahme der Arbeitslosigkeit (A) notiert werden als „ $F\bar{S} + M$ “. Die Schreibweise der komplexen Ursache entspricht der in der Booleschen Algebra gängigen verkürzten Notation der konjunktiven Verknüpfung, die im Fall der Konjunktion $F\bar{S}$ genau dann wahr ist, wenn sowohl F als auch $\bar{S}$ gegeben sind, das heisst, wenn der Faktor F den Wert 1 und der Faktor S den Wert 0 annimmt.⁴ Von einzelnen Literalen einer komplexen Ursache wird gesagt, dass sie *kausal relevant* für eine Wirkung sind, womit man dem Umstand Rechnung trägt, dass ein Literal alleine typischerweise noch keine vollständige Ursache einer Wirkung darstellt, jedoch zusammen mit bestimmten anderen Literalen einen wesentlichen Teil zur Hervorbringung der Wirkung beiträgt (vgl. u. a. Mackie (1980), Kapitel 3). Die Darstellung von Alternativursachen ist ebenfalls der Booleschen Algebra entnommen und entspricht der disjunktiven Verknüpfung, die im Fall der Disjunktion $F\bar{S} + M$ genau dann wahr ist, wenn $F\bar{S}$, M oder beide gegeben sind.⁵ Wiederum mittels Wertezuordnungen ausgedrückt

⁴Gebräuchliche alternative Schreibweisen sind unter anderem „ $F \cdot \bar{S}$ “, „ $F \wedge \bar{S}$ “ oder „ $F \cap \bar{S}$ “ (vgl. u. a. Schneider & Wagemann (2012), S. 54).

⁵Weitere gebräuchliche Notationen sind „ $F\bar{S} \vee M$ “ oder „ $F\bar{S} \cup M$ “ (vgl. u. a. Schneider & Wagemann (2012), S. 54).

ist $F\bar{S} + M$ also genau dann wahr, wenn F den Wert 1 und S den Wert 0 oder wenn M den Wert 1 annimmt.

Die bis hierhin erfassten Kausalstrukturen, bestehend aus an- und abwesenden Faktoren, die zusammen als komplexe Ursachen oder als Teile von Alternativursachen eine Wirkung hervorrufen, entsprechen jener kausalen Komplexität, die in der Regel von QCA erfasst wird. Wie im einleitenden Beispiel aus Tabelle 1.1 gezeigt, interessiert man sich dabei üblicherweise für das Zusammenspiel von Literalen, die das Auftreten einer Wirkung begünstigen oder behindern. Bei QCA wird diesbezüglich von „equifinality“ und „conjunctural causation“ gesprochen, welche die Komplexität der untersuchten Kausalstrukturen charakterisieren (vgl. Schneider & Wagemann (2012), S. 78 f.). Während das erste dieser beiden Prädikate besagt, dass mehrere alternative Ursachen dieselbe Wirkung hervorrufen können, wird mit dem zweiten Prädikat ausgedrückt, dass eine solche Ursache aus mehreren Literalen besteht, die ihren Effekt auf eine Wirkung nur bei gemeinsamer Anwesenheit entfalten.

Die von QCA betrachteten Strukturen haben immer nur eine Wirkung, Kausalstrukturen in der Welt, in der wir leben, können aber auch mehrere Wirkungen haben. Erstens können verschiedene Wirkungen eine *gemeinsame Ursache* besitzen. Die Automatisierung (M) ist zum Beispiel nicht nur kausal relevant für die Zunahme der Arbeitslosigkeit (A), sondern auch für die Steigerung der Produktivität (P). In diesem Fall ist M sowohl kausal relevant für A als auch kausal relevant für P . Besteht zusätzlich kein kausaler Zusammenhang zwischen A und P , ergibt sich daraus insgesamt eine *Common-Cause-Struktur*.

Weiter enthält eine Kausalstruktur auch mehrere Wirkungen, wenn sie eine *Wechselwirkung* beschreibt: Die anhaltende Erwerbslosigkeit (E) kann zu Resignation (R) führen und umgekehrt verringert die Resignation die Aussichten auf einen Arbeitsplatz. In diesem Fall wäre also E kausal relevant für R und R wiederum kausal relevant für E .

Ferner kann sich die kausale Relevanz eines Literals über mehrere, als *Kausalkette* angeordnete Literale erstrecken. Literale können also sowohl direkte wie auch indirekte Wirkungen haben. So ist beispielsweise die Automatisierung (M) kausal relevant für die Zunahme der Arbeitslosigkeit (A) und A seinerseits kausal relevant für die steigende Nachfrage bei Arbeitslosenprogrammen (N). Neben der direkten kausalen Relevanz zwischen M und A sowie A und N besteht weiter eine indirekte kausale Relevanz zwischen M und N .

Schliesslich können sich in Bezug auf Kausalstrukturen mit mehreren Wirkungen, ähnlich wie bei den Wechselwirkungen, Verkettungen zu *kausalen Zyklen* zusammenschliessen. Um das zu illustrieren, wird das obige Beispiel der Wechselwirkung zwischen der anhaltenden Erwerbslosigkeit (E) und der Resignation (R) um die „Dequalifikation“ (D) erweitert, die den Verlust

der beruflichen Fähigkeiten durch fehlende Praxis und Weiterbildung beschreiben soll. Diese drei Literale lassen sich derart zu einem Zyklus zusammenschließen, dass E direkt kausal relevant für R sowie indirekt kausal relevant für D ist, R direkt kausal relevant für D sowie indirekt kausal relevant für E ist und D direkt kausal relevant für E sowie indirekt kausal relevant für R ist. Wechselwirkungen und kausale Zyklen werden künftig zusammenfassend auch als *Feedbackstrukturen* bezeichnet.

2.1.4. Graphennotation

Ein praktisches Hilfsmittel zur Darstellung von kausalen Abhängigkeiten sind Graphen. In der Graphentheorie wird ein Graph durch eine Menge von Knoten und eine Menge von Knoten verbindende Kanten definiert (vgl. u. a. Diestel (2010); Gross & Yellen (2004)). Knoten und Kanten bilden die Grundbausteine für Graphen, die vor allem zur Darstellung von Entitäten und deren Verbindungen in den unterschiedlichsten Gebieten zum Einsatz kommen. Zur Darstellung von Kausalstrukturen eignen sich die sogenannten *gerichteten Hypergraphen* (vgl. u. a. Ausiello et al. (2001)). Gerichtete Hypergraphen stellen eine Verallgemeinerung der gerichteten Graphen, kurz *Digraphen*, dar. Letztere zeichnen sich durch ihre gerichteten Kanten aus, die jeweils zwei Knoten derart miteinander verbinden, dass einer der verbundenen Knoten den sogenannten Fuss und der andere den sogenannten Kopf darstellt, wobei die Kante vom Fuss weg hin zum Kopf gerichtet ist. Graphisch werden diese Kanten üblicherweise durch Pfeile repräsentiert, deren Pfeilspitze beim Kopf der Verbindung angesiedelt ist. Interpretiert man die Knoten eines Graphen als Literale und die Pfeile als „ist direkt kausal relevant für“, lässt sich mit solchen Digraphen die kausale Beziehung zwischen einer Ursache A und der Wirkung B darstellen, wobei A dem Fuss und B dem Kopf der Kante entspricht. Graph (a) aus Abbildung 2.1 zeigt den dazugehörigen Digraphen. Ist anstelle des positiven Literals A das negative Literal $\bar{A}$ eine Ursache von B , lässt sich dies auf ähnliche Weise darstellen, mit dem Unterschied, dass das im Kreis enthaltene „ A “ durch „ $\bar{A}$ “ ersetzt wird. Eine derartige graphische Separierung von positivem und negativem Literal kann jedoch teils zu sehr unübersichtlichen Graphen führen. Stattdessen wird zur Darstellung einer Kausalbeziehung, bei der ein negatives Literal $\bar{X}$ involviert ist, ein ungefülltes Quadrat „ $\diamond$ “ zwischen den X enthaltenden Knoten und die entsprechende Kante gesetzt. So stellen die Digraphen (b) aus Abbildung 2.1 oben beginnend den direkten Kausalzusammenhang zwischen der Ursache $\bar{A}$ und der Wirkung B , der Ursache A und der Wirkung $\bar{B}$ sowie der Ursache $\bar{A}$ und der Wirkung $\bar{B}$ dar. Die horizontale Ausrichtung der Pfeile (bzw. die vertikale Ausrichtung im Digraphen (a)) hat dabei keine Bedeutung.

In einem Digraphen kann ein Knoten in mehreren gerichteten Kanten vorkommen, wobei jede Kante jeweils genau einen Knoten aufweist, der dem Fuss entspricht, und genau einen Knoten,

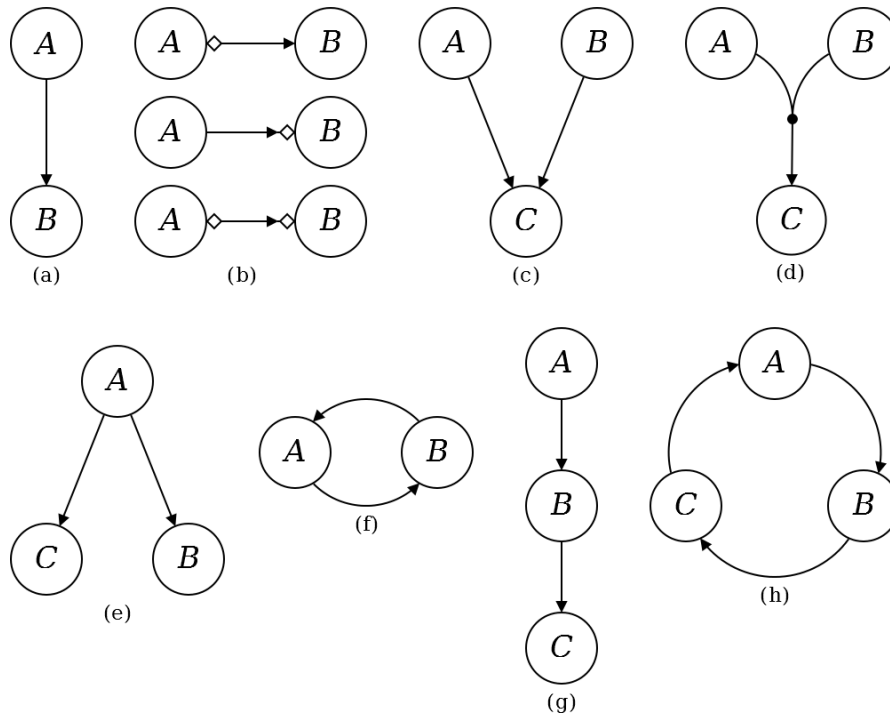


Abbildung 2.1.: Katalog von Kausalgraphen, bestehend aus (a) der direkten Verursachung, (b) dem Kausalzusammenhang mit Einbezug von abwesenden Faktoren, (c) Alternativursachen, (d) der komplexen Verursachung, (e) der Common-Cause-Struktur, (f) der Wechselwirkung, (g) der Kausalkette und (h) dem kausalen Zyklus.

der den Kopf darstellt (vgl. Gross & Yellen (2004), S. 1 und S. 6). In einem gerichteten Hypergraphen kann im Gegensatz dazu ein Fuss einer Kante nicht nur aus genau einem, sondern auch aus mehreren Knoten bestehen. Solche Kanten werden als *gerichtete Hyperkanten* bezeichnet, wobei jede gerichtete Kante auch eine gerichtete Hyperkante ist, aber nicht umgekehrt (vgl. Ausiello et al. (2001), S. 314). Das in der vorliegenden Arbeit verwendete graphische Merkmal zur Darstellung einer gerichteten Hyperkante, die keine gerichtete Kante ist, sind die von jedem Knoten am Fuss ausgehenden Linien, die in einem Bündelungspunkt „•“ zusammenlaufen und als einzelne Pfeilspitze in den als Kopf auftretenden Knoten münden. Die Verwendung von Hyperkanten ermöglicht eine klare Unterscheidung zwischen alternativen und komplexen Ursachen einer Wirkung, die, wie etwa am obigen Beispiel der Arbeitslosigkeit veranschaulicht, beide in ein und derselben Kausalstruktur vorkommen können. So wird der Hypergraph (c) aus Abbildung 2.1, der auch ein Digraph ist, als Kausalstruktur interpretiert, die zwei Alternativursachen A und B von C darstellt. Der Hypergraph (d), der kein Digraph ist, gibt den Zusammenhang zwischen der komplexen Ursache AB und ihrer Wirkung C wieder. Weiter stellt der Hypergraph (bzw. Digraph) (e) eine Common-Cause-Struktur mit gemeinsamer Ursache A von B und C dar.

Die graphische Notation eines indirekten Kausalzusammenhangs zeichnet sich durch einen sogenannten *Pfad* aus, der mehr als zwei Knoten enthält. Aus graphentheoretischer Sicht ist ein Pfad eine Sequenz von Knoten und Kanten, sodass der jeweils zwischen zwei Kanten e_i und e_j positionierte Knoten einer solchen Sequenz Kopf von e_i und Fuss beziehungsweise Teil des Fusses von e_j ist. Graph (g) aus Abbildung 2.1 zeigt einen Pfad, der kausal interpretiert eine Kausalkette beschreibt, bei der A direkt kausal relevant für B , B direkt kausal relevant für C und A indirekt kausal relevant für C ist. Ist der erste Knoten eines Pfades identisch mit dem letzten, ergibt sich, wie in (f) und (h) graphisch dargestellt, eine zyklische Verkettung. Kausal interpretiert handelt es sich dabei um Feedbackstrukturen, wobei (f) einer kausalen Wechselwirkung zwischen A und B entspricht und (h) einen kausalen Zyklus darstellt.

Ein gerichteter Hypergraph, der kausale Abhängigkeiten repräsentiert, wird als *Kausalgraph* bezeichnet. Jeder Kausalgraph ist ein gerichteter Hypergraph, nicht jeder gerichtete Hypergraph jedoch auch ein Kausalgraph. Das liegt daran, dass ein Kausalgraph gewissen Einschränkungen unterliegt, die ein Hypergraph im Allgemeinen nicht aufweist. So ist etwa ein gerichteter Hypergraph mit einer Kante, deren Kopf X und deren Fuss $\bar{X}$ ist, kein Kausalgraph. Eine solche Schleife würde kausal interpretiert zum Ausdruck bringen, dass die Anwesenheit eines Faktors kausal relevant für seine eigene Abwesenheit ist. Weshalb diese und weitere Einschränkungen gelten, wird sich in den folgenden Abschnitten im Zusammenhang mit den Merkmalen des Verursachungsbegriffs zeigen.

Abbildung 2.1 kann als Zusammenstellung der jeweils einfachsten Form der verschiedenen Typen von kausalen Verknüpfungen angesehen werden.⁶ Jeder Kausalgraph beliebiger Komplexität besteht aus einer endlichen Anzahl dieser Grundelemente. Beispielsweise lässt sich die Zunahme der Arbeitslosigkeit (A), die neben der Automatisierung von Arbeitsprozessen (M) durch den Wegzug grosser Firmen (F) in Kombination mit dem Fehlen eines Sozialplanes für die Entlassenen ($\bar{S}$) verursacht wird, durch den Kausalgraphen aus Abbildung 2.2 darstellen. Verwendet werden dafür die drei Grundstrukturen (b), (c) und (d) aus Abbildung 2.1.

2.1.5. Faktormengen

Kausale Erklärungen von Wirkungen sind typischerweise unvollständig und kausale Abhängigkeiten werden in den meisten Fällen relativ zu einer Menge von in Betracht gezogenen Faktoren, von nun an als *Faktormenge* bezeichnet, formuliert und untersucht.⁷ Auch die Unterteilung der

⁶Zu einem späteren Zeitpunkt wird noch ein weiterer Knotentyp zur Darstellung von kausalen Abhängigkeiten eingeführt. Wie dieser Knoten aussieht und weshalb er eingeführt wird, ist Thema von Abschnitt 2.11.2.

⁷Die Faktormenge wird auch „factor frame“ genannt (vgl. bspw. Baumgartner (2013a), S. 15).

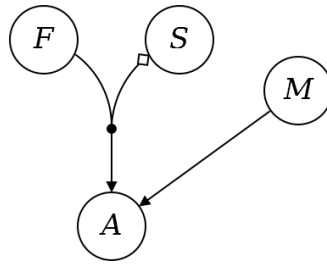


Abbildung 2.2.: Graphische Darstellung der Kausalzusammenhänge für die Zunahme der Arbeitslosigkeit.

kausalen Relevanz in die direkte und die indirekte kausale Relevanz ist jeweils in Beziehung zur Faktormenge zu setzen. Für die Kausalkette (g) aus Abbildung 2.1 gilt beispielsweise, dass *A* relativ zur Faktormenge $\{A, B, C\}$ indirekt kausal relevant und relativ zur Faktormenge $\{A, C\}$ direkt kausal relevant für *C* ist. Faktormengen werden durch fette, nicht-kursive Grossbuchstaben **A**, **B**, **B**₁, **B**₂ etc. dargestellt.⁸

Nicht jede beliebige Faktormenge eignet sich für die Formulierung und Untersuchung von Kausalzusammenhängen. Von den Faktoren einer betrachteten Faktormenge muss unter anderem verlangt werden, dass die natürlichen Eigenschaften, die von den Faktoren repräsentiert werden, logisch und konzeptionell unabhängig voneinander sind sowie keine metaphysischen Abhängigkeiten wie Supervenienz oder konstitutive Zusammenhänge zwischen ihnen bestehen (vgl. u. a. Baumgartner (2013b), S. 93 f.). Welchen Kriterien diese von Faktoren einer Faktormenge repräsentierten Eigenschaften allerdings im Detail zu genügen haben, ist allgemein noch ungeklärt. Deshalb wird typischerweise auf die Angabe eines umfassenden Kriterienkataloges verzichtet und lediglich darauf hingewiesen, dass es ungeeignete Faktormengen geben kann (vgl. bspw. Spirtes et al. (2000), S. 92). Bei dieser pragmatischen Herangehensweise wird stattdessen einfach vorausgesetzt, dass sich die betrachteten Faktormengen für kausales Modellieren eignen (vgl. u. a. Hitchcock (2007), S. 502 f.). Angesichts des Schwerpunktes der vorliegenden Arbeit wird diese Annahme auch hier getroffen. Neben der Natürlichkeit von durch Faktoren einer Faktormenge repräsentierten Eigenschaften wird davon ausgegangen, dass diese Eigenschaften keine problematischen Abhängigkeiten wie jene der obigen Form aufweisen oder kurz, dass sie *modal unabhängig* voneinander sind.

Kausalstrukturen werden typischerweise relativ zu Faktormengen formuliert und untersucht, die nicht alle Faktoren enthalten, deren positive und/oder negative Literale kausal relevant für Literale anderer Faktoren aus der Faktormenge sind. Dabei können Faktoren, deren Literale kausal relevant für Literale anderer Faktoren der betrachteten Faktormenge sind, aus verschiedenen

⁸Der Buchstabe **F** wird als Stellvertreter für eine beliebige Faktormenge zur Formulierung allgemeiner Aussagen reserviert.

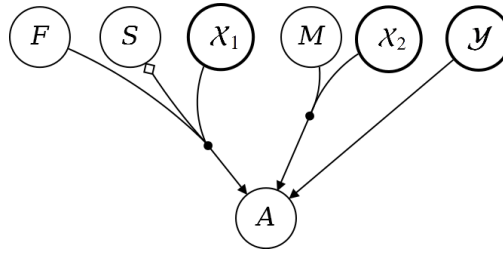


Abbildung 2.3.: Verallgemeinerte Darstellung der Kausalbeziehung aus Abbildung 2.2.

Gründen selbst nicht in der Faktormenge enthalten sein, das heisst latent sein. Beispielsweise dann, wenn sie noch unbekannt sind oder wenn sie aufgrund ihrer geringen Bedeutung bezüglich eines bestimmten Forschungsinteresses nicht explizit aufgeführt werden. Die relevanten Literale solcher Faktoren werden in Platzhaltern X_i und Y_j zusammengefasst (vgl. u. a. Mackie (1980), S. 66 f.). X_i repräsentiert weitere Konjunktionen von Literalen, die zusammen mit Literalen der bereits im Fokus stehenden Faktoren eine Ursache bilden. Solche mit anderen Literalen eine Ursache bildende Literale werden auch *Ko-Faktorliterale* oder kurz *Ko-Literale* genannt. Y_j repräsentiert eine Disjunktion von weiteren Konjunktionen, die latenten Alternativursachen entsprechen. Der Kausalgraph aus Abbildung 2.2 kann damit zum Kausalgraph aus Abbildung 2.3 verallgemeinert werden, wobei die Platzhalter, die weitere Ko-Literale und latente Alternativursachen darstellen, im Graphen durch die dicke Umrandung zusätzlich hervorgehoben werden.

Der so entstehende Kausalgraph repräsentiert im Unterschied zum Kausalgraph aus Abbildung 2.2 nicht nur kausal relevante Literale der Wirkung A, sondern vollständige Ursachen. X_1 , X_2 und Y können dabei mehrere alternative Bündel von Literalen repräsentieren, sodass M beispielsweise im Fall von X_2 zusammen mit I oder zusammen mit anderen Literalen J und K eine komplexe Ursache bildet. Löst man X_2 entsprechend auf, führt dies zum Kausalgraphen aus Abbildung 2.4, welcher der Vollständigkeit halber zwei neue Platzhalter X_3 und X_4 enthält. Die komplexe Kausalbeziehung ist dann vollständig aufgedeckt, wenn jeder Platzhalter des entsprechenden Kausalgraphen die leere Menge repräsentiert.

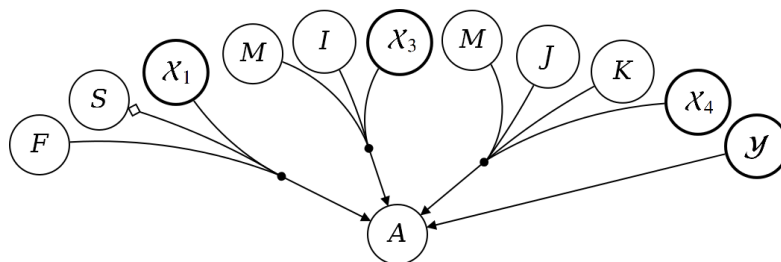


Abbildung 2.4.: Darstellung der Kausalbeziehung aus Abbildung 2.3 mit weiterer Spezifizierung von X_2 .

Wie im Graph aus Abbildung 2.4 bereits erkennbar, werden Darstellungen von Kausalzusammenhängen mit der Angabe jedes Platzhalters $\mathcal{X}_i$ und $\mathcal{Y}_j$ schnell unübersichtlich. Im weiteren Verlauf wird deshalb oft auf eine jeweilige explizite Darstellung dieser Platzhalter verzichtet und stattdessen angenommen, dass Kausalstrukturen erweiterbar sind. Damit ist gemeint, dass aus einer relativ zur Faktormenge $\mathbf{F}_1$ bestehenden Kausalstruktur mit dem Einbezug von weiteren, in $\mathbf{F}_1$ nicht enthaltenen Faktoren aus $\mathbf{F}_2$ relativ zu $\mathbf{F}_3 = \mathbf{F}_1 \cup \mathbf{F}_2$ eine Kausalstruktur resultieren kann, in der zusätzlich Literale der Faktoren aus $\mathbf{F}_2$ verortet sind (vgl. u. a. Baumgartner (2008b)).

2.2. Suffizienz, Notwendigkeit und Redundanzfreiheit

Die sich aus Abschnitt 2.1.2 ergebende Komplexität von Kausalstrukturen gibt einen Überblick über die grundlegenden Typen der mittels konfiguralen Methoden untersuchten Kausalzusammenhänge. Vom Verursachungsbegriff wird gefordert, dass er in der Lage ist, diese kausalen Beziehungen zu erfassen. Konfigurale Methoden wie QCA oder CNA vergleichen zur Aufdeckung solcher Kausalbeziehungen zwischen Literalen von Faktoren einer Menge $\mathbf{F} = \{A, B, C, \dots\}$ die verschiedenen, empirisch möglichen Konfigurationen dieser Faktoren miteinander. Eine Konfiguration besteht aus Faktoren, die in geeigneter räumlicher und zeitlicher Nähe gemeinsam an- oder abwesend sind. In der metaphysischen Auffassung eines Aktualismus verortet, sind dabei jene Konfiguration empirisch möglich, die mindestens einmal im Laufe der Existenz des analysierten Systems (in einem atemporalen Sinn) auftreten.⁹ Welche Konfigurationen das genau sind, ist primitiv, das heisst nicht weiter erklärbar. Beschreibt beispielsweise die Kausalstruktur aus Abbildung 2.2 die zwischen A , F , S und M bestehenden Kausalzusammenhänge, entspricht die Anwesenheit aller vier Faktoren A , F , S und M einer empirisch möglichen Konfiguration. Diese Konfiguration widerspiegelt eine Zusammenfassung von Fällen, wie beispielsweise Länder oder Wirtschaftsregionen, die alle durch A , F , S und M repräsentierten Eigenschaften aufweisen. Im Gegensatz dazu beschreibt vor dem Hintergrund derselben Kausalstruktur und unter der Voraussetzung, dass $\mathcal{X}_2$ gegeben ist, etwa der anwesende Faktor M zusammen mit der Abwesenheit des Faktors A keine empirisch mögliche Konfiguration: Da M zusammen mit $\mathcal{X}_2$ eine vollständige Ursache von A bildet, ist bei gegebenem M und gegebenen Ko-Literalen $\mathcal{X}_2$ in geeigneter raumzeitlicher Nähe auch A anwesend.

Die geeignete raumzeitliche Nähe ist ein unscharfer Begriff und variiert je nach Kontext von unmittelbarer räumlicher und zeitlicher Angrenzung bis hin zu evolutionären Abständen (vgl. bspw. Hume (1739), S. 75 f.; Graßhoff & May (2001), S. 92; Baumgartner (2008b), S. 349 f.).

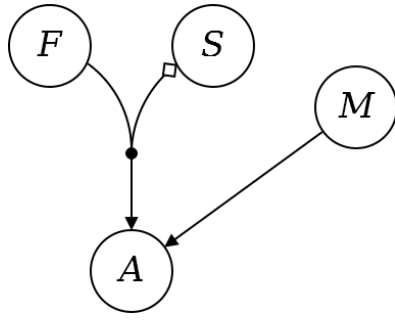
⁹Zum Aktualismus vgl. bspw. Adams (1974), Menzel (1990).

Für die kausalen Abhängigkeiten aus Abbildung 2.2 wäre beispielsweise die Instanz von M , die eine zu Beginn des Jahres 2018 durchgeführte Automatisierung in einer Firma X am Standort Y beschreibt, in geeigneter raumzeitlicher Nähe zur Instanz von A , die der Zunahme der Arbeitslosigkeit desselben Jahres in der Region um X entspricht. Nicht in geeigneter raumzeitlicher Nähe wären etwa die gleiche Instanz von M zusammen mit der Zunahme der Arbeitslosigkeit um das Jahr 2000 in einer Region, die nie in einer Geschäftsbeziehung mit der Firma X stand. Aufgrund der starken kontextabhängigen Variation ist fraglich, ob sich etwas Generelles über den Begriff der Nähe sagen lässt. Obwohl wünschenswert, wird auch hier aufgrund des analytischen Schwerpunktes der vorliegenden Arbeit auf eine vertiefte Analyse des Begriffs der geeigneten Nähe verzichtet. Stattdessen soll davon ausgegangen werden, dass bei den Beispielen dieser Arbeit jeweils klar ist, wie der Begriff der Nähe zu verstehen ist, und dass die Beispiele diesem genügen. Dementsprechend wird die raumzeitliche Nähe im weiteren Verlauf der Arbeit oftmals auch gar nicht mehr explizit erwähnt.

Konfigurationale Methoden nutzen den Umstand, dass sich bei kausalen Abhängigkeiten in den dazugehörigen Konfigurationen sichtbare Regelmässigkeiten im An- und Abwesenheitsverhalten der Faktoren zeigen. Ist beispielsweise X die alleinige Ursache von W , zeigt sich relativ zur Faktormenge $\mathbf{A} = \{X, W\}$, dass immer wenn X gegeben ist, in geeigneter raumzeitlicher Nähe auch W gegeben ist, das heisst X ist *hinreichend* für W . Umgekehrt ist immer wenn W gegeben ist, in geeigneter raumzeitlicher Nähe auch X gegeben, das heisst X ist *notwendig* für W . Ist X eine hinreichende Bedingung von W , wird dies durch die Boolesche Operation der Implikation „ $X \rightarrow W$ “ ausgedrückt.¹⁰ Sind X und W so miteinander verknüpft, dass X eine notwendige Bedingung von W ist, wird dies notiert als „ $X \leftarrow W$ “. Ist X genau dann in geeigneter raumzeitlicher Nähe gegeben, wenn W gegeben ist, das heisst ist X sowohl notwendig als auch hinreichend für W , wird dies unter Verwendung des Bikonditionals dargestellt als „ $X \leftrightarrow W$ “. Mithilfe der Wertezuordnungen ausgedrückt, ist $X \rightarrow W$ genau dann wahr, wenn X den Wert 0 oder W den Wert 1 annimmt. Falsch ist der Ausdruck genau dann, wenn X den Wert 1 und W den Wert 0 annimmt. Analoges gilt für die Notwendigkeit von X für W , wobei die Wertezuordnungen von X und W vertauscht werden. $X \leftarrow W$ schliesslich ist genau dann wahr, wenn X und W entweder beide den Wert 0 oder beide den Wert 1 annehmen.

Angesichts der kausalen Komplexität wird es kaum der Fall sein, dass jeweils mit einer Instanz von X auch eine Instanz von W gegeben ist oder immer wenn W instantiiert auch X instantiiert ist. Der kausale Zusammenhang zwischen X und der Wirkung W ist nicht immer dadurch gekennzeichnet, dass X hinreichend und/oder notwendig für W ist. Zu sehen ist dies bereits am einfachen Beispiel der steigenden Arbeitslosigkeit aus Abbildung 2.2. Abbildung 2.5 zeigt den

¹⁰ Alternativ wird diese Verknüpfung auch als „ $X \subset W$ “ oder „ $X \Rightarrow W$ “ dargestellt (vgl. Schneider & Wagemann (2012), S. 53 f.).



$\mathcal{K}_{2,5}$	F	S	M	A
c_1	1	1	1	1
c_2	1	1	0	0
c_3	1	0	1	1
c_4	1	0	0	1
c_5	0	1	1	1
c_6	0	1	0	0
c_7	0	0	1	1
c_8	0	0	0	0

Abbildung 2.5.: Kausalgraph mit der entsprechenden Konfigurationstabelle $\mathcal{K}_{2,5}$ über $\mathbf{B} = \{A, F, M, S\}$ der Kausalstruktur aus Abbildung 2.2.

entsprechenden Kausalgraphen zusammen mit allen dazugehörigen Konfigurationen, wobei davon ausgegangen wird, dass die Ko-Literale X_1 von $F\bar{S}$ beziehungsweise X_2 von M jeweils gegeben sind (vgl. Abbildung 2.3). Des Weiteren soll zur besseren Veranschaulichung angenommen werden, dass $\mathcal{Y}$ leer ist und somit keine weiteren Alternativursachen von A existieren. Unter diesen Voraussetzungen gibt es insgesamt acht verschiedene Konfigurationen der Faktoren aus der Menge $\mathbf{B} = \{A, F, M, S\}$, die mit den dargestellten kausalen Abhängigkeiten vereinbar sind. Diese Konfigurationen sind in der durch den kalligraphischen Kurzbezeichner „ $\mathcal{K}_{2,5}$ “ gekennzeichneten Tabelle aus Abbildung 2.5 aufgelistet.¹¹ Jede Zeile c_i , $1 \leq i \leq 8$, aus $\mathcal{K}_{2,5}$ über $\mathbf{B} = \{A, F, M, S\}$ entspricht einer der empirisch möglichen Konfigurationen der Faktoren A , F , M und S , wobei die An- beziehungsweise Abwesenheit eines Faktors mit einer 1 beziehungsweise 0 in der entsprechenden Spalte wiedergegeben wird.

F bildet zusammen mit $\bar{S}$ eine komplexe Ursache $F\bar{S}$ von A , weshalb mit F auch $\bar{S}$ gegeben sein muss, damit A auftritt. In $\mathcal{K}_{2,5}$ zeigt sich dies unter anderem bei der Konfiguration c_2 , bei der zwar F , nicht aber A gegeben ist, weil der Faktor S nicht abwesend ist. Des Weiteren existieren aufgrund der Alternativursache M von A zwei Konfigurationen c_5 und c_7 , bei denen zwar A , nicht aber F gegeben ist. Es zeigt sich, dass, obwohl F kausal relevant für A ist, F weder hinreichend

¹¹Das kalligraphische $\mathcal{K}$ steht für „Konfigurationen“ und ist mit jenem Index versehen, welcher der Nummer der Tabelle oder der Abbildung entspricht, in der die Konfigurationen zu finden sind. Diese Schreibweise für Konfigurationstabellen als Ganzes wird künftig zugunsten einer besseren Übersicht regelmässig eingesetzt. Weil Konfigurationstabellen dabei jeweils in Tabellen- oder Abbildungsumgebungen eingebettet sein können, werden Tabellen und Abbildungen über denselben Zähler durchnummeriert. Auf diese Weise wird sichergestellt, dass jede Konfigurationstabelle „ $\mathcal{K}_{x,y}$ “ eindeutig einer entsprechenden Tabelle oder Abbildung x,y zugeordnet werden kann. Für die Formulierungen von allgemeinen Aussagen, die sich auf kein spezifisches Beispiel einer Konfigurationstabelle beziehen, wird im weiteren Verlauf oft auch der Bezeichner „ $\mathcal{K}_F$ “ verwendet. Diese Kurzform steht für eine Konfigurationstabelle über die Faktormenge $\mathbf{F}$.

noch notwendig für A ist. Umgekehrt bestehen Bedingungen, die hinreichend und notwendig für A sind, nicht zwangsläufig nur aus kausal relevanten Literalen. So sind etwa gemäss $\mathcal{K}_{2.5}$ die Bedingungen FSM (vgl. c_1), $F\bar{S}M$ (c_3), $F\bar{S}\bar{M}$ (c_4), $\bar{F}SM$ (c_5) und $\bar{F}\bar{S}M$ (c_7) alle hinreichend für A , da immer, wenn eine davon gegeben ist, auch A gegeben ist. Des Weiteren gilt, dass immer, wenn A gegeben ist, auch FSM , $F\bar{S}M$, $F\bar{S}\bar{M}$, $\bar{F}SM$ oder $\bar{F}\bar{S}M$ gegeben sind. Insgesamt ergeben sich also folgende Suffizienz- und Notwendigkeitsbeziehungen, wobei beispielsweise mit $\bar{F}$ und S Literale in den hinreichenden und notwendigen Bedingungen von A enthalten sind, die nicht kausal relevant für A sind:

$$FSM + F\bar{S}M + F\bar{S}\bar{M} + \bar{F}SM + \bar{F}\bar{S}M \leftrightarrow A \quad (2.1)$$

Kausale Relevanz kann also nicht einfach über Mitgliedschaft in einer hinreichenden oder notwendigen Bedingung definiert werden. Damit stellt sich die Frage, wie kausale Relevanz mit auf Suffizienz und Notwendigkeit basierenden Regularitäten in Anbetracht der Komplexität von Kausalstrukturen aus Abschnitt 2.1.2 zu definieren ist.

Ein erneuter Blick auf die notwendige Bedingung von A aus (2.1) zeigt, dass jede darin enthaltene hinreichende Bedingung von A auch mindestens eine tatsächliche Ursache von A als Teil aufweist. In jeder Bedingung kommt $F\bar{S}$ oder M vor. So enthält etwa die eben explizit erwähnte hinreichende Bedingung $\bar{F}SM$ von A die tatsächliche Ursache M von A . Offensichtlich sind also die Ursachen von A in der Regelmässigkeit aus (2.1) enthalten, zusätzlich dazu sind jedoch noch überflüssige Teile vorhanden. Überflüssig deshalb, weil etwa die hinreichende Bedingung $\bar{F}SM$ auch dann noch hinreichend ist, wenn man $F\bar{S}$ daraus entfernt, sodass nur M übrigbleibt. Denn gemäss den Konfigurationen c_1 , c_3 , c_5 und c_7 aus $\mathcal{K}_{2.5}$ ist immer, wenn M gegeben ist, auch A gegeben. Daraus folgt, dass etwa auch die folgende Regelmässigkeit das Verhalten der Faktoren aus $\mathcal{K}_{2.5}$ einfängt:

$$FSM + F\bar{S}M + F\bar{S}\bar{M} + M + \bar{F}\bar{S}M \leftrightarrow A \quad (2.2)$$

Genau darin liegt der entscheidende Unterschied zwischen Teilen einer auf Suffizienz und Notwendigkeit basierenden Regelmässigkeit, die tatsächlich kausale Abhängigkeiten einfangen, und Teilen, die dies nicht tun: Auf letztere kann für die Beschreibung des Instantiierungsverhaltens verzichtet werden, während erstere einen unverzichtbaren Beitrag leisten. Denn würde man umgekehrt aus der hinreichenden Bedingung $\bar{F}SM$ von A das tatsächlich für A kausal relevante Literal M entfernen, ergibt sich folgende Regelmässigkeit:

$$FSM + F\bar{S}M + F\bar{S}\bar{M} + \bar{F}S + \bar{F}\bar{S}M \leftrightarrow A \quad (2.3)$$

Mit $\bar{F}S$ ist in (2.3) eine Bedingung enthalten, die gemäss c_6 aus $\mathcal{K}_{2.5}$ gegeben ist, ohne dass auch A anwesend ist. Anders als von (2.3) behauptet, ist $\bar{F}S$ also nicht hinreichend für A , weshalb mit dieser Regelmässigkeit das Auftrittsverhalten der Faktoren F , S , M und A nicht mehr korrekt wiedergegeben wird.

Kausal relevante Literale leisten unverzichtbare Beiträge im Hinblick auf die Instanziierung ihrer Wirkungen. Das Instanzierungsverhalten von A in den Konfigurationen c_5 und c_6 beispielsweise kann nur mit M erklärt werden, das in c_5 gegeben und in c_6 nicht gegeben ist. M ist also in mindestens einem Kontext entscheidend für die Beschreibung des Auftrittsverhaltens von A – kurz: M ist ein sogenannter *Difference-Maker* (vgl. u. a. Baumgartner (2015), S. 841). Welche Eigenschaften ein Literal zu einem solchen Difference-Maker machen, wird sich in der Folge dieses Kapitels zeigen. Für W kausal irrelevante Literale hingegen haben keinen Einfluss auf das Auftrittsverhalten von W . Egal ob diese Literale gegeben oder nicht gegeben sind, es bleibt unter allen Umständen ohne Konsequenzen für W – kurz: Diese Faktoren sind *redundant*. Insgesamt ergibt sich daraus das folgende fundamentale Prinzip der Redundanzfreiheit, dem auf Suffizienz und Notwendigkeit basierende Regelmässigkeiten zu genügen haben, um kausal interpretierbar zu sein:

Prinzip der Redundanzfreiheit (RF): Nur Regelmässigkeiten, deren Teile unverzichtbar für die Beschreibung des Instanzierungsverhaltens von Faktoren sind, fangen kausale Abhängigkeiten ein.

2.3. INUS-Theorie

Die bekannteste Theorie, die Ursache-Wirkung-Beziehungen beliebiger Komplexität mithilfe von notwendigen und hinreichenden Bedingungen einzufangen versucht, ist die INUS-Theorie des Philosophen John Leslie Mackie (vgl. Mackie (1980)). Gemäss der INUS-Theorie ist jede Ursache eines Kausalzusammenhangs mindestens ein „Insufficient but Non-redundant part of an Unnecessary but Sufficient condition“ oder kurz: mindestens eine INUS-Bedingung (vgl. Mackie (1980), S. 62). Um zu verstehen, was genau eine INUS-Bedingung ist, wird erneut das Beispiel der steigenden Arbeitslosigkeit aus Abbildung 2.5 betrachtet.

Es wurde bereits festgestellt, dass F nicht hinreichend für A ist (vgl. Konfiguration c_2 aus $\mathcal{K}_{2.5}$). Analog dazu verhält es sich nicht so, dass immer, wenn $\bar{S}$ gegeben ist, auch A gegeben ist: Bei der Konfiguration c_8 steigt die Arbeitslosigkeit trotz dem Ausbleiben von S nicht. Anders verhält es sich jedoch, wenn zusammen mit F ein Sozialplan fehlt ($\bar{S}$), das heisst, wenn die Konjunktion $F\bar{S}$ gegeben ist (vgl. Konfigurationen c_3 und c_4). Hier wird gemäss der Konfigurationstabelle $\mathcal{K}_{2.5}$ auch immer die Arbeitslosigkeit zunehmen. Während also F und $\bar{S}$ alleine nicht hinreichend für A sind, ist immer, wenn $F\bar{S}$ gegeben ist, auch A anwesend. Sowohl F als auch $\bar{S}$ sind dabei nicht-redundante Teile der hinreichenden Bedingung $F\bar{S}$ von A . Das bedeutet, dass, wenn eines der beiden Literale aus der Konjunktion $F\bar{S}$ entfernt wird, der verbleibende Teil nicht mehr

hinreichend für A ist. Weiter ist aus der Konfigurationstabelle $\mathcal{K}_{2.5}$ zu entnehmen, dass $F\bar{S}$ nicht notwendig für A ist, da es mit c_1 , c_5 sowie c_7 Konfigurationen gibt, bei denen zwar A , nicht aber $F\bar{S}$ gegeben ist. Für F und $\bar{S}$ gilt damit insgesamt, dass sie für sich genommen nicht hinreichend, aber nicht-redundante Teile der hinreichenden Bedingung $F\bar{S}$ von A sind, wobei die Konjunktion $F\bar{S}$ nicht notwendig für A ist – kurz: F und $\bar{S}$ sind INUS-Bedingungen von A .

Zur anschaulichen Beschreibung der INUS-Bedingungen anhand der Kausalstruktur aus Abbildung 2.5 sind wir oben der Einfachheit halber davon ausgegangen, dass alle Ko-Literale von $F\bar{S}$ und M jeweils gegeben sind. In der Praxis kann es jedoch durchaus vorkommen, dass in einigen Fällen beispielsweise trotz gegebener Konjunktion $F\bar{S}$ die Wirkung A nicht auftritt, weil mindestens ein unbeachtetes Literal aus X_1 nicht gegeben ist. Zwischen den eine Ursache bildenden Literalen von Faktoren einer Menge F und deren Wirkung W besteht nicht unter allen Umständen eine Regularität. Die Regularität besteht vielmehr nur im Kontext von gleich bleibenden Hintergrundbedingungen, die im Beispiel zur Arbeitslosigkeit etwa dann vorliegen, wenn bei allen Konfigurationen jeweils alle Literale aus X_1 gegeben sind. Diese Hintergrundbedingungen bilden gemäss Mackie das sogenannte *kausale Feld*, relativ zu dem kausale Hypothesen jeweils formuliert werden (vgl. Mackie (1980), S. 34 f.).¹²

Der INUS-Theorie zufolge ist X (bzw. $\bar{X}$) genau dann kausal relevant für W , wenn X (bzw. $\bar{X}$) im kausalen Feld $\hat{F}$ eine INUS-Bedingung oder selbst eine notwendige oder hinreichende Bedingung von W ist, wobei die Faktoren X und W verschieden sind (vgl. Mackie (1980), S. 62 ff.; Baumgartner & Graßhoff (2004), S. 94 f.). Eine Konjunktion von Literalen, die alle INUS-Bedingungen von W sind, wird auch als *minimal hinreichende Bedingung* von W bezeichnet (vgl. Mackie (1980), S. 62). Ein gültiger Spezialfall stellt dabei jener dar, bei dem eine minimal hinreichende Bedingung nur aus einem Literal besteht. In diesem Fall ist das Literal keine INUS-Bedingung, sondern für sich genommen bereits hinreichend oder notwendig für W . Damit werden kausale Abhängigkeiten eingefangen, bei denen einzelne Literale vollständige Ursachen von Wirkungen darstellen. Von der INUS-Theorie ausgeschlossen wird jedoch die Selbstverursachung, was aus der geforderten Verschiedenheit von X und W folgt.¹³ Ohne Ausschluss der

¹²Die Voraussetzung des kausalen Feldes ist auch bekannt als *ceteris-paribus*-Klausel (vgl. Graßhoff & May (2001), S. 89). Was genau die Merkmale eines solchen kausalen Feldes sind, wird in Kapitel 3 unter dem Begriff „Homogenität“ beschrieben.

¹³Die Selbstverursachung ist zu unterscheiden von Feedbackstrukturen, wie bspw. einer Wechselwirkung zwischen zwei Literalen A und B . Bei der Selbstverursachung ruft ein Literal C sich direkt selbst hervor und ist damit nur von sich selbst kausal abhängig. Im Gegensatz dazu tritt A (bzw. B) bei der Wechselwirkung nur über den kausalen Effekt von B (bzw. A) auf.

Selbstverursachung würde sich kein gehaltvoller Verursachungsbegriff definieren lassen, zumal jedes Literal trivialerweise hinreichend und notwendig für sich selbst ist.¹⁴

Im Beispiel zur Arbeitslosigkeit aus Abbildung 2.5 kann neben den beiden INUS-Bedingungen F und $\bar{S}$ von A gefolgert werden, dass M hinreichend ist für A . Denn immer wenn M gegeben ist, ist auch A gegeben (vgl. Konfigurationen c_1, c_3, c_5 sowie c_7). Ferner ist M nicht notwendig für A , da es mit c_4 aus $\mathcal{K}_{2.5}$ eine Konfiguration gibt, bei der zwar A , nicht aber M gegeben ist. Ist jedoch A gegeben, ist immer $F\bar{S}$ oder M gegeben, womit die Disjunktion aller minimal hinreichenden Bedingungen (im kausalen Feld $\hat{F}$) hinreichend und notwendig für die Wirkung A ist. Für die notwendige Disjunktion von minimal hinreichenden Bedingungen nutzt Mackie den Begriff der „full cause“ (vgl. Mackie (1980), S. 64). Unter Berücksichtigung des kausalen Feldes $\hat{F}$, das durch den Zusatz „In $\hat{F}$:“ dargestellt wird, ergibt sich aus $\mathcal{K}_{2.5}$ das folgende Bikonditional:

$$\text{In } \hat{F}: F\bar{S} + M \leftrightarrow A \quad (2.4)$$

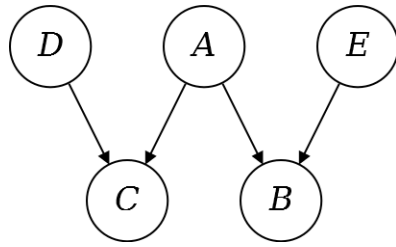
(2.4) erfasst die untersuchte Kausalstruktur korrekt und beinhaltet nur jene Literale mindestens als INUS-Bedingungen von A , die auch tatsächlich kausal relevant für A sind.

Mackie setzt die von (RF) geforderte Redundanzfreiheit für hinreichende Bedingungen um, indem er letztere derart von redundanten Literalen befreit, dass aus der übrigbleibenden Bedingung kein Literal mehr entfernt werden kann, ohne das die resultierende Bedingung ihre Suffizienz verliert. Obwohl dies zwar dazu führt, dass (2.4) die Ursachen von A aus Abbildung 2.5 korrekt einfängt, zeigt das folgende Beispiel aus Abbildung 2.6, dass die INUS-Theorie im Allgemeinen nicht in der Lage ist, alle Typen von komplexen Kausalstrukturen korrekt abzubilden. Im Unterschied zum Beispiel aus Abbildung 2.5 sind in dieser Common-Cause-Struktur nicht nur eine, sondern zwei Wirkungen involviert, wobei A sowohl für C als auch für B kausal relevant ist, zwischen B und C aber keine kausale Abhängigkeit besteht.

Aus Symmetriegründen ist es für den Moment ausreichend, sich auf die Analyse der Ursachen von B zu beschränken. Gemäss dem Kausalgraphen handelt es sich dabei um A und E , die zwei Alternativursachen von B darstellen.

Die Konfigurationstabelle $\mathcal{K}_{2.6}$ über $\mathbf{C} = \{A, B, C, D, E\}$ zeigt, dass immer, wenn A gegeben ist, auch B gegeben ist (vgl. Konfigurationen c_1 bis c_4). Das heisst, dass A hinreichend für B ist. Weiter ist immer, wenn E anwesend ist, B ebenfalls anwesend (c_1, c_3, c_5 und c_6). Auch E ist somit hinreichend für B . Ferner sind sowohl A als auch E minimal hinreichende Bedingungen

¹⁴Andere Kausaltheorien, die nicht auf Regularitäten fussen, wären von dieser Trivialität in gleichem Masse betroffen. Beispielsweise erhöht im Rahmen einer probabilistischen Kausaltheorie ein Ereignis trivialerweise die Eintrittswahrscheinlichkeit dieses Ereignisses (vgl. Baumgartner (2006), S. 25).



$\mathcal{K}_{2,6}$	A	B	C	D	E
c_1	1	1	1	1	1
c_2	1	1	1	1	0
c_3	1	1	1	0	1
c_4	1	1	1	0	0
c_5	0	1	1	1	1
c_6	0	1	0	0	1
c_7	0	0	1	1	0
c_8	0	0	0	0	0

Abbildung 2.6.: Common-Cause-Struktur zusammen mit der entsprechenden Konfigurationstabelle $\mathcal{K}_{2,6}$, die alle Konfigurationen der Faktoren aus $\mathbf{C} = \{A, B, C, D, E\}$ auflistet.

(vgl. c_7 und c_8).¹⁵ Schliesslich gibt es mit $C\bar{D}$ eine letzte hinreichende Bedingung von B , da immer, wenn die Konjunktion $C\bar{D}$ gegeben ist (vgl. c_3 und c_4), auch B gegeben ist. Dabei ist weder C (vgl. c_7) noch $\bar{D}$ (vgl. c_8) hinreichend für B , womit auch $C\bar{D}$ minimal hinreichend für B ist. Diese drei minimal hinreichenden Bedingungen sind disjunktiv verknüpft notwendig für B . Immer, wenn B gegeben ist, ist auch A , E oder $C\bar{D}$ gegeben. Insgesamt resultiert (im kausalen Feld $\hat{F}$) das folgende Bikonditional:

$$A + E + C\bar{D} \leftrightarrow B \quad (2.5)$$

(2.5) gibt die kausalen Abhängigkeiten aus Abbildung 2.6 nicht korrekt wieder. Zusätzlich zu A und E sind auch C und $\bar{D}$ mindestens INUS-Bedingungen von B und stellen gemäss der INUS-Theorie Ursachen von B dar. Mackie selbst erkennt dieses Problem und beschreibt es basierend auf einem Beispiel des Philosophen C. D. Broad (1887–1971), wonach zum durch heulende Sirenen signalisierten Zeitpunkt 17.00 (A) sowohl die Fabrikarbeiter in Manchester ihre Arbeit niederlegen (C) als auch die Fabrikarbeiter in London (B) (vgl. Mackie (1980), S. 81 ff.; Broad (1925), S. 455 f.). Unter Verwendung dieser Interpretationen von A , B und C ist gemäss der INUS-Theorie die Arbeitsniederlegung in Manchester kausal relevant für die Arbeitsniederlegung in London.

Notwendige und hinreichende Bedingungen, deren Bestandteile der INUS-Theorie genügen, werden oftmals als Suchziel von QCA angegeben (vgl. Schneider & Wagemann (2012), S. 79 f.; Ragin (2008), S. 154). Mit dem Beispiel aus Abbildung 2.6 zeigt sich jedoch, dass komplexe Kausalstrukturen existieren, für die gemäss der INUS-Theorie Literale kausal relevant für

¹⁵Existiert mindestens eine Konfiguration, bei der die untersuchte Wirkung nicht auftritt, gilt allgemein, dass hinreichende Bedingungen, die aus nur einem Literal bestehen, auch minimal hinreichend sind.

andere sind, obwohl tatsächlich keine kausale Abhängigkeit zwischen diesen besteht. Die INUS-Theorie liefert demnach keinen adäquaten Verursachungsbegriff für das konfigurationale kausale Modellieren.

2.4. RDN-Bikonditional

Die von der INUS-Theorie geforderte Redundanzfreiheit bezweckt, dass jede hinreichende Bedingung einer Wirkung frei von redundanten Literalen ist. Dabei bleibt jedoch jener Fall unberücksichtigt, bei dem nicht nur Teile relativ zu einer hinreichenden Bedingung redundant sind, sondern die gesamte hinreichende Bedingung an sich redundant ist. Die Konjunktion $C\bar{D}$ des Bikonditionals (2.5) stellt genau eine solche redundante hinreichende Bedingung dar. Immer wenn $C\bar{D}$ gegeben ist (vgl. c_3 und c_4 aus $\mathcal{K}_{2.6}$), ist auch die minimal hinreichende Bedingung A vorhanden. Im Gegensatz zu $C\bar{D}$ kann jedoch für die Beschreibung des Instantiierungsverhaltens auf A nicht verzichtet werden, zumal sonst die Anwesenheit von B in c_2 unbegründet bleiben würde. Denn in c_2 ist B gegeben, ohne dass auch $C\bar{D}$ oder E gegeben sind, womit bei dieser Konfiguration nur A das Auftreten von B erklären kann. Konjunkte einer minimal hinreichenden Bedingung können nur dann unverzichtbar für die Beschreibung des Instantiierungsverhaltens von Faktoren sein, wenn die minimal hinreichende Bedingung selbst kein redundanter Teil einer notwendigen Bedingung von B darstellt. Dabei ist eine redundanzfreie notwendige Bedingung genau dann gegeben, wenn sie notwendig ist und kein Disjunkt aus der Bedingung entfernt werden kann, ohne dass die daraus resultierende Bedingung ihre Notwendigkeit für die Wirkung verliert. Die notwendige Bedingung $A + E + C\bar{D}$ von B ist nicht redundanzfrei, da $A + E$ bereits notwendig für B ist: Immer wenn B gegeben ist, ist A oder E gegeben. Die Disjunktion $A + E$ ihrerseits enthält kein redundantes Disjunkt, und es lässt sich weder A (sichtbar bei c_2 und c_4) noch E (c_5 , c_6) eliminieren, ohne dass die resultierende Bedingung ihre Notwendigkeit für B verliert. Fordert man neben der Redundanzfreiheit von hinreichenden Bedingungen auch die Redundanzfreiheit von notwendigen Bedingungen, ergibt sich aus $\mathcal{K}_{2.6}$ folgendes Bikonditional (vgl. Graßhoff & May (2001), S. 95 f. und S. 101 ff.):

$$A + E \leftrightarrow B \quad (2.6)$$

Durch die Hinzunahme der Redundanzfreiheit von notwendigen Bedingungen fällt $C\bar{D}$ aus dem Bikonditional und es bleiben jene Literale übrig, die gemäss Abbildung 2.6 tatsächlich Ursachen von B sind.

Aus Konfigurationen gefolgerte Regelmässigkeiten fangen nur dann Kausalzusammenhänge ein, wenn sie frei von jeglichen Redundanzen sind. Um dieser von (RF) geforderten maximalen Redundanzfreiheit, auch maximale Sparsamkeit oder Minimalität genannt, gerecht zu werden, gilt

es nicht nur von hinreichenden, sondern auch von notwendigen Bedingungen Redundanzfreiheit zu fordern. Im Folgenden werden zwei Definitionen für die beiden Bedingungsformen präsentiert. Dazu werden jeweils auch nochmals explizit die in Abschnitt 2.1 angenommene Natürlichkeit sowie modale Unabhängigkeit der Eigenschaften aufgeführt, die die Faktoren repräsentieren (vgl. S. 15 und S. 22). Mit der modalen Unabhängigkeit der Eigenschaften geht auch der Ausschluss der Selbstverursachung einher, die bereits von der INUS-Theorie ausgeschlossen wird.

Definition 2.1 (Minimal hinreichende Bedingung)

Sei ϵ ein Konjunktionsterm von Literalen $Z_1 \cdots Z_n$ mit $n \geq 1$. ϵ ist genau dann eine minimal hinreichende Bedingung von W , wenn

- (a) die Faktoren aus ϵ und der Faktor W natürliche und modal voneinander unabhängige Eigenschaften repräsentieren,
- (b) $\epsilon \rightarrow W$ gilt und
- (c) für keinen echten Teil ϵ' von ϵ gilt: $\epsilon' \rightarrow W$.¹⁶

Ein aus Literalen bestehender Konjunktionsterm ist eine Konjunktion, die vom gleichen Faktor entweder das positive oder das negative Literal enthält (vgl. u. a. Richter et al. (1997), S. 93). Entspricht l der Anzahl Literale eines Konjunktionsterms ϵ , ist ein echter Teil ϵ' von ϵ ein in ϵ enthaltener Konjunktionsterm mit k Literalen, wobei $0 \leq k < l$ ist. Ist AB eine hinreichende Bedingung von W , kann ABC keine minimal hinreichende Bedingung von W sein und enthält mindestens den redundanten Teil C (vgl. z. B. Graßhoff & May (2001), S. 94). Ein einzelnes Literal besitzt mit dem Konjunktionsterm, der kein Literal enthält, einen einzigen echten Teil. Dieser sogenannte 0-stellige Konjunktionsterm kann gemäss Definition 2.1 keine minimal hinreichende Bedingung ϵ sein, da ϵ definiert ist als Konjunktionsterm mit mindestens einem Literal ($n \geq 1$).¹⁷

¹⁶Angesichts der in Abschnitt 2.1.2 eingeführten Konvention (vgl. S. 16), gemäss der kursive Grossbuchstaben wie Z_i , $1 \leq i \leq n$, oder W positive Literale beschreiben, sind mit obiger Definition streng genommen jene Fälle, bei denen eine minimal hinreichende Bedingung ausschliesslich aus negativen Literalen besteht oder eine Bedingung minimal hinreichend für ein negatives Literal ist, nicht eingeschlossen. Diese Fälle sind (und werden auch künftig) alleine aus Gründen einer möglichst kompakten und übersichtlichen Darstellung nicht explizit aufgeführt. Jedes Literal Z_i bzw. W lässt sich ohne Weiteres durch $\bar{Z}_i$ bzw. $\bar{W}$ ersetzen, womit sich u. a. minimal hinreichende Bedingungen für die ausdrückliche Abwesenheit eines Faktors W beschreiben lassen.

¹⁷Zum Begriff des 0-stelligen Konjunktionsterms vgl. u. a. Richter et al. (1997), S. 93. Der 0-stellige Konjunktionsterm kommt bspw. dann zum Einsatz, wenn ein Bikonditional der Form $X + \bar{X} \leftrightarrow W$ vorliegt: Wird ein Literal X bzw. $\bar{X}$ auf dessen kausale Relevanz für eine Wirkung W untersucht und liegen für die Kausalanalyse Konfigurationen vor, bei denen die Wirkung W immer auftritt, kann $X + \bar{X} \leftrightarrow W$ gefolgert werden. Sowohl X als auch $\bar{X}$ sind

Definition 2.2 (Minimal notwendige Bedingung)

Sei δ eine Disjunktion (in disjunktiver Normalform) von Literalen $Z_1 \cdots Z_k + \dots + Z_m \cdots Z_n$ mit $n \geq 1$. δ ist genau dann eine minimal notwendige Bedingung von W , wenn

- (a) die Faktoren aus δ und der Faktor W natürliche und modal voneinander unabhängige Eigenschaften repräsentieren,
- (b) $\delta \leftarrow W$ gilt und
- (c) für keinen echten Teil δ' von δ gilt: $\delta' \leftarrow W$.

Eine notwendige Bedingung ist in disjunktiver Normalform (DNF), wenn es sich um eine Disjunktion von Konjunktionstermen handelt (vgl. bspw. Büning & Lettmann (1994), S. 42). Entspricht l der Anzahl Disjunkte einer disjunktiven Normalform δ , ist ein echter Teil δ' von δ eine aus k in δ enthaltenen Disjunkten bestehende Disjunktion, wobei $0 < k < l$ ist. Ist $AB + C$ eine notwendige Bedingung von W , kann $AB + C + \bar{D}E$ keine minimal notwendige Bedingung von W sein und enthält mindestens das redundante Disjunkt $\bar{D}E$ (vgl. z.B. Graßhoff & May (2001), S. 96).

Die beiden Bedingungen aus den Definitionen 2.1 und 2.2 werden weiter zur *redundanzfreien Bedingung in disjunktiver Normalform*, kurz *RDN-Bedingung*, zusammengefasst:

Definition 2.3 (Redundanzfreie Bedingung in disjunktiver Normalform (RDN-Bedingung))

Eine redundanzfreie Bedingung in disjunktiver Normalform (RDN-Bedingung) ϕ_W von W ist eine aus minimal hinreichenden Bedingungen bestehende minimal notwendige Bedingung von W .

Kommt X als Konjunkt in mindestens einer minimal hinreichenden Bedingung einer RDN-Bedingung von W vor, ist garantiert, dass X relativ zu den restlichen Literalen der RDN-Bedingung von W unverzichtbar für die Beschreibung des Instantiierungsverhaltens von W ist. Wird X aus der RDN-Bedingung von W entfernt, ist der übriggebliebene Term nicht mehr notwendig für W oder enthält mindestens eine Bedingung, die nicht mehr hinreichend für W ist.

Schliesslich wird ein Bikonditional der Form von (2.6), das eine RDN-Bedingung von B über das Bikonditional mit B verknüpft und damit die aus $\mathcal{K}_{2,6}$ ableitbare redundanzfreie Regulari-

hinreichend, aber nicht minimal hinreichend für W , zumal es mit dem 0-stelligen Konjunktionsterm einen echten Teil von X bzw. $\bar{X}$ gibt, der hinreichend ist. Der Faktor W tritt immer auf und somit auch unabhängig davon, ob der Faktor X anwesend ist oder nicht. Unter diesen Umständen ist X bzw. $\bar{X}$ redundant und eine allfällige kausale Relevanz von X bzw. $\bar{X}$ für W nicht nachweisbar.

tät zwischen diesen Faktoren wiedergibt, als *RDN-förmiges Bikonditional* (*RDN-Bikonditional*) von B bezeichnet. Ob ein RDN-Bikonditional $\phi_W \leftrightarrow W$ das Instantiierungsverhalten von Faktoren einer Menge $\mathbf{F}$ korrekt einfängt, wird dabei relativ zu einer Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ bestimmt. Beschreibt $\phi_W \leftrightarrow W$ das Instantiierungsverhalten aus $\mathcal{K}_{\mathbf{F}}$ korrekt, wird gesagt, dass $\phi_W \leftrightarrow W$ von $\mathcal{K}_{\mathbf{F}}$ impliziert wird beziehungsweise $\phi_W \leftrightarrow W$ aus $\mathcal{K}_{\mathbf{F}}$ folgt oder $\phi_W \leftrightarrow W$ ein RDN-Bikonditional relativ zu $\mathcal{K}_{\mathbf{F}}$ ist. Letzteres lässt sich mit einer sogenannten *Wahrheitstabelle* oder kurz *Wahrheitstabelle* beziehungsweise *Wertetabelle* demonstrieren, die in der Booleschen Algebra oft zum Einsatz kommt. Tabelle 2.7 zeigt zwei solche, zugunsten des vorliegenden Zwecks etwas angepasste Wertetabellen für $F\bar{S} + M \leftrightarrow A$ und $A + E \leftrightarrow B$. In Spalte (i) sind die Konfigurationen aus den entsprechenden Konfigurationstabellen aufgelistet, die jenes An- und Abwesenheitsverhalten der Faktoren wiedergeben, das auf derselben Zeile in Spalte (ii) dargestellt ist. Eine 1 beziehungsweise 0 aus Spalte (ii) lässt sich wie bis anhin als „ist anwesend/gegeben“ beziehungsweise „ist nicht anwesend/gegeben“ interpretieren. Auch das RDN-Bikonditional in Spalte (iii) der jeweiligen Tabelle nimmt die Werte 1 oder 0 an. Hier wird eine 1 beziehungsweise 0 interpretiert als „ist gegeben/gilt“ oder „ist wahr“ beziehungsweise „ist nicht gegeben/gilt nicht“ oder „ist falsch“. Die Werte dieser letzten Spalte (iii) ergeben sich, indem man das Bikonditional für die jeweilige Zeile Spalte (ii) auswertet. Relativ zu einem Ausdruck oder Teilausdruck werden die Zeilen aus Spalte (ii) oftmals auch als *Belegungen* bezeichnet. Das RDN-Bikonditional $F\bar{S} + M \leftrightarrow A$ aus Tabelle 2.7a nimmt also beispielsweise für jene Belegung, bei der allen Faktoren der Wert 1 zugewiesen wird, auch den Wert 1 an. Indem $F\bar{S} + M \leftrightarrow A$ auch für alle anderen Belegungen der Faktoren, die einer Konfiguration aus der jeweiligen Konfigurationstabelle entsprechen, wahr ist, stellt sich insgesamt heraus, dass $F\bar{S} + M \leftrightarrow A$ das Instantiierungsverhalten der Faktoren F , S , M und A aus $\mathcal{K}_{2,5}$ korrekt wiedergibt.¹⁸ Dasselbe gilt gemäss der Wertetabelle 2.7b für $A + E \leftrightarrow B$ relativ zu $\mathcal{K}_{2,6}$ und dem darin abgebildeten Instantiierungsverhalten der Faktoren A , B und E .

Allgemein lässt sich ein RDN-Bikonditional folgendermassen definieren:

Definition 2.4 (*RDN-förmiges Bikonditional* (*RDN-Bikonditional*))

Sei $\mathcal{K}_{\mathbf{F}}$ eine Konfigurationstabelle über eine Faktormenge $\mathbf{F}$. Ein von $\mathcal{K}_{\mathbf{F}}$ impliziertes Bikonditional $\phi_W \leftrightarrow W$ ist genau dann ein RDN-Bikonditional von W relativ zu $\mathcal{K}_{\mathbf{F}}$, wenn ϕ_W eine RDN-Bedingung von W ist.

X ist genau dann Teil eines RDN-Bikonditionals von W , wenn X als Konjunkt in mindestens einer minimal hinreichenden Bedingung einer minimal notwendigen Bedingung von W vorkommt oder kurz, wenn X in einer RDN-Bedingung von W enthalten ist.

¹⁸Mehr zur Wahrheitswertanalyse und dem damit verbundenen mathematischen Fundament findet man im Kapitel 4.

(i)	(ii)				(iii)				
$\mathcal{K}_{2.5}$	F	S	M	A	$F\bar{S} + M \leftrightarrow A$				
c_1	1	1	1	1	1				
c_2	1	1	0	0	1				
c_3	1	0	1	1	1				
c_4	1	0	0	1	1				
c_5	0	1	1	1	1				
c_6	0	1	0	0	1				
c_7	0	0	1	1	1				
c_8	0	0	0	0	1				

(a)

(i)	(ii)			(iii)
$\mathcal{K}_{2.6}$	A	E	B	$A + E \leftrightarrow B$
c_1, c_3	1	1	1	1
c_2, c_4	1	0	1	1
c_5, c_6	0	1	1	1
c_7, c_8	0	0	0	1

(b)

Tabelle 2.7.: Wertetabellen 2.7a und 2.7b mit allen in der jeweiligen Spalte (ii) aufgelisteten Belegungen, für die das entsprechende Bikonditional aus Spalte (iii) wahr ist.

Ein hier als RDN-Bikonditional von W bezeichneter Ausdruck $\phi_W \leftrightarrow W$ wird in der bestehenden Literatur üblicherweise als *einfache Minimale Theorie* von W bezeichnet (vgl. Graßhoff & May (2001), S. 96; Baumgartner (2013b), S. 90 f.). Der Begriff der einfachen Minimalen Theorie ist dabei derart eng an den Begriff der kausalen Relevanz geknüpft, dass es sich etabliert hat, letzteren unter Verwendung von ersterem zu definieren. Konkret wird typischerweise gesagt, dass X notwendigerweise Teil einer einfachen Minimalen Theorie von W sein muss, damit man X kausale Relevanz für W zusprechen kann. Die vorliegende Arbeit soll bezüglich dieser Handhabung, den Begriff der kausalen Relevanz basierend auf dem Begriff der einfachen Minimalen Theorie zu definieren, keine Ausnahme darstellen. Das Bikonditional $\phi_W \leftrightarrow W$ aus Definition 2.4 wird in der vorliegenden Arbeit nun aber deshalb nicht als einfache Minimale Theorie von W bezeichnet, weil die bisher geforderte Redundanzfreiheit noch nicht ausreicht, um die kausale Relevanz eines Literals basierend auf Konfigurationen einzufangen. Weshalb genau das so ist, wird sich gleich im nächsten Abschnitt zeigen. Nichtsdestotrotz bringen Literale, die Teil eines RDN-Bikonditionals von W sind, eine für die kausale Interpretierbarkeit notwendige Eigenschaft mit sich: Für all diese Literale gibt es sogenannte *Difference-Making-Paare*. Betrachtet man beispielsweise das RDN-Bikonditional $F\bar{S} + M \leftrightarrow A$, zeigt sich in den Konfigurationen aus $\mathcal{K}_{2.5}$ etwa für F , dass es mit c_4 und c_8 zwei Konfigurationen gibt, in denen M jeweils nicht gegeben und $\bar{S}$ jeweils gegeben ist, sodass die Variation im Auftretensverhalten des Faktors A alleine auf die Variation des Faktors F zurückgeht. Das heisst, relativ zu c_4 und c_8 ist A deshalb in c_4 gegeben und in c_8 nicht, weil F in c_4 gegeben und in c_8 nicht gegeben ist. Auch für M gibt es zum Beispiel mit c_6 und c_8 in $\mathcal{K}_{2.5}$ Konfigurationen, für die die Bedingung $F\bar{S}$ jeweils nicht gegeben ist, indem der Faktor F abwesend ist, sodass die Variation des Faktors A alleine auf die Variation des Faktors M zurückgeht.

Ein Difference-Making-Paar kann folgendermassen definiert werden, wobei zur übersichtlicheren Darstellung die beiden bereits bekannten Platzhalter $\mathcal{X}$ und $\mathcal{Y}$ verwendet werden, mit deren Hilfe sich ein Bikonditional $AZ_1 \cdots Z_i + Z_j \cdots Z_k + \cdots + Z_m \cdots Z_n \leftrightarrow W$ via Substitutionen $\mathcal{X} := Z_1 \cdots Z_i$ und $\mathcal{Y} := Z_j \cdots Z_k + \cdots + Z_m \cdots Z_n$ darstellen lässt als $A\mathcal{X} + \mathcal{Y} \leftrightarrow W$ (vgl. S. 22):

Definition 2.5 (Difference-Making-Paar)

Sei $\delta \leftrightarrow W$ ein von $\mathcal{K}_F$ impliziertes Bikonditional und A derart Teil von δ , dass $\delta := A\mathcal{X} + \mathcal{Y}$ ist. Ein Difference-Making-Paar von A bezüglich W relativ zu $A\mathcal{X} + \mathcal{Y}$ ist ein Paar von Konfigurationen c_i und c_j aus $\mathcal{K}_F$, sodass A und W in c_i gegeben und in c_j nicht gegeben sind, während $\mathcal{X}$ in beiden Konfigurationen gegeben und $\mathcal{Y}$ in beiden Konfigurationen nicht gegeben ist.

$\mathcal{X} := Z_1 \cdots Z_i$ ist genau dann gegeben, wenn alle darin enthaltenen Literale $Z_1, \dots, Z_i$ gegeben sind. $\mathcal{Y} := Z_j \cdots Z_k + \cdots + Z_m \cdots Z_n$ ist genau dann gegeben, wenn mindestens ein Disjunkt gegeben ist, das seinerseits genau dann gegeben ist, wenn alle im Disjunkt enthaltenen Literale gegeben sind.

Notwendige und hinreichende Bedingungen fangen dann kausale Abhängigkeiten ein, wenn sie frei von redundanten Teilen sind. Diesem Prinzip der Redundanzfreiheit (RF) zufolge soll A nur dann kausale Relevanz für W zugesprochen werden, wenn A unverzichtbar für die Beschreibung des Instantiierungsverhalten der W beinhaltenden Faktormenge ist. Bei Literalen, die Teil eines RDN-Bikonditionals sind, zeigt sich diese Unverzichtbarkeit in Form von Difference-Making-Paaren. Für jedes Literal A eines RDN-Bikonditionals von W gilt, dass es ein Difference-Making-Paar von A bezüglich W gibt oder anders ausgedrückt: Ist

- (i) $\phi_W \leftrightarrow W$ ein RDN-Bikonditional und
- (ii) für jedes Literal aus ϕ_W existiert ein Difference-Making-Paar bezüglich W ,

gilt die Äquivalenz (i) $\Leftrightarrow$ (ii). Dieser Link zwischen Redundanzfreiheit und Difference-Making-Paar wurde bislang in der Literatur noch nicht explizit gemacht. Die Äquivalenz (i) $\Leftrightarrow$ (ii) formalisiert erstmals zu Teilen jene grundlegende Idee von Regularitätstheorien, dass jede Ursache in mindestens einem Kontext einen entscheidenden Unterschied betreffend die Anwesenheit der Wirkung macht.

Die Äquivalenz (i) $\Leftrightarrow$ (ii) lässt sich durch die folgenden beiden Teilbeweise nachweisen. Dabei sei $\phi_W \leftrightarrow W$ ein von $\mathcal{K}_F$ impliziertes Bikonditional und A ein beliebiges Literal eines beliebigen Disjunktes aus ϕ_W , sodass $\phi_W := A\mathcal{X} + \mathcal{Y}$ ist.

(i) $\Rightarrow$ (ii): Sei $\phi_W \leftrightarrow W$ ein RDN-Bikonditional.

Indem $\phi_W \leftrightarrow W$ mit $\phi_W := A\mathcal{X} + \mathcal{Y}$ von $\mathcal{K}_F$ impliziert wird, gilt (1) für jede Konfiguration entweder (1a) $W = 1$ und ($A\mathcal{X} = 1$ oder $\mathcal{Y} = 1$) oder (1b) $W = 0$ und $A\mathcal{X} = 0$ und $\mathcal{Y} = 0$. Dabei gilt, dass es sowohl Konfigurationen gibt, die (1a) genügen, als auch Konfigurationen, die (1b) genügen. Andernfalls würde entweder $1 \leftrightarrow W$ oder $0 \leftrightarrow W$ gelten, was der Annahme widerspricht, dass $A\mathcal{X} + \mathcal{Y} \leftrightarrow W$ ein RDN-Bikonditional ist.

Weil $A\mathcal{X} + \mathcal{Y}$ eine minimal notwendige Bedingung von W ist, ist $\mathcal{Y}$ nicht notwendig für W . Das heisst, (2) es gibt Konfigurationen, für die gilt: $W = 1$ und $\mathcal{Y} = 0$. Aus (1) und (2) folgt: (3) Es gibt Konfigurationen, für die $W = 1$ und $A\mathcal{X} = 1$ und $\mathcal{Y} = 0$ ist.

Indem $A\mathcal{X}$ ein Disjunkt des RDN-Bikonditionals $A\mathcal{X} + \mathcal{Y} \leftrightarrow W$ ist, ist $A\mathcal{X}$ eine minimal hinreichende Bedingung von W . Das heisst, es gilt, dass (4) wenn $A\mathcal{X} = 1$, dann $W = 1$ und (5) es existieren Konfigurationen, sodass $\mathcal{X} = 1$ und $W = 0$. Bezüglich (5) gilt aufgrund von $W = 0$ und der Notwendigkeit von $A\mathcal{X} + \mathcal{Y}$ für W , dass für diese Konfigurationen $\mathcal{Y} = 0$ ist. Das heisst, (6) es gibt Konfigurationen, für die gilt: $\mathcal{X} = 1$ und $\mathcal{Y} = 0$ und $W = 0$. Für diese Konfigurationen muss aufgrund der Suffizienz von $A\mathcal{X}$ für W weiter gelten, dass $A = 0$. Daraus folgt zusammen mit (3), dass es für A mindestens ein Difference-Making-Paar gibt.

(i) $\Leftarrow$ (ii): Existiere für jedes Literal aus ϕ_W ein Difference-Making-Paar bezüglich W .

Indem $\phi_W \leftrightarrow W$ mit $\phi_W := A\mathcal{X} + \mathcal{Y}$ von $\mathcal{K}_F$ impliziert wird, ist $A\mathcal{X} + \mathcal{Y}$ notwendig und $A\mathcal{X}$ hinreichend für W . Weiter gibt es aufgrund der Existenz eines Difference-Making-Paars von A bezüglich W mindestens ein Paar von Konfigurationen c_i und c_j , sodass $A = W = 1$ in c_i und $A = W = 0$ in c_j , während $\mathcal{X} = 1$ und $\mathcal{Y} = 0$ in c_i und c_j . Aus c_j folgt, dass $\mathcal{X}$ nicht hinreichend ist für W und weil A ein beliebiges Literal aus einer hinreichenden Bedingung $A\mathcal{X}$ von W ist, wobei auch für jedes Literal aus $\mathcal{X}$ ein Difference-Making-Paar existiert, ist $A\mathcal{X}$ minimal hinreichend für W . Weiter folgt aus c_i , dass es Konfigurationen gibt, bei denen W gegeben ist und nur die minimal hinreichende Bedingung $A\mathcal{X}$ aus $A\mathcal{X} + \mathcal{Y}$ gegeben ist. Indem dies auch für jedes weitere Disjunkt aus $\mathcal{Y}$ gilt, kann kein Disjunkt aus der notwendigen Bedingung $A\mathcal{X} + \mathcal{Y}$ entfernt werden, ohne dass der resultierende Ausdruck seine Notwendigkeit für W verliert. Die aus minimal hinreichenden Bedingungen bestehende notwendigen Bedingung $A\mathcal{X} + \mathcal{Y}$ ist demnach auch minimal notwendig und $A\mathcal{X} + \mathcal{Y} \leftrightarrow W$ somit ein RDN-Bikonditional.

2.5. Strukturelle Redundanz

Ist A Teil eines RDN-Bikonditionals von W , ist garantiert, dass A relativ zu den restlichen Literalen der RDN-Bedingung von W in mindestens einem Fall den Unterschied im Hinblick auf das Vorhandensein von W ausmacht. Wird A aus der RDN-Bedingung von W entfernt, ist der übrigbleibende Term nicht mehr notwendig für W oder enthält mindestens eine Bedingung, die nicht mehr hinreichend für W ist. Mit dem Begriff des RDN-Bikonditionals ist somit sichergestellt, dass darin keine redundanten Teile mehr enthalten sind und für alle Literale Difference-Making-Paare existieren. Es bleibt aber ein weiterer Typ von Redundanz, der, sofern er unberücksichtigt bleibt, kausale Fehlschlüsse produzieren kann und deshalb ausgeschlossen werden muss. Dieser Typ von Redundanz wurde bisher von keiner Regularitätstheorie der Kausalität berücksichtigt.

Analog zu minimal hinreichenden Bedingungen, die als Ganzes redundant sind und damit den Begriff der minimalen Notwendigkeit motiviert haben, besteht die Möglichkeit, dass ein RDN-Bikonditional zwar relativ zu den darin enthaltenen Literalen keine redundanten Teile aufweist, das RDN-Bikonditional jedoch immer noch als Ganzes relativ zum Instantiierungsverhalten aller betrachteten Faktoren redundant ist. Was genau mit dieser Form der Redundanz gemeint ist und weshalb man sie zu berücksichtigen hat, soll anhand des Beispiels aus Abbildung 2.8 veranschaulicht werden. Es zeigt eine elektrische Schaltung mit den entsprechenden Faktordefinitionen.

Der Einfachheit halber soll davon ausgegangen werden, dass diese Schaltung bis auf die Schalterbewegungen von äusseren Einflüssen isoliert ist. Die Schaltung ist so aufgebaut, dass das Brennen der Glühlampe D verursacht wird durch den geschlossenen Schalter B zusammen mit dem Wechselschalter A , der nach oben zeigt. Neben dieser komplexen Ursache AB wird das

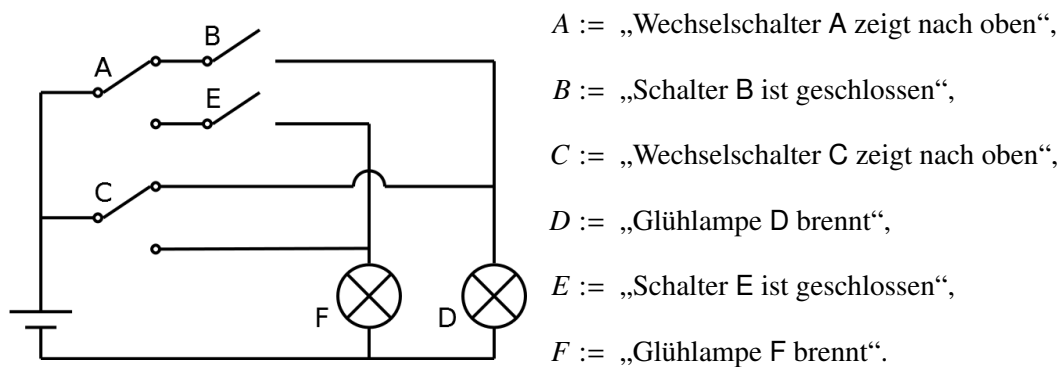
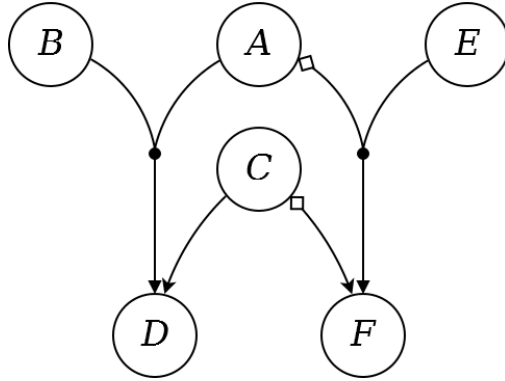


Abbildung 2.8.: Beispiel einer elektrischen Schaltung, die den kausalen Abhängigkeiten aus Abbildung 2.9 folgt.



$\mathcal{K}_{2,9}$	A	B	C	D	E	F
c_1	1	1	1	1	1	0
c_2	1	1	1	1	0	0
c_3	1	1	0	1	1	1
c_4	1	1	0	1	0	1
c_5	1	0	1	1	1	0
c_6	1	0	1	1	0	0
c_7	1	0	0	0	1	1
c_8	1	0	0	0	0	1
c_9	0	1	1	1	1	1
c_{10}	0	1	1	1	0	0
c_{11}	0	1	0	0	1	1
c_{12}	0	1	0	0	0	1
c_{13}	0	0	1	1	1	1
c_{14}	0	0	1	1	0	0
c_{15}	0	0	0	0	1	1
c_{16}	0	0	0	0	0	1

Abbildung 2.9.: Die der Schaltung aus Abbildung 2.8 zugrunde liegende Kausalstruktur zusammen mit der dazugehörigen Konfigurationstabelle $\mathcal{K}_{2,9}$ über $\mathbf{D} = \{A, B, C, D, E, F\}$.

Licht der Lampe D alternativ durch den nach oben gerichteten Wechselschalter C verursacht. Die Wechselschalter zeichnen sich dadurch aus, dass sie entweder nach oben oder nach unten zeigen, wobei beide Positionen einen geschlossenen Schalter ergeben. Das Brennen der Glühlampe F wird durch den nach unten gerichteten Wechselschalter C oder alternativ durch den geschlossenen Schalter E zusammen mit dem nach unten gerichteten Wechselschalter A verursacht. Stellt man diese kausalen Abhängigkeiten unter Berücksichtigung der obigen Faktordefinitionen als Kausalgraph dar, ergibt sich die Kausalstruktur aus Abbildung 2.9. Zusätzlich ist die zu dieser Kausalstruktur gehörige Konfigurationstabelle $\mathcal{K}_{2,9}$ dargestellt, die alle Konfigurationen der Faktoren aus $\mathbf{D} = \{A, B, C, D, E, F\}$ beinhaltet.

Aus $\mathcal{K}_{2,9}$ können die folgenden RDN-Bikonditionale abgeleitet werden:

$$AB + C \leftrightarrow D \quad (2.7)$$

$$\bar{C} + \bar{A}E \leftrightarrow F \quad (2.8)$$

$$\bar{F} + \bar{A}D \leftrightarrow C \quad (2.9)$$

Dieses Resultat kann mit dem entsprechenden Abschnitt des im Anhang A.3 dargestellten R-Skripts (vgl. S. 257 ff.), das durch Verwendung des Packages „cna“ innerhalb der Softwareumge-

bung „R“ ausgeführt werden kann (vgl. R Core Team (2016); Baumgartner & Ambühl (2019)), repliziert werden. Ähnliche noch folgende Berechnungen, die in der vorliegenden Arbeit nicht im Detail ausgeführt werden, mit denen jedoch ein gewisser Rechenaufwand einhergeht, sind ebenfalls in diesem R-Replikationsskript zu finden.

Ein Vergleich mit der Kausalstruktur aus Abbildung 2.9 zeigt, dass das RDN-Bikonditional (2.7) die Ursachen von D und das RDN-Bikonditional (2.8) die Ursachen von F korrekt wiedergibt. Beide RDN-Bikonditionale (2.7) und (2.8) sind isoliert betrachtet redundanzfrei. Nichtsdestotrotz greift die Forderung, dass nur Teile von RDN-Bikonditionalen kausal relevant sind, nach wie vor zu kurz, da es sich auch bei (2.9) um ein RDN-Bikonditional handelt. Immer wenn $\bar{A}D$ gegeben ist, ist in $\mathcal{K}_{2.9}$ auch C gegeben (vgl. $c_9, c_{10}, c_{13}, c_{14}$). Darüber hinaus gilt weder für $\bar{A}$ alleine (vgl. u. a. c_{11}) noch für D alleine (vgl. u. a. c_3), dass sie hinreichend für C sind, womit die Konjunktion $\bar{A}D$ minimal hinreichend für C ist. Dasselbe gilt für das negative Literal $\bar{F}$, das eine weitere minimal hinreichende Bedingung von C darstellt. Weiter gilt, dass immer, wenn C gegeben ist, auch $\bar{A}D$ oder $\bar{F}$ gegeben ist (vgl. $c_1, c_2, c_5, c_6, c_9, c_{10}, c_{13}$ und c_{14}), weshalb die Disjunktion $\bar{F} + \bar{A}D$ eine notwendige Bedingung von C ist. Diese notwendige Bedingung ist ferner auch minimal, weil keines der beiden Disjunkte $\bar{F}$ (vgl. u. a. c_1) und $\bar{A}D$ (vgl. u. a. c_{13}) aus $\bar{F} + \bar{A}D$ entfernt werden kann, ohne dass der resultierende Ausdruck seine Notwendigkeit für C verliert. Diese (innere) Redundanzfreiheit von (2.9) zeigt sich auch dadurch, dass es von allen Literalen $\bar{F}$ (vgl. u. a. c_2 und c_4 aus $\mathcal{K}_{2.9}$), $\bar{A}$ (u. a. c_3, c_9) und D (u. a. c_{13}, c_{15}) Difference-Making-Paare relativ zu C gibt.

Im Gegensatz zu den beiden RDN-Bikonditionalen (2.7) und (2.8) möchte man das RDN-Bikonditional (2.9) nicht kausal interpretieren, zumal man damit etwa behaupten würde, dass die brennende Glühlampe D kausal relevant für den nach unten gerichteten Wechselschalter C ist, obwohl der tatsächliche Kausalzusammenhang genau in umgekehrter Richtung verläuft. Vielmehr ist das RDN-Bikonditional (2.9) eine Regelmässigkeit, die sich aus den beiden anderen RDN-Bikonditionalen (2.7) und (2.8) ergibt. (2.9) ist derart von den beiden Ausdrücken (2.7) und (2.8) abhängig, dass (2.7) und (2.8) das RDN-Bikonditional (2.9) implizieren. Das heisst, die Konjunktion

$$(AB + C \leftrightarrow D) \cdot (\bar{C} + \bar{A}E \leftrightarrow F) \quad (2.10)$$

impliziert $\bar{F} + \bar{A}D \leftrightarrow C$.¹⁹ Der Nachweis dieser logischen Abhängigkeit liefert die Wahrheitstabelle 2.10, wobei aus Platzgründen nur genau jene der $2^6 = 64$ Belegungen der Faktoren A, B, C, D, E und F aufgelistet sind, für die der Ausdruck (2.10) wahr ist.

¹⁹Im Unterschied zur konjunktiven Verknüpfung von Literalen innerhalb eines RDN-Bikonditionals wird die konjunktive Verknüpfung von mehreren RDN-Bikonditionalen durch „ $\cdot$ “ symbolisiert. Dabei handelt es sich jedoch lediglich um einen Unterschied in der Notation, deren Einführung die bessere Lesbarkeit solcher Konjunktionen bezweckt.

A	B	C	D	E	F	$(AB + C \leftrightarrow D) \cdot (\bar{C} + \bar{A}E \leftrightarrow F)$	$\bar{F} + \bar{A}D \leftrightarrow C$
1	1	1	1	1	0	1	1
1	1	1	1	0	0	1	1
1	1	0	1	1	1	1	1
1	1	0	1	0	1	1	1
1	0	1	1	1	0	1	1
1	0	1	1	0	0	1	1
1	0	0	0	1	1	1	1
1	0	0	0	0	1	1	1
0	1	1	1	1	1	1	1
0	1	1	1	0	0	1	1
0	1	0	0	1	1	1	1
0	1	0	0	0	1	1	1
0	0	1	1	1	1	1	1
0	0	1	1	0	0	1	1
0	0	0	0	1	1	1	1
0	0	0	0	0	1	1	1

Tabelle 2.10.: Wahrheitstabelle mit allen Belegungen der Faktoren aus $\mathbf{F} = \{A, B, C, D, E, F\}$, für die $(AB + C \leftrightarrow D) \cdot (\bar{C} + \bar{A}E \leftrightarrow F)$ wahr ist.

Das Beispiel der elektrischen Schaltung zeigt, dass nicht alle aus Konfigurationen folgenden RDN-Bikonditionale kausal interpretierbar sind. RDN-Bikonditionale wie (2.9) verletzen (RF), weil die durch die Konjunktion der drei RDN-Bikonditionale

$$(AB + C \leftrightarrow D) \cdot (\bar{C} + \bar{A}E \leftrightarrow F) \cdot (\bar{F} + \bar{A}D \leftrightarrow C), \quad (2.11)$$

erfassten Regelmässigkeiten bereits durch den nur aus $(AB + C \leftrightarrow D) \cdot (\bar{C} + \bar{A}E \leftrightarrow F)$ bestehenden echten Teil von (2.11) erfasst wird. Das heisst, (2.10) ist genau für dieselben Belegungen wahr wie (2.11) – (2.10) und (2.11) sind logisch äquivalent. Die Konjunktion (2.10) ist der einzige echte Teil der Konjunktion (2.11), der äquivalent zu (2.11) ist. Letzteres gilt weder für die aus (2.7) und (2.9) noch für die aus (2.8) und (2.9) bestehende Konjunktion. Auch die aus nur einem der drei RDN-Bikonditionalen bestehenden echten Teile von (2.11) sind nicht äquivalent zu (2.11). Für (2.10) seinerseits gilt folglich, dass daraus kein RDN-Bikonditional entfernt werden kann, ohne dass der resultierende Ausdruck seine Äquivalenz zu (2.10) und damit auch zu (2.11) verliert.

Für die Konjunktion (2.10) gilt, dass sie das Instantiierungsverhalten aus $\mathcal{K}_{2.9}$ als Ganzes wiedergibt und dies nicht mehr tut, sobald (mindestens) ein RDN-Bikonditional aus (2.10) entfernt wird. Dabei gibt eine Konjunktion Ψ von RDN-Bikonditionalen das in einer Konfigurationsta-

belle $\mathcal{K}_F$ abgebildete Instantiierungsverhalten der Faktoren aus $\mathbf{F}$ genau dann als Ganzes wieder, wenn Ψ logisch äquivalent zur Konjunktion aller RDN-Bikonditionale ist, die von $\mathcal{K}_F$ impliziert werden. Für einen derartigen Ausdruck Ψ gilt, dass er für alle Belegungen wahr ist, die den Konfigurationen aus $\mathcal{K}_F$ entsprechen und Ψ allen Regularitäten Rechnung trägt, die sich in $\mathcal{K}_F$ manifestieren. Anhand des Beispiels des Stromkreises veranschaulicht, gilt es Letzteres deshalb zu berücksichtigen, weil etwa auch der folgende, das RDN-Bikonditional von C beinhaltende Ausdruck für alle Belegungen wahr ist, die einer Konfiguration aus $\mathcal{K}_{2,9}$ entsprechen:

$$(\bar{C} + \bar{A}E \leftrightarrow F) \cdot (\bar{F} + \bar{A}D \leftrightarrow C) \quad (2.12)$$

Für (2.12) gilt darüber hinaus, dass es keinen echten Teil davon gibt, der äquivalent zu (2.12) ist. In diesem Ausdruck werden jedoch nicht alle Regularitäten aus dem ebenfalls wahren RDN-Bikonditional (2.7) berücksichtigt. Das RDN-Bikonditional (2.7) hält unter anderem jenes aus den Konfigurationen ableitbare Instantiierungsverhalten der Faktoren fest, gemäss dem immer wenn C gegeben, auch D gegeben ist. Diese Regelmässigkeit wird von (2.12) nicht eingefangen, da der Ausdruck auch dann wahr sein kann, wenn C gegeben, D aber nicht gegeben ist. Berücksichtigt man ebenfalls alle Regularitäten aus (2.7), indem man es als Konjunkt in (2.12) aufnimmt, woraus sich die Konjunktion (2.11) ergibt, gibt es mit (2.10) nur genau eine Konjunktion von RDN-Bikonditionalen, die (a) äquivalent zu (2.11) ist und damit alle aus $\mathcal{K}_{2,9}$ gefolgerten Regelmässigkeiten erfasst sowie (b) kein RDN-Bikonditional aus der Konjunktion entfernt werden kann, ohne dass (a) verletzt wird. Von der Konjunktion (2.10) wird gesagt, dass sie *strukturell minimal* ist:

Definition 2.6 (Strukturelle Minimalität)

Sei $\mathcal{K}_F$ eine Konfigurationstabelle über eine Faktormenge $\mathbf{F}$. Seien weiter $\beta_1, \dots, \beta_n$, $n \geq 1$, alle RDN-Bikonditionale, die von $\mathcal{K}_F$ impliziert werden und $\Phi = \beta_1 \cdots \beta_n$ die konjunktive Verknüpfung all dieser RDN-Bikonditionale.

Eine Konjunktion $\Psi = \beta_1 \cdots \beta_m$, $1 \leq m \leq n$, ist genau dann strukturell minimal relativ zu $\mathcal{K}_F$, wenn

- (a) Ψ logisch äquivalent zu Φ ist und
- (b) es keine zu Ψ logisch äquivalente Konjunktion Ψ' von RDN-Bikonditionalen gibt, die aus Ψ resultiert, wenn mindestens ein RDN-Bikonditional aus letzterer entfernt wird (d. h., Ψ' ist ein echter Teil von Ψ).

Das RDN-Bikonditional (2.9) genügt (RF) nicht, weil es in keiner strukturell minimalen Konjunktion von RDN-Bikonditionalen enthalten ist, die das Instantiierungsverhalten aus $\mathcal{K}_{2,9}$ als

Ganzes einfängt. (2.9) ist damit nicht unverzichtbar für die Beschreibung des Instantiierungsverhaltens der Faktoren aus $\mathbf{F}$. Von (2.9) wird gesagt, dass es *strukturell redundant* ist:

Definition 2.7 (Strukturelle Redundanz)

Sei $\mathcal{K}_{\mathbf{F}}$ eine Konfigurationstabelle über eine Faktormenge $\mathbf{F}$. Ein von $\mathcal{K}_{\mathbf{F}}$ impliziertes RDN-Bikonditional β ist genau dann strukturell redundant relativ zu $\mathcal{K}_{\mathbf{F}}$, wenn es keine relativ zu $\mathcal{K}_{\mathbf{F}}$ strukturell minimale Konjunktion von RDN-Bikonditionalen gibt, die β enthält.

Kausale Abhängigkeiten zwischen n Faktoren können nur dann vorliegen, wenn es weniger als die 2^n logisch möglichen Konfigurationen dieser Faktoren gibt. Dabei leistet jedes RDN-Bikonditional, das Teil jener Regelmässigkeit ist, die das Verhalten dieser Faktoren beschreibt, einen wesentlichen Beitrag zur Erklärung, weshalb eben genau diese Konfigurationen und nicht andere oder weitere vorliegen. Ist ein RDN-Bikonditional strukturell redundant, weist es diese Eigenschaft nicht auf und es gibt entsprechend auch keine Rechtfertigung dafür zu sagen, dass man die in diesem Bikonditional abgebildeten Regelmässigkeiten als Kausalzusammenhänge interpretieren muss.

2.6. Minimale Theorie

Sei $\mathcal{K}_{\mathbf{F}}$ eine Konfigurationstabelle über eine Faktormenge $\mathbf{F}$. Damit die Teile eines aus $\mathcal{K}_{\mathbf{F}}$ gefolgerten RDN-Bikonditionals kausal interpretiert werden können, muss das RDN-Bikonditional selbst Teil einer strukturell minimalen Konjunktion Ψ von RDN-Bikonditionalen sein, die das gesamte in $\mathcal{K}_{\mathbf{F}}$ abgebildete Instantiierungsverhalten der Faktoren aus $\mathbf{F}$ einfängt. Es muss also gelten, dass (a) Ψ äquivalent zu $\Phi = \beta_1 \cdots \beta_n$ ist, die ihrerseits die Konjunktion aller aus $\mathcal{K}_{\mathbf{F}}$ gefolgerten RDN-Bikonditionale $\beta_1, \dots, \beta_n$, $n \geq 1$, darstellt, und dass (b) es keinen echten Teil von Ψ gibt, der logisch äquivalent zu Ψ ist.

Neben diesen Bedingungen (a) und (b) muss für eine Konjunktion Ψ einer Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ über $\mathbf{F}$ ferner gelten, dass (c) keine zwei RDN-Bikonditionale von Literalen desselben Faktors in Ψ enthalten sind. Diese Bedingung ist zu fordern, weil sich eine entsprechende, strukturell minimale Konjunktion nicht als eine Kausalstruktur interpretieren lässt. Würde man eine strukturell minimale Konjunktion kausal interpretieren, die mindestens zwei RDN-Bikonditionale vom gleichen Literal X enthält, hätte man zwei sich voneinander unterscheidende kausale Erklärungen für das Instantiierungsverhalten des Faktors X . Würden die RDN-Bikonditionale das Verhalten von X nicht unterschiedlich beschreiben, wären sie logisch äquivalent und die entsprechende Konjunktion, die diese RDN-Bikonditionale von X enthält, nicht strukturell minimal. Eine Kausalstruktur kann jedoch nicht mehrere nicht äquivalente kausale Erklärungen für

$\mathcal{K}_{2,11}$	A	B	C	D	E
c_1	1	1	1	0	1
c_2	1	1	0	1	1
c_3	1	0	1	1	1
c_4	1	0	0	1	1
c_5	0	1	1	1	1
c_6	0	0	0	1	0

Tabelle 2.11.: Konfigurationstabelle $\mathcal{K}_{2,11}$ über $\mathbf{E} = \{A, B, C, D, E\}$.

das Instantiierungsverhalten desselben Faktors X aufweisen. Diese Erklärungen können nicht zu derselben Struktur, sondern höchstens zu sich voneinander unterscheidenden Strukturen gehören. Dies soll anhand der Konfigurationstabelle $\mathcal{K}_{2,11}$ über die Faktormenge $\mathbf{E} = \{A, B, C, D, E\}$ illustriert werden.

Die sich aus $\mathcal{K}_{2,11}$ ergebenden Konjunktionen von RDN-Bikonditionalen, die den obigen Bedingungen (a) und (b) genügen, sind die folgenden:

$$(\bar{A} + \bar{B} + \bar{C} \leftrightarrow D) \cdot (A + B \leftrightarrow E) \cdot (A + C \leftrightarrow E) \quad (2.13)$$

$$(\bar{D} + \bar{B}C + \bar{C}E \leftrightarrow A) \cdot (\bar{A} + \bar{B} + \bar{C} \leftrightarrow D) \cdot (A + B \leftrightarrow E) \quad (2.14)$$

$$(\bar{D} + \bar{B}\bar{C} + \bar{B}E \leftrightarrow A) \cdot (\bar{A} + \bar{B} + \bar{C} \leftrightarrow D) \cdot (A + C \leftrightarrow E) \quad (2.15)$$

Eine aus einer Konfigurationstabelle gefolgerte strukturell minimale Konjunktion der Form (2.13), die (mindestens) zwei RDN-Bikonditionale desselben Literals enthält, lässt sich nicht als eine Kausalstruktur interpretieren. Würde man (2.13) kausal interpretieren, folgt daraus eine vermeintliche Kausalstruktur mit drei Ursachen A , B und C der Wirkung E . Interpretiert man diese drei Ursachen weiter als drei Alternativursachen von E , ist entweder B oder C redundant, da bereits $A + C$ beziehungsweise $A + B$ eine minimal notwendige Bedingung von E bilden. Konkret lassen sich keine Konfigurationen in $\mathcal{K}_{2,11}$ finden, bei denen eine Variation von E auf eine Variation von B zurückgeführt werden kann, während A und C konstant abwesend sind. Dasselbe gilt für Variationen von C und E mit konstant abwesendem A und B . Auch eine allfällige Interpretation als Kette, sodass B ausschliesslich über C indirekt kausal relevant für E oder C ausschliesslich über B indirekt kausal relevant für E ist, würde vor dem Hintergrund der geforderten Difference-Making-Paare eine den Daten widersprechende Struktur zur Folge haben. Eine derartige Verkettung würde implizieren, dass B kausal relevant für C beziehungsweise C kausal relevant für B ist, was sich nur dann rechtfertigen lässt, wenn es ein Difference-Making-Paar von B bezüglich C beziehungsweise von C bezüglich B gibt. Derartige Paare existieren jedoch nicht, da es weder RDN-Bikonditionale von C gibt, die B enthalten, noch RDN-Bikonditionale von B , die C enthalten. Es bleibt die Möglichkeit jener kausalen Interpretation, wonach B oder C aus-

schliesslich indirekt über A kausal relevant für E sind. Doch auch dies würde zu einer Struktur führen, die vor dem Hintergrund von $\mathcal{K}_{2.11}$ der grundlegenden Difference-Making-Anforderung nicht genügen kann. Möchte man behaupten, dass B (bzw. C) über A indirekt kausal relevant für E ist, lassen sich keine zwei Konfigurationen generieren, bei denen B (bzw. C) und E variieren, während A konstant abwesend ist. Ein solches durch $A + B \leftrightarrow E$ (bzw. $A + C \leftrightarrow E$) gefordertes Difference-Making-Paar von B (bzw. C) bezüglich E kann es in diesem Fall deshalb nicht geben, weil aufgrund der angenommenen Verkettung eine Variation von E , die man auf eine Variation von B zurückführen möchte, zwangsläufig eine Variation des Zwischenglieds A zur Folge hätte.

Eine Kausalstruktur kann also nicht gleichzeitig beiden Regularitäten $A + B \leftrightarrow E$ und $A + C \leftrightarrow E$ aus derselben strukturell minimalen Konjunktion genügen. Die beiden RDN-Bikonditionale von E können höchstens zwei unterschiedlichen Strukturen zugeordnet werden. Mit (2.14) und (2.15) gibt es zwei entsprechende Konjunktionen, die kausal interpretierbare Strukturen repräsentieren. Letztere werden als *Minimale Theorien* von $\mathcal{K}_{2.11}$ bezeichnet. Bevor im nächsten Abschnitt 2.7 konkreter auf diese beiden Minimalen Theorien von $\mathcal{K}_{2.11}$ und die damit einhergehende kausale Mehrdeutigkeit von $\mathcal{K}_{2.11}$ eingegangen wird, wird der vorliegende Abschnitt unter anderem mit einer entsprechenden Definition der Minimalen Theorie abgeschlossen. Eine Minimale Theorie einer Konfigurationstabelle $\mathcal{K}_F$ ist eine relativ zu $\mathcal{K}_F$ strukturell minimale Konjunktion von RDN-Bikonditionalen, die keine zwei RDN-Bikonditionale von Literalen desselben Faktors enthält:

Definition 2.8 (Minimale Theorie)

Sei $\mathcal{K}_F$ eine Konfigurationstabelle über eine Faktormenge F . Seien weiter $\beta_1, \dots, \beta_n$, $n \geq 1$, alle RDN-Bikonditionale, die von $\mathcal{K}_F$ impliziert werden und $\Phi = \beta_1 \cdots \beta_n$ die konjunktive Verknüpfung all dieser RDN-Bikonditionale.

Eine Minimale Theorie von $\mathcal{K}_F$ über F ist eine Konjunktion $\Psi = \beta_1 \cdots \beta_m$, $1 \leq m \leq n$, von RDN-Bikonditionalen, wobei

- (a) Ψ äquivalent zu Φ ist,
- (b) es keinen echten Teil von Ψ gibt, der äquivalent zu Ψ ist und
- (c) keine zwei RDN-Bikonditionale β_i und β_j von Literalen desselben Faktors $W \in F$ in Ψ enthalten sind.

Basierend auf der Definition 2.8 lässt sich schliesslich der Begriff der einfachen Minimale Theorie definieren, der später für den Begriff der kausalen Relevanz eingesetzt wird. Jedes RDN-Bikonditional einer Minimalen Theorie ist eine einfache Minimale Theorie:

Definition 2.9 (Einfache Minimale Theorie)

Ein RDN-Bikonditional β von W , das in einer Minimalen Theorie von $\mathcal{K}_F$ über F enthalten ist, ist eine einfache Minimale Theorie von W relativ zu $\mathcal{K}_F$ über F .

X ist genau dann Teil einer einfachen Minimalen Theorie von W relativ zu einer Konfigurationstabelle $\mathcal{K}_F$ über F , wenn X als Konjunkt in mindestens einer minimal hinreichenden Bedingung einer RDN-Bedingung von W vorkommt und die RDN-Bedingung zusammen mit W ein RDN-Bikonditional bildet, das Teil einer Minimalen Theorie von $\mathcal{K}_F$ ist. Damit ist garantiert, dass X relativ zum gesamten Instantiierungsverhalten aller Faktoren aus $\mathcal{K}_F$ in mindestens einem Fall den Unterschied im Hinblick auf das Vorhandensein von W ausmacht – kurz: X ist ein Difference-Maker von W . Allgemein wird ein Difference-Maker folgendermassen definiert, wobei die Definition in dem Sinne flexibel formuliert ist, dass man sowohl bei einem Konjunktionsterm ϵ , der als solcher in mindestens einer minimal hinreichenden Bedingung einer einfachen Minimalen Theorie von W vorkommt, als auch bei jedem echten Teil von ϵ bis hin zu den einzelnen Literalen jeweils von einem Difference-Maker von W sprechen kann:

Definition 2.10 (Difference-Maker)

Sei $\phi_W \leftrightarrow W$ eine einfache Minimale Theorie von W relativ zu $\mathcal{K}_F$ über F . Jeder Konjunktionsterm ϵ von Literalen $Z_1, \dots, Z_n$ mit $n \geq 1$, der als solcher Teil der einfachen Minimalen Theorie von W ist, ist ein Difference-Maker von W relativ zu $\mathcal{K}_F$.

Gemäss Definition 2.10 ist ein Konjunktionsterm ϵ von Literalen $Z_1, \dots, Z_n$ mit $n \geq 1$ genau dann ein Difference-Maker von W relativ zu $\mathcal{K}_F$, wenn es von jedem in ϵ enthaltenen Literal ein Difference-Making-Paar bezüglich W relativ zu $\phi_W \leftrightarrow W$ gibt, wobei $\phi_W \leftrightarrow W$ nicht strukturell redundant ist.

Genügt ein Bikonditional der Definition 2.9, genügt es aus rein funktionaler Sicht auch allen Anforderungen, die von Mackie (1980) oder von Graßhoff & May (2001) gefordert werden. So ist das Bikonditional aus (2.4), das die Kausalstruktur aus Abbildung 2.5 einfängt und mit dem die INUS-Bedingung motiviert wurde, ebenfalls eine einfache Minimale Theorie (vgl. S. 26 und S. 30). Auch das zur Common-Cause-Struktur aus Abbildung 2.6 gehörige Bikonditional (2.6), mit dem die zusätzlich zu fordernde Redundanzfreiheit der notwendigen Bedingung veranschaulicht wurde, ist als Teil der folgenden (eindeutigen) Minimalen Theorie von $\mathcal{K}_{2.6}$ eine einfache Minimale Theorie (vgl. S. 31 f.):

$$(A + E \leftrightarrow B) \cdot (A + D \leftrightarrow C) \quad (2.16)$$

2.7. Ambiguitäten

Die Konfigurationstabelle $\mathcal{K}_{2.11}$ über $\mathbf{E} = \{A, B, C, D, E\}$ impliziert mit (2.14) und (2.15) zwei Minimale Theorien. Kausal interpretiert sind die Ursachen von E gemäss (2.14) A und B , während (2.15) nahelegt, dass E durch A oder C verursacht wird. Wie oben bereits im Zusammenhang mit der Bedingung (c) aus Definition 2.8 festgestellt, lassen sich diese beiden Kausalaussagen jedoch nicht derart vereinen, dass für die dahinterliegende Kausalstruktur beides gleichzeitig zutrifft. Im vorliegenden Fall gilt also, dass die kausale Interpretation entweder von (2.14) oder von (2.15) die den Konfigurationen aus $\mathcal{K}_{2.11}$ zugrundeliegende Kausalstruktur einfängt. Zumal die beiden Ausdrücke logisch äquivalent sind, lässt sich alleine basierend auf $\mathcal{K}_{2.11}$ aber nicht entscheiden, welche der beiden Strukturen es tatsächlich ist – es ergibt sich eine *Ambiguität*.

Ambiguitäten sind ein gängiges Phänomen des kausalen Schliessens, das nicht nur an die hier verwendete Kausaltheorie und den damit einhergehenden methodologischen Ansatz gebunden ist (vgl. u. a. Simon (1954); Spirtes et al. (2000), Kapitel 4; Eberhardt (2014)). Nicht selten ist man mit der Schwierigkeit konfrontiert, dass aus den untersuchten Daten mehrere alternative Kausalhypothesen hervorgehen, die untereinander nicht verträglich, für sich genommen aber alle genau gleich gut mit den Daten vereinbar sind. Überwunden wird eine derartige Unterterminiertheit typischerweise dadurch, dass man nach weiterer empirischer Evidenz sucht, die im Idealfall eine Kausalhypothese stützt, während sie die übrigen widerlegt. Indem beim konfigurationsalen kausalen Modellieren jede Kausalanalyse relativ zu einer begrenzten Menge von Faktoren durchgeführt wird, erfolgt eine derartige Verbesserung der Datenbasis unter anderem durch Erweiterungen dieser Faktormenge (vgl. S. 24). So kann sich im obigen Beispiel von $\mathcal{K}_{2.11}$ etwa mit der Hinzunahme eines relativ zu $\mathbf{E}$ konstant anwesenden Faktors F herausstellen, dass es noch weitere empirisch mögliche Konfigurationen gibt und $\mathcal{K}_{2.11}$ nur ein Teil der Konfigurationstabelle $\mathcal{K}_{2.12}$ über $\mathbf{G} = \mathbf{E} \cup \{F\} = \{A, B, C, D, E, F\}$ ist (vgl. c_1 bis c_9 sowie c_{12} und c_{13} aus $\mathcal{K}_{2.12}$).

Im Gegensatz zu $\mathcal{K}_{2.11}$ ergeben sich aus $\mathcal{K}_{2.12}$ keine Mehrdeutigkeiten. Die von $\mathcal{K}_{2.11}$ noch gestützte Kausalhypothese, wonach B eine Ursache von E ist, wird mit der Hinzunahme von F zugunsten der alternativen Kausalhypothese, wonach neben A auch C eine Ursache von E darstellt, verworfen. Dieser Umstand zeigt sich in der eindeutigen Minimalen Theorie von $\mathcal{K}_{2.12}$:

$$(\bar{D} + B\bar{C} + \bar{B}E \leftrightarrow A) \cdot (\bar{A} + \bar{B} + \bar{C} \leftrightarrow D) \cdot (A + CF \leftrightarrow E) \quad (2.17)$$

Aus Minimalen Theorien $\Phi_1, \Phi_2, \dots, \Phi_n$ einer Konfigurationstabelle $\mathcal{K}_F$ resultierende Kausalhypothesen sind so zu lesen, dass Φ_1 oder Φ_2 oder ... oder Φ_n die zwischen den Faktoren aus $\mathbf{F}$ bestehenden Kausalzusammenhänge einfangen. Wie sich später im Zusammenhang mit den

$\mathcal{K}_{2.12}$	A	B	C	D	E	F
c_1	1	1	1	0	1	1
c_2	1	1	1	0	1	0
c_3	1	1	0	1	1	1
c_4	1	1	0	1	1	0
c_5	1	0	1	1	1	1
c_6	1	0	1	1	1	0
c_7	1	0	0	1	1	1
c_8	1	0	0	1	1	0
c_9	0	1	1	1	1	1
c_{10}	0	1	1	1	0	0
c_{11}	0	0	1	1	0	0
c_{12}	0	0	0	1	0	1
c_{13}	0	0	0	1	0	0

Tabelle 2.12.: Konfigurationstabelle $\mathcal{K}_{2.12}$ über $\mathbf{G} = \{A, B, C, D, E, F\}$, die $\mathcal{K}_{2.11}$ als Teil enthält (vgl. c_1 bis c_9 sowie c_{12} und c_{13}).

Kausalketten noch zeigen wird, handelt es sich dabei im Allgemeinen aber nicht um ein exklusives Oder. Mehrere Minimale Theorien lassen sich teils durchaus so interpretieren, dass die von ihnen eingefangenen Zusammenhänge in ein und derselben Kausalstruktur integriert sind (vgl. Abschnitt 2.11). Wie anhand von $\mathcal{K}_{2.11}$ gezeigt, ist dies jedoch nicht immer der Fall, weshalb sich die Bedingung, dass ein Literal X Teil einer einfachen Minimalen Theorie von W ist, zwar als notwendig, nicht aber als hinreichend dafür herausstellt, A als Ursache von W zu interpretieren. Indem auch (2.14) eine Minimale Theorie ist, wäre man allein mit der Forderung dieser Bedingung gezwungen, B und C als Ursachen von E zu interpretieren, obwohl sich mit der Verbesserung der Datenbasis herausstellen kann, dass nur C kausal relevant für E ist. Der nachfolgende Abschnitt erklärt, was neben der Mitgliedschaft in einer Minimalen Theorie zusätzlich gelten muss, um eine hinreichende Bedingung für kausale Relevanz zu erhalten.

2.8. Vollständigkeits- und Erweiterungsklausel

Durch die auf mehreren Ebenen geforderte Minimalität ist eine Minimale Theorie frei von Redundanzen und jedes Literal, das Teil einer darin enthaltenen einfachen Minimalen Theorie ist, ein Difference-Maker. Diese Eigenschaft wird jeweils relativ zu einer spezifischen Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ über eine Faktormenge $\mathbf{F}$ bestimmt. Eine solche Konfigurationstabelle kann wie im Fall von $\mathcal{K}_{2.11}$ betreffend die daraus ableitbaren Kausalaussagen mehrdeutig sein. Wie sich

gleich zeigen wird, kann eine Konfigurationstabelle auch lückenhaft sein, woraus sich ebenfalls Kausalhypothesen ergeben können, die den tatsächlichen Kausalzusammenhängen widersprechen. Um den Verursachungsbegriff von der Relativierung auf Konfigurationstabellen zu lösen und die von Konfigurationen implizierten Minimalen Theorien kausal zu interpretieren, sind zwei weitere Bedingungen notwendig.

2.8.1. Vollständigkeitsklausel

Für einen adäquaten, auf Minimalen Theorien von Konfigurationstabellen basierenden Begriff der kausalen Relevanz muss zusätzlich gelten, dass die Konfigurationstabellen jeweils alle empirisch möglichen Konfigurationen der jeweiligen Faktoren enthalten. In der Hume'schen Metaphysik verortet, ist eine Konfiguration genau dann empirisch möglich, wenn sie mindestens einmal im Laufe der Existenz des analysierten Systems (in einem atemporalen Sinn) auftritt. Um welche Konfigurationen es sich dabei genau handelt, ist eine nackte Tatsache (brute fact), die nicht weiter erklärt werden kann.

Die Forderung, dass relativ zu den jeweiligen Faktormengen alle (und nur die) empirisch möglichen Konfigurationen vorliegen, wird als *Vollständigkeitsklausel* bezeichnet. Weshalb diese Klausel gelten muss, soll anhand eines auf Tabelle 2.13 basierenden Szenarios illustriert werden. Dabei soll davon ausgegangen werden, dass die drei Faktoren der Menge $\mathbf{H} = \{X, Y, Z\}$ tatsächlich in keiner Weise kausal miteinander verknüpft sind und es insgesamt die acht empirisch möglichen Konfigurationen aus $\mathcal{K}_{2.13a}$ gibt. Es soll weiter davon ausgegangen werden, dass für Kausalanalysen, aus welchem Grund auch immer, jeweils nur die vier Konfigurationen aus $\mathcal{K}_{2.13b}$ vorliegen. $\mathcal{K}_{2.13b}$ stellt eine echte Teilmenge der empirisch möglichen Konfigurationen aus $\mathcal{K}_{2.13a}$ über $\mathbf{H}$ dar.

Relativ zu den Konfigurationen aus $\mathcal{K}_{2.13b}$ resultiert die einfache Minimale Theorie $X + Y \leftrightarrow Z$. Würde man also für die kausale Interpretation von Minimalen Theorien nicht zusätzlich die Vollständigkeitsklausel fordern, wäre der Ausschluss eines derartigen Szenarios nicht garantiert, bei dem man X und Y als Ursachen von Z betrachten würde, obwohl keine Kausalzusammenhänge existieren. Gilt jedoch die Vollständigkeitsklausel, womit relativ zu $\mathbf{H} = \{X, Y, Z\}$ die Minimalen Theorien von $\mathcal{K}_{2.13a}$ bestimmt werden, lassen sich relativ zu $\mathbf{H}$ auch keine RDN-Bikonditionale ableiten, weshalb basierend auf dem Begriff der Minimalen Theorie korrekterweise auch keine Kausalzusammenhänge gefolgert werden.

Im weiteren Verlauf dieses Kapitels wird die Gültigkeit der Vollständigkeitsklausel vorausgesetzt, weshalb anstelle von einer einfachen Minimalen Theorie „relativ zu $\mathcal{K}_{\mathbf{F}}$ über $\mathbf{F}$ “ oft kurz

$\mathcal{K}_{2.13a}$	X	Y	Z
c_1	1	1	1
c_2	1	1	0
c_3	1	0	1
c_4	1	0	0
c_5	0	1	1
c_6	0	1	0
c_7	0	0	1
c_8	0	0	0

(a)

$\mathcal{K}_{2.13b}$	X	Y	Z
c_1	1	1	1
c_2	1	0	1
c_3	0	1	1
c_4	0	0	0

(b)

Tabelle 2.13.: Konfigurationstabellen $\mathcal{K}_{2.13a}$ und $\mathcal{K}_{2.13b}$ über $\mathbf{H} = \{X, Y, Z\}$, wobei $\mathcal{K}_{2.13a}$ alle empirisch möglichen Konfigurationen enthält.

von einer einfachen Minimalen Theorie „relativ zu $\mathbf{F}$ “ gesprochen wird. Wenn nicht explizit anders erwähnt, ist damit jeweils die Rede von einer einfachen Minimalen Theorie relativ zu $\mathcal{K}_{\mathbf{F}}$ über $\mathbf{F}$, wobei $\mathcal{K}_{\mathbf{F}}$ alle empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}$ enthält.

2.8.2. Erweiterungsklausel

Neben der Vollständigkeitsklausel ist eine letzte Bedingung notwendig, damit sich ein auf Minimalen Theorien basierender Verursachungsbegriff ergibt, der adäquat ist: die *Erweiterungsklausel* (vgl. u. a. Baumgartner (2008b), S. 340 ff.). Um zu zeigen, was diese Klausel besagt und weshalb ihr notwendigerweise zu genügen ist, damit Minimale Theorien kausal interpretiert werden können, soll nochmals die Kausalstruktur aus Abbildung 2.9 herangezogen werden (vgl. S. 40). Anstelle aller empirisch möglichen Konfigurationen über die Faktormenge $\mathbf{D} = \{A, B, C, D, E, F\}$ werden vor dem Hintergrund derselben Kausalstruktur jedoch alle empirisch möglichen Konfigurationen relativ zur Faktormenge $\mathbf{I} = \{A, B, C, D, F\}$ betrachtet. Daraus ergibt sich die Konfigurationstabelle $\mathcal{K}_{2.14}$ über $\mathbf{I} = \{A, B, C, D, F\}$ aus Tabelle 2.14.²⁰

Aus $\mathcal{K}_{2.14}$ über $\mathbf{I} = \{A, B, C, D, F\}$ ergeben sich die beiden RDN-Bikonditionale $AB + C \leftrightarrow D$ und $\bar{F} + \bar{A}D \leftrightarrow C$, die konjunktiv verknüpft die folgende Minimale Theorie von $\mathcal{K}_{2.14}$ bilden:

$$(AB + C \leftrightarrow D) \cdot (\bar{F} + \bar{A}D \leftrightarrow C) \quad (2.18)$$

²⁰Wird der Faktor E nicht berücksichtigt, resultiert eine Konfigurationstabelle mit lediglich zehn anstelle von 16 Konfigurationen. Der Grund dafür liegt darin, dass durch das Ausblenden von E einige ohne E nicht unterscheidbare Konfigurationen resultieren (vgl. bspw. c_1 und c_2 aus $\mathcal{K}_{2.9}$).

$\mathcal{K}_{2,14}$	A	B	C	D	F
c_1	1	1	1	1	0
c_2	1	1	0	1	1
c_3	1	0	1	1	0
c_4	1	0	0	0	1
c_5	0	1	1	1	1
c_6	0	1	1	1	0
c_7	0	1	0	0	1
c_8	0	0	1	1	1
c_9	0	0	1	1	0
c_{10}	0	0	0	0	1

Tabelle 2.14.: Auflistung aller empirisch möglichen Konfigurationen der Kausalstruktur aus Abbildung 2.9 relativ zur Faktormenge $\mathbf{I} = \{A, B, C, D, F\}$.

$\bar{F} + \bar{A}D \leftrightarrow C$ ist relativ zu $\mathcal{K}_{2,14}$ eine einfache Minimale Theorie. Kausal interpretiert ergibt sich damit unter anderem die Kausalhypothese, wonach D kausal relevant für C ist. Diese relativ zu $\mathbf{I} = \{A, B, C, D, F\}$ formulierte Kausalhypothese ist vor dem Hintergrund der tatsächlichen Kausalstruktur aus Abbildung 2.9 bekanntlich falsch. Wird die Faktormenge $\mathbf{I}$ um den Faktor E zu $\mathbf{D} = \mathbf{I} \cup \{E\} = \{A, B, C, D, E, F\}$ erweitert, resultieren die in (2.10) enthaltenen einfachen Minimalen Theorien $AB + C \leftrightarrow D$ sowie $\bar{C} + \bar{A}E \leftrightarrow F$. Das RDN-Bikonditional $\bar{F} + \bar{A}D \leftrightarrow C$ stellt sich als strukturell redundant heraus, womit relativ zu $\mathbf{D}$ insgesamt die tatsächlichen Kausalzusammenhänge korrekt eingefangen werden.

Die Erweiterung zu $\mathbf{D}$ der Faktormenge $\mathbf{I}$ führt dazu, dass das unter $\mathbf{I}$ noch als einfache Minimale Theorie betrachtete RDN-Bikonditional $\bar{F} + \bar{A}D \leftrightarrow C$ relativ zu $\mathbf{D}$ redundant wird und folglich keine einfache Minimale Theorie mehr darstellt. Im Gegensatz dazu bleibt $\bar{C} + \bar{A}E \leftrightarrow F$ in Anbetracht der tatsächlichen Kausalstruktur korrekterweise sowohl unter der Faktormenge $\mathbf{I}$ als auch unter $\mathbf{D}$ als einfache Minimale Theorie bestehen. Die kausale Relevanz von $\bar{C}$, $\bar{A}$ und E für F ist nicht an eine spezifische Faktormenge gebunden und die Literale sind unabhängig davon, welche weiteren Faktoren betrachtet werden, jeweils Difference-Maker.²¹ Ein Literal X eines Faktors aus einer Menge $\mathbf{F}$, das Teil einer einfachen Minimalen Theorie eines Literals W eines

²¹Eine Faktormenge kann sich im Übrigen neben Erweiterungen auch durch verschiedene Spezifizierungsgrade ändern. Gleiche kausale Abhängigkeiten können je nach Kontext mittels verschiedenen Spezifikationen von Faktoren beschrieben werden: Ein Faktor wie „stabile Demokratie“ (D) kann durch eine detaillierte Konjunktion von Kriterien beschrieben werden, die eine stabile Demokratie ausmachen. In einem solchen Fall wird ein Literal des Faktors D einer Faktormenge $\mathbf{F}$ aber nicht Opfer der Redundanz, sondern wird ersetzt und eine Analyse relativ zu einer anderen Faktormenge $\mathbf{G} \setminus \{D\}$ durchgeführt (vgl. Hofmann & Baumgartner (2011) betreffend die möglichen Konsequenzen bei der Verwendung verschiedener Spezifizierungsgrade eines Faktors).

anderen Faktors aus $\mathbf{F}$ ist, kann nur dann als kausal relevant für W interpretiert werden, wenn X auch bei jeder Erweiterung $\mathbf{F}' \supset \mathbf{F}$ Teil einer einfachen Minimalen Theorie von W bleibt. Diese Bedingung wird als *Erweiterungsklausel* bezeichnet. Fordert man diese Klausel nicht, ist nicht ausgeschlossen, dass jeweils ein wie das oben anhand von $\mathcal{K}_{2.14}$ beschriebenes Szenario vorliegt. Analog wie bei $\bar{F} + \bar{A}D \leftrightarrow C$ könnten sich die relativ zu $\mathbf{F}$ gefolgerten Regelmässigkeiten zwischen X und W , die X als Teil einer einfachen Minimalen Theorie von W auszeichnen, bei Erweiterungen auflösen und die relativ zu $\mathbf{F}$ gestützte Kausalhypothese, wonach X kausal relevant für W ist, sich mit Einbezug weiterer Faktoren als falsch herausstellen.

Bezüglich der Erweiterungen gilt es abschliessend zu bemerken, dass man bei der Hinzunahme weiterer Faktoren zu einer Menge die modale Unabhängigkeit der von den Faktoren repräsentierten (natürlichen) Eigenschaften berücksichtigen muss (vgl. S. 22). Wie bei jeder Faktormenge $\mathbf{F}$ muss ebenfalls bei einer Erweiterung von $\mathbf{F}$ zu $\mathbf{F}'$ sichergestellt werden, dass mit den hinzukommenden Faktoren keine modalen Abhängigkeiten eingeführt werden. Aus diesem Grund wird in der vorliegenden Arbeit künftig jeweils von *geeigneten* Erweiterungen gesprochen.

2.9. Verursachungsbegriff

Damit X kausal relevant für W ist, muss X nicht nur relativ zu einer Faktormenge $\mathbf{F}$ Teil einer einfachen Minimalen Theorie von W sein, sondern relativ zu jeder geeigneten Erweiterung $\mathbf{F}'$ von $\mathbf{F}$ Teil einer einfachen Minimalen Theorie von W bleiben – kurz: X muss *permanenter* Teil einer einfachen Minimalen Theorie von W sein.

Die Erweiterungsklausel – in Kombination mit der Vollständigkeitsklausel – ersetzt das von Mackie in seiner INUS-Theorie geforderte kausale Feld (vgl. S. 29 f.). Beide Forderungen verfolgen den Zweck, die sich relativ zu einer begrenzten Faktormenge manifestierenden Regelmässigkeiten von dieser Relativierung zu lösen und zu verallgemeinern. Die Erweiterungsklausel überwindet dabei zwei Mängel, die ein Verursachungsbegriff mit Einbezug des kausalen Feldes aufweist: Einerseits bricht die Erweiterungsklausel den in gewisser Hinsicht bestehenden definitorischen Regress auf, der mit dem kausalen Feld einhergeht. Einem Verursachungsbegriff mit kausalem Feld zufolge ist X genau dann eine Ursache von Y , wenn X und Y im kausalen Feld einer bestimmten Regelmässigkeit folgen. Das kausale Feld beschreibt dabei gleichbleibende Hintergrundbedingungen, die ihrerseits genau dann gegeben sind, wenn der kausale Einfluss der nicht betrachteten Faktoren zwischen den verschiedenen Konfigurationen jeweils derselbe bleibt. Um also X als Ursache von Y auszuweisen, wird der Verursachungsbegriff über das kausale Feld bereits indirekt verwendet, indem Bedingungen an das Verhalten der nicht betrachteten kausal

relevanten Faktoren gestellt werden. Andererseits bietet die Erweiterungsklausel einen Umgang mit Ambiguitäten. Würde man anstelle der Erweiterungsklausel fordern, dass die Minimalen Theorien in einem kausalen Feld eingebettet sind, wäre man gezwungen, sowohl (2.14) als auch (2.15) kausal zu interpretieren (vgl. S. 45). Geht man davon aus, dass es neben F keine weiteren Faktoren gibt, die mit den Faktoren aus $\mathbf{E} = \{A, B, C, D, E\}$ kausal verknüpft sind, würde die dazugehörige Konfigurationstabelle $\mathcal{K}_{2.11}$ mit konstant anwesendem Faktor F den vom kausalen Feld beschriebenen gleichbleibenden Hintergrundbedingungen genügen.

2.9.1. Kausale Relevanz

Unter Berücksichtigung all der obigen Forderungen ergibt sich insgesamt folgende Definition der kausalen Relevanz, wobei der Ausdruck „kausal relevant“ bezüglich der Direkt- und Indirektheit unbestimmt und nicht identisch mit der möglicherweise suggerierten direkten kausalen Relevanz ist. Das heisst, ein im Sinne der folgenden Definition für eine Wirkung W kausal relevantes Literal X kann direkt oder indirekt kausal relevant für W sein²²:

Definition 2.11 (*Kausale Relevanz*)

X ist genau dann kausal relevant für W , wenn es eine die Faktoren X und W enthaltende Faktormenge $\mathbf{F}$ derart gibt, dass

- (a) X relativ zu allen empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}$ Teil einer einfachen Minimalen Theorie von W ist,
- (b) für jede geeignete Erweiterung $\mathbf{F}' \supset \mathbf{F}$ von $\mathbf{F}$ X relativ zu allen empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}'$ Teil einer einfachen Minimalen Theorie von W bleibt.

Von X wird abkürzend gesagt, dass es permanenter Teil einer einfachen Minimalen Theorie von W ist.

Definition 2.11 liefert eine vollständige und reduktive Definition von kausaler Relevanz, die einem angemessenen Verursachungsbegriff zur konzeptuellen Fundierung von konfigurationalen Methoden wie QCA oder CNA entspricht. Aus Konfigurationen gefolgerte Suffizienz- und Notwendigkeitsbeziehungen sind genau dann kausal interpretierbar, wenn es sich bei diesen Regelmässigkeiten um einfache Minimale Theorien handelt, deren Teile der Erweiterungsklausel

²²Die Aufspaltung der kausalen Relevanz in direkte und indirekte kausale Relevanz wird im Zusammenhang mit der Definition der Kausalkette in Abschnitt 2.11 vorgenommen.

genügen. Das wesentliche Charakteristikum der einfachen Minimalen Theorien ist dabei ihre strikte Redundanzfreiheit: Ist eine notwendige Disjunktion von hinreichenden Konjunktionen eines Literals W nicht minimal oder mit W zu einem Bikonditional verknüpft nicht Teil einer strukturell minimalen Konjunktion von Bikonditionalen, können die darin enthaltenen Literale auch nicht kausal interpretiert werden. QCA fordert diese maximale Redundanzfreiheit nicht. Nach aktueller Ansicht gehen Vertreter dieses Analyseverfahrens sogar in die entgegengesetzte Richtung und vertreten die Haltung, auf „sparsame“ Darstellungen von hinreichenden und notwendigen Bedingungen zugunsten von Mittelwegen („intermediate solutions“) zu verzichten (vgl. bspw. Schneider & Wagemann (2012), Kapitel 8). Weshalb das so ist, wird in den folgenden Kapiteln ersichtlich (vgl. u. a. Kapitel 5, S. 192). Das bisherige kausaltheoretische Fundament von CNA hingegen fordert zwar Minimalität in Form von RDN-Bikonditionalen, berücksichtigt aber die strukturellen Redundanzen nicht. Letztere führen dazu, dass man für die kausale Interpretation von RDN-Bikonditionalen, die von Konfigurationen von Faktoren einer Menge F impliziert werden, das Instantiierungsverhalten der Faktoren aus F als Ganzes berücksichtigen muss und sich komplexe Kausalstrukturen mit mehreren Wirkungen aufgrund möglicher struktureller Redundanzen nicht einfach modular durch Verknüpfungen von einzelnen RDN-Bikonditionalen aufbauen lassen.

Aufgrund der Erweiterungsklausel (b) aus Definition 2.11 ist der Schluss von einfachen Minimalen Theorien auf Kausalbeziehungen inhärent induktiv. Bei der kausalen Interpretation von einfachen Minimalen Theorien bleibt somit immer ein „induktives Risiko“, das durch Folgeanalysen mit erweiterten Faktormengen kontinuierlich sinkt (vgl. Baumgartner (2015), S. 845). Möchte man von einfachen Minimalen Theorien deduktiv auf kausale Abhängigkeiten schließen, müsste die entsprechende, alle empirisch möglichen Konfigurationen enthaltende Konfigurationstabelle $\mathcal{K}_F$ (1) die universelle Menge von allen möglichen (geeigneten) Faktoren oder mindestens (2) alle Faktoren, die kausal miteinander verknüpft sind, berücksichtigen. Variante (2) würde einen definitorischen Zirkel produzieren, indem der Begriff der kausalen Relevanz in der Definition ebendieser vorkommen würde. Mit Variante (1) würde sich ein Verursachungsbegriff ergeben, der sich nicht in ein praktikables Verfahren für das konfigurationale kausale Modellieren implementieren liesse. Angesichts des Ziels der vorliegenden Arbeit, genau letzteres zu liefern, stellt ein solcher Verursachungsbegriff keine Option dar, auch wenn er aus rein theoretischer Sicht angemessen wäre.

2.9.2. Ursachen einer Wirkung

Mit Definition 2.11 liegt der zentrale Baustein für die Beschreibung von über Konfigurationen ermittelten kausalen Abhängigkeiten vor. Darauf basierend lassen sich sukzessiv komplexere

Kausalstrukturen wiedergeben, bis schliesslich die gesamte Komplexität aus Abschnitt 2.1.2 eingefangen ist. Konkret können mit diesem Begriff der kausalen Relevanz alle kausalen Verknüpfungen aus (a) und (b) der Abbildung 2.1 beschrieben werden (vgl. S. 20). Im Folgenden werden zwei Definitionen für die Begriffe der komplexen sowie alternativen Ursachen präsentiert (vgl. Strukturtypen (c) und (d) aus Abbildung 2.1). Definition 2.11 liefert dazu zwar die fundamentale Komponente, reicht jedoch alleine für eine vollständige Beschreibung noch nicht aus. Genügen etwa drei Literale A , B und C der Definition 2.11 derart, dass A und B kausal relevant für C sind, ist mit dieser Definition noch nicht bestimmt, ob nun A und B zusammen Teil derselben komplexen Ursache von C , zweier Alternativursachen oder beides sind. Solche komplexe und alternative Ursachen eines Literals W lassen sich mittels einfacher Minimaler Theorien folgendermassen definieren:

Definition 2.12 (*Komplexe Ursache*)

Sei ϵ ein Konjunktionsterm von Literalen $Z_1 \cdots Z_n$ mit $n \geq 2$. ϵ ist genau dann eine komplexe Ursache von W , wenn es eine die Faktoren $Z_1, \dots, Z_n, W$ enthaltende Faktormenge $\mathbf{F}$ derart gibt, dass

- (a) ϵ als solcher relativ zu allen empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}$ Teil einer einfachen Minimalen Theorie von W ist und
- (b) für jede geeignete Erweiterung $\mathbf{F}' \supset \mathbf{F}$ von $\mathbf{F}$ ϵ als solcher relativ zu allen empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}'$ Teil einer einfachen Minimalen Theorie von W bleibt.

Definition 2.13 (*Alternative Ursache*)

Seien $\epsilon_1, \dots, \epsilon_n$ mit $n \geq 2$ verschiedene, aus jeweils mindestens einem Literal bestehende Konjunktionsterme. $\epsilon_1, \dots, \epsilon_n$ sind genau dann verschiedene Alternativursachen von W , wenn es eine die Faktoren aus $\epsilon_1, \dots, \epsilon_n$ und den Faktor W enthaltende Faktormenge $\mathbf{F}$ derart gibt, dass

- (a) $\epsilon_1, \dots, \epsilon_n$ relativ zu allen empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}$ Teil unterschiedlicher Disjunkte einer einfachen Minimalen Theorie von W sind und
- (b) es für jede geeignete Erweiterung $\mathbf{F}' \supset \mathbf{F}$ von $\mathbf{F}$ eine einfache Minimale Theorie von W gibt, die relativ zu $\mathbf{F}'$ (a) erfüllt.

Die Definitionen 2.11 bis 2.13 ermöglichen eine lückenlose kausale Interpretation einer einfachen Minimalen Theorie, deren Teile bei Hinzunahme weiterer geeigneter Faktoren in ihren Verknüpfungen zueinander bestehen bleiben. Damit ist gemeint, dass etwa zwei Literale A und

B , die relativ zu einer Faktormenge $\mathbf{F}$ als konjunktive Verknüpfung AB Teil einer einfachen Minimalen Theorie eines Literals C sind, auch relativ zu jeder Erweiterung der Faktormenge $\mathbf{F}$ als konjunktive Verknüpfung AB Teil einer einfachen Minimalen Theorie von C bleiben, wobei nicht ausgeschlossen wird, dass AB durch Erweiterungen der Faktormenge um zusätzliche Literale ergänzt wird. Es bedeutet auch, dass zwei Literale E und F , die relativ zu $\mathbf{F}$ als Disjunktion $E + F$ Teil einer einfachen Minimalen Theorie von C sind, auch relativ zu jeder Erweiterung disjunktiv verknüpft Teil einer einfachen Minimalen Theorie von C bleiben, wobei nicht ausgeschlossen wird, dass sich E und F zum Beispiel durch die Hinzunahme gewisser Faktoren zusätzlich auch als Teil eines Konjunktionstermes einer einfachen Minimalen Theorie von C herausstellen.

Mit Ausnahme einer unter Abschnitt 2.11 noch folgenden Unterscheidung zwischen direkter und indirekter kausaler Relevanz fangen die Definitionen 2.11, 2.12 und 2.13 die verschiedenen Typen von Ursachen einer Wirkung ein. Das soll nochmals anhand der aus den Konfigurationen aus Abbildung 2.5 resultierenden Minimalen Theorie $F\bar{S} + M \leftrightarrow A$ veranschaulicht werden. Bleiben F und $\bar{S}$ bei Hinzunahme weiterer Faktoren Teil der gleichen und M Teil einer weiteren, F und $\bar{S}$ nicht enthaltenden minimal hinreichenden Bedingung derselben einfachen Minimalen Theorie von A , folgt aus obigen Definitionen der Reihe nach, dass

- (1) F , M und $\bar{S}$ gemäss Definition 2.11 kausal relevant für A sind,
- (2) F und $\bar{S}$ zusammen gemäss Definition 2.12 eine komplexe Ursache $F\bar{S}$ von A bilden und schliesslich
- (3) $F\bar{S}$ und M gemäss Definition 2.13 zwei alternative Ursachen von A sind.

2.10. Eigenschaften des Verursachungsbegriffs

Definition 2.11 beschreibt das Kernstück des kausalthoretischen Fundaments des konfigurationsalen kausalen Modellierens, indem sie den Schluss von Booleschen Abhängigkeiten auf Kausalität reguliert. Damit soll nicht behauptet werden, dass dies der einzige valable Verursachungsbegriff ist. In anderen Kontexten, die unter anderem von den Fragestellungen oder den dazugehörigen Datengrundlagen abhängig sind, kann es durchaus andere Verursachungsbegriffe geben, die angemessen sind. Im Kontext konfiguratorischer Methoden wie QCA und CNA jedoch ist es Definition 2.11, die eine adäquate Definition von Kausalität darstellt.

2.10.1. Die kausalen Prinzipien

Der über Minimale Theorien definierte Verursachungsbegriff bringt gewisse Voraussetzungen mit sich und hat bestimmte Implikationen, die als grundlegende Eigenschaften der in dieser Arbeit betrachteten kausalen Abhängigkeiten zu verstehen sind. Einige Eigenschaften können in Form von kausalen Prinzipien dargestellt werden, die den groben Rahmen jenes Bereiches abstecken, innerhalb dessen der hier verwendete Verursachungsbegriff zur Anwendung kommt (vgl. u. a. Graßhoff & May (2001), S. 88 ff.). Hinsichtlich des eingangs dieses Kapitels angesprochenen kausalen Pluralismus können diese Prinzipien als Facetten verstanden werden, die von der Kausaltheorie berücksichtigt werden und denen die untersuchten Kausalzusammenhänge somit genügen. Das bekannteste dieser Prinzipien ist das Determinismusprinzip (vgl. u. a. Graßhoff & May (2001), S. 88):

Determinismusprinzip: Immer wenn dieselben Ursachen auftreten, treten auch dieselben Wirkungen auf.

Das Determinismusprinzip ist fundamental für eine Analyse von Kausalzusammenhängen mittels notwendiger und hinreichender Bedingungen. Liegen zwei Situationen S_1 und S_2 vor, bei denen die gleichen Ursachen instantiiert sind, werden gemäß dem Determinismusprinzip in beiden Situationen auch dieselben Wirkungen instantiiert sein. Es ist somit ausgeschlossen, dass gleiche Ursachen in Situationen vom selben Typ unterschiedliche Wirkungen nach sich ziehen. Sind die Wirkungen zweier Situationen verschieden, unterscheiden sich auch ihre Ursachen voneinander. Kausalbeziehungen, die dem Determinismusprinzip genügen, werden auch als *deterministische* Kausalbeziehungen bezeichnet.²³

Das zweite Prinzip ist das Kausalitätsprinzip (vgl. u. a. Graßhoff & May (2001), S. 88 f.). Dieses Prinzip besagt, dass jedes Ereignis eine Ursache hat. Die im Folgenden präsentierte Version des Kausalitätsprinzips weicht jedoch etwas von dieser üblicherweise verwendeten Formulierung ab. Der Grund ist, dass die in der vorliegenden Arbeit entwickelte Kausalitätstheorie primär auf Typen-Ebene angesiedelt ist, deren kausal verknüpfte Entitäten Literale sind. Bei obiger Formulierung handelt es sich aufgrund des verwendeten Begriffs des Ereignisses aber um eine Token-Version. Eine entsprechende Typen-Version des Prinzips ist die folgende Formulierung:

²³Wie bereits bemerkt, wird damit nicht behauptet, dass die Gesamtheit aller natürlichen Prozesse in der Welt, in der wir leben, deterministisch sind. Gerade das Determinismusprinzip wird u. a. im Zusammenhang mit physikalischen Theorien wie der Quantenmechanik sowie im Zusammenhang mit den aus dem Prinzip zu folgernden Konsequenzen bezüglich der Willensfreiheit kontrovers diskutiert (vgl. u. a. Wheeler & Zurek (1983); Pereboom (2009)). Es wird jedoch behauptet bzw. vorausgesetzt, dass kausale Prozesse im Sinne von Definition 2.11 deterministisch sind.

Kausalitätsprinzip: Jedes Literal hat eine Ursache.

Gemäss dem Kausalitätsprinzip sind spontan instantiierte Literale ausgeschlossen. Ist in einer Situation ein Literal instantiiert, kann darauf geschlossen werden, dass in dieser Situation auch eine Ursache dieses Literals gegeben ist.

Im Zusammenhang mit den in dieser Arbeit für eine Kausalanalyse herangezogenen Daten in Form von Konfigurationen spielen das Determinismus- und Kausalitätsprinzip aus methodologischer Sicht wichtige Rollen. Um das zu zeigen, soll von einem stark idealisierten Szenario ausgegangen werden, das aus zwei Situationen S_1 und S_2 besteht. S_1 und S_2 seien bis auf die Variation einer Wirkung B und eines weiteren Faktors A , die beide in S_1 instantiiert sind und in S_2 nicht, identisch. Auf der Suche nach den Ursachen von B lässt sich mit dem Kausalitätsprinzip folgern, dass in S_1 eine Ursache von B instantiiert sein muss. Weiter tritt B in S_2 nicht auf, womit gemäss dem Determinismusprinzip in S_1 kausal relevante Literale von B gegeben sein müssen, die in S_2 nicht gegeben sind. Da S_1 und S_2 abgesehen von B bis auf die Variation im Instantiierungsverhalten des Faktors A identisch sind, kann schliesslich gefolgert werden, dass in S_1 A Ursache von B ist. Während also unter diesen idealisierten Bedingungen mit dem Kausalitätsprinzip gefolgert werden kann, dass in S_1 mindestens eine Ursache von A instantiiert ist, kann unter anderem mit Hilfe des Determinismusprinzips festgestellt werden, welches Literal in S_1 seine kausale Relevanz offenbart. Dieser auf derart idealisierte Bedingungen basierende Schluss auf kausale Abhängigkeiten ist bekannt als *Differenzmethode* und geht auf den Philosophen, Politiker und Ökonomen John Stuart Mill (1806–1873) zurück (vgl. Mill (1843), S. 256).

Eine dritte Voraussetzung, die es zu explizieren gilt, ist das Prinzip der kausalen Eindeutigkeit. Dazu muss vorab der Begriff der *kompletten* Konfigurationstabelle $\mathcal{K}_F$ über F definiert werden. $\mathcal{K}_F$ über F ist dann komplett, wenn $\mathcal{K}_F$ alle empirisch möglichen Konfigurationen der Faktoren einer nicht unterspezifizierten Menge F enthält. Dabei ist F nicht unterspezifiziert, wenn jeder Pfad zu Wirkungen in F durch Literale mindestens eines Faktors aus F repräsentiert wird, sodass immer, wenn eine Wirkung aus F gegeben ist, relativ zu F auch jeweils mindestens eine vollständige Ursache dieser Wirkung gegeben ist. Bei der Komplettheit wird neben der Vollständigkeit einer Konfigurationstabelle über F nun also auch eine Anforderung an die Faktormenge F gestellt. Wie sich gezeigt hat (vgl. bspw. S. 51 f.), reflektieren gewisse Faktormengen die unterliegende Kausalstruktur richtig und andere nicht. Bei der Komplettheit geht es nun darum, Kriterien für Konfigurationstabellen anzugeben, die korrekte Reflektion von Kausalität garantieren und denen dennoch nicht nur eine derart grosse Konfigurationstabelle genügt, die alle empirisch möglichen Konfigurationen über die universelle Menge von allen möglichen (geeigneten) Faktoren enthält.

Definition 2.14 (Komplettheit)

Sei $\mathcal{K}_F$ eine Konfigurationstabelle über F . Seien weiter $F_1 \subseteq F$ und $F_2 \subseteq F$ die beiden folgenden Mengen von Faktoren aus F :

- Ein Faktor $W \in F$ ist genau dann in F_1 enthalten, wenn es (mindestens) ein Literal eines Faktors aus F gibt, das eine Ursache von W ist.
- Ein Faktor $X \in F$ ist genau dann in F_2 enthalten, wenn es (mindestens) ein Literal eines Faktors aus F gibt, das eine Wirkung von X ist.

$\mathcal{K}_F$ ist genau dann komplett, wenn gilt:

- (a) Jedes Literal eines Faktors $Z \notin F$, das kausal relevant für ein Literal W eines Faktors aus F_1 ist, verursacht W ausschliesslich über Literale von Faktoren aus F_2 ,
- (b) alle empirisch möglichen Konfigurationen der Faktoren aus F sind in $\mathcal{K}_F$ enthalten.

Nimmt man beispielsweise an, dass die Konfigurationstabelle $\mathcal{K}_{2.6}$ über $C = \{A, B, C, D, E\}$ aus Abbildung 2.6 komplett ist (vgl. S. 31), folgt daraus, dass in der damit korrespondierenden Kausalstruktur aus derselben Abbildung alle Pfade hin zu B und C in C repräsentiert sind und allfällige, weitere für B und C kausal relevante Literale jeweils nur über A , D oder E die Wirkung verursachen. In $\mathcal{K}_{2.6}$ ist also das gesamte Auftritsverhalten der vollständigen Ursachen von B und C abgebildet, wobei es keine weiteren Ko-Literale der Ursachen oder Alternativursachen von B und C gibt, deren Verhalten nicht in $\mathcal{K}_{2.6}$ wiedergegeben ist. $\mathcal{K}_{2.6}$ ist nicht mehr komplett, wenn eine Zeile und/oder eine der Spalten A , D oder E gestrichen wird. Vor dem Hintergrund derart vollständiger Informationen muss sich eindeutig die Struktur aus Abbildung 2.6 ableiten lassen. Allgemein wird diese Voraussetzung vom Prinzip der kausalen Eindeutigkeit festgehalten:

Prinzip der kausalen Eindeutigkeit: Jede komplette Konfigurationstabelle korrespondiert genau mit einer Kausalstruktur.

Das Prinzip der kausalen Eindeutigkeit ist neu und unterscheidet sich von den obigen beiden altbekannten Prinzipien. Motiviert durch Ambiguitäten, die sich bei Kausalanalysen ergeben können (vgl. Abschnitt 2.7), muss dieses Prinzip angenommen werden, um die Difference-Maker-Idee regularitätstheoretisch konsistent umzusetzen. Denn wie bereits beispielhaft anhand der Konfigurationstabelle $\mathcal{K}_{2.11}$ gesehen, wird man oft mit der Schwierigkeit konfrontiert, dass aus Konfigurationen $\mathcal{K}_F$ mehrere, kausal interpretiert voneinander zu unterscheidende, aber logisch äquivalente Minimale Theorien von $\mathcal{K}_F$ abgeleitet werden können. Bei $\mathcal{K}_{2.11}$ über

$\mathbf{E} = \{A, B, C, D, E\}$ wurden diese Ambiguitäten über die Erweiterungsklausel aufgelöst: Durch eine Analyse der Konfigurationstabelle $\mathcal{K}_{2.12}$ über $\mathbf{G} = \{A, B, C, D, E, F\}$, die im Gegensatz zu $\mathcal{K}_{2.11}$ zusätzlich die Ursache F von E berücksichtigt, hat sich ergeben, dass die mit (2.14) einhergehende Kausalhypothese zugunsten von jener, die von (2.15) zum Ausdruck gebracht wird, zurückgewiesen werden kann (vgl. S. 45 bzw. S. 49). Dies wiederum bedingt jedoch, dass ein derartiger Faktor F auch tatsächlich existiert und die relativ zu $\mathbf{E}$ untersuchte Konfigurationstabelle $\mathcal{K}_{2.11}$ nicht alle empirisch möglichen Konfigurationen enthält. Das Prinzip der kausalen Eindeutigkeit fängt genau diesen Umstand ein und besagt, dass allfällige Mehrdeutigkeiten das Produkt von Mängeln in den Daten sind – sind diese Mängel behoben, lösen sich auch die Mehrdeutigkeiten auf. Ohne dieses Prinzip wäre Letzteres nicht garantiert, sodass eine Regularitätstheorie gezwungen wäre, sowohl (2.14) als auch (2.15) kausal zu interpretieren. Weder das Determinismus- noch das Kausalitätsprinzip fangen dieses Problem ab, zumal (2.14) und (2.15) für sich kausal interpretiert keines dieser beiden Prinzipien verletzen und die dazugehörige Konfigurationstabelle $\mathcal{K}_{2.11}$ von den beiden Prinzipien somit auch nicht als mangelhaft ausgezeichnet wird.

Es bleibt zu bemerken, dass aus dem Prinzip der kausalen Eindeutigkeit weder folgt, dass mit jeder Konfigurationstabelle genau eine Kausalhypothese vereinbar sein muss, noch dass der hier verwendete Verursachungsbegriff nur genau dann für kausales Modellieren eingesetzt werden kann, wenn die untersuchte Konfigurationstabelle frei von Mängeln wie Fragmentierung ist. Das Prinzip besagt, dass die Phänomene, die Gegenstand des konfiguralen kausalen Modellierens sind, letztlich einen eindeutigen kausalen Aufbau haben, auch wenn dieser Aufbau nicht vollständig erfassbar ist. Das Prinzip besagt anders formuliert, dass es in unserer Welt keine Struktur gibt, die eine Konfigurationstabelle $\mathcal{K}_F$ generiert, aus der auch dann kausale Mehrdeutigkeiten folgen, wenn $\mathcal{K}_F$ komplett ist. Vor dem Hintergrund der in dieser Arbeit präsentierten Kausaltheorie werden solche Strukturen als künstlich konstruiert betrachtet, das heisst, sie entsprechen nicht Kausalstrukturen, die in der Welt, in der wir leben, existieren.

Die obigen drei Prinzipien werden vom Verursachungsbegriff im Sinne von Definition 2.11 vorausgesetzt. Die beiden nun folgenden Prinzipien – das Difference-Making-Prinzip und das Permanenz-Prinzip – stellen Implikationen des Verursachungsbegriffs dar (vgl. bspw. Graßhoff & May (2001), S. 89 f.).²⁴ Sie verhindern, dass Scheinursachen in eine Kausalstruktur aufgenommen werden. Als Scheinursachen werden solche Literale betrachtet, die (1) zwar Teil einer notwendigen oder hinreichenden Bedingung einer Wirkung sind, jedoch keine Difference-Maker dieser Wirkung sind, oder Literale, die (2) relativ zu einer Faktormenge Teil einer einfachen Mi-

²⁴In Graßhoff & May (2001) und Baumgartner & Graßhoff (2004) wird anstelle vom „Difference-Making-Prinzip“ bzw. „Permanenz-Prinzip“ vom „Prinzip der (kausalen) Relevanz“ bzw. „Prinzip der persistenten Relevanz“ gesprochen (vgl. S. 89 f. bzw. S. 68 ff.).

nimalen Theorie einer Wirkung sind, bei geeigneten Erweiterungen aber nicht permanenter Teil einer einfachen Minimalen Theorie dieser Wirkung bleiben.

Difference-Making-Prinzip: Jede Ursache ist ein Difference-Maker der Wirkung.

Permanenz-Prinzip: Eine Ursache bleibt ein Difference-Maker der Wirkung, wenn zusätzliche Faktoren in die Betrachtung miteinbezogen werden.

Während das Difference-Making-Prinzip aus der Bedingung (a) der Definition 2.11 des Verursachungsbegriffs folgt, lässt sich das Permanenz-Prinzip aus der Bedingung (b) derselben Definition ableiten. Beide Prinzipien wurden in den vorhergehenden Abschnitten unter anderem bereits basierend auf Beispielen wie der Common-Cause-Struktur aus Abbildung 2.6 oder der Struktur aus Abbildung 2.9 diskutiert und deren methodologische Relevanz aufgezeigt (vgl. S. 31 ff. bzw. S. 52 f.): Sie unterscheiden kausale von nicht kausalen Regularitäten.

2.10.2. Kausale Wohlgeformtheit und minimaler Grad an kausaler Komplexität

Die kausalen Prinzipien liefern wesentliche Kriterien dafür, ob ein Hypergraph das Abbild einer Kausalstruktur sein kann oder nicht. Nicht jede beliebige, einen Hypergraphen bildende Verknüpfung von Literalen stellt kausale Abhängigkeiten dar. Damit ein Graph ein Abbild einer Kausalstruktur sein kann, muss er *kausal wohlgeformt* sein. Ein Graph ist nur genau dann kausal wohlgeformt, wenn er dem Determinismus-, dem Kausalitäts-, dem Difference-Making- sowie dem Permanenz-Prinzip genügt:

Definition 2.15 (Kausale Wohlgeformtheit)

Ein gerichteter Hypergraph ist genau dann kausal wohlgeformt, wenn er dem Determinismus-, dem Kausalitäts-, dem Difference-Making- sowie dem Permanenz-Prinzip genügt.

Das Prinzip der kausalen Eindeutigkeit wird für die Definition der kausalen Wohlgeformtheit bewusst ausgeklammert. Der Grund dafür sind Ambiguitäten, die aus Konfigurationstabellen über eine Faktormenge folgen können. Beispielsweise lässt sich entweder (2.14) oder (2.15) als Kausalstruktur interpretieren, die den Konfigurationen $\mathcal{K}_{2,11}$ zugrunde liegt (vgl. S. 45). Auch wenn sich mit Hinzunahme weiterer Faktoren die aus (2.15) hervorgehende Struktur nicht als die den Konfigurationen tatsächlich zugrundeliegende Kausalstruktur herausstellt, spricht relativ zu $\mathcal{K}_{2,11}$ nichts dagegen, diese als mögliche Kausalstruktur zu betrachten. Das Prinzip der kausalen Eindeutigkeit sagt nichts über die kausale Wohlgeformtheit aus, aber etwas darüber,

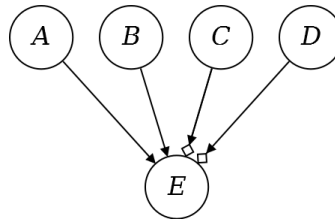


Abbildung 2.15.: Beispiel eines Hypergraphens, der nicht kausal wohlgeformt ist.

ob ein Hypergraph eine *allumfassende* Kausalstruktur darstellen kann. Was das genau bedeutet, wird sich sogleich zeigen. Zuerst soll jedoch mit dem Beispiel aus Abbildung 2.15 ein Hypergraph betrachtet werden, der nicht kausal wohlgeformt und in dieser Form kein Repräsentant von kausalen Abhängigkeiten ist.

Würde man diesen Graphen kausal interpretieren, gäbe er zwei Alternativursachen A und B von E und zwei Alternativursachen C und D von $\bar{E}$ wieder, wobei jede der vier Ursachen unabhängig von den drei anderen E beziehungsweise $\bar{E}$ verursachen würde. Dass dieser Graph jedoch keine kausale Struktur repräsentiert, wird ersichtlich, wenn man die dazugehörigen Konfigurationen betrachtet. Konkret ergibt sich unter anderem ein Widerspruch bei jener Konfiguration, bei der die Faktoren A , B , C und D alle anwesend sind. Bei dieser Konfiguration ist sowohl A als auch B gegeben, womit gemäss dem Graphen Ursachen von E gegeben sind und folglich auch E anwesend ist. Andererseits sind sowohl C als auch D anwesend, womit Ursachen dafür gegeben sind, dass E ausbleibt. Insgesamt folgt, dass E bei konstant anwesenden Faktoren A , B , C und D sowohl an- als auch abwesend ist. Dieser Widerspruch wird über das Determinismusprinzip abgefangen: Kann die Wirkung E bei konstant anwesenden Faktoren A , B , C und D sowohl gegeben als auch nicht gegeben sein, verletzt dies den einzuhaltenden Grundsatz, wonach gleiche Ursachen von gleichen Wirkungen begleitet werden.

Mit der Folgerung, dass der Hypergraph aus Abbildung 2.15 nicht kausal wohlgeformt ist, geht nicht einher, dass keine Kausalaussagen sowohl zur An- als auch zur Abwesenheit eines Faktors W gemacht werden können. Es bedeutet lediglich, dass dies nur unter Einsatz entweder des positiven Literals W oder des negativen Literals $\bar{W}$ dargestellt werden kann. Sind die kausalen Abhängigkeiten für W oder $\bar{W}$ dargestellt, implizieren erstere aufgrund der Zweiwertigkeit die kausalen Abhängigkeiten für letztere und umgekehrt. Dies zeigt sich auch in den jeweiligen einfachen Minimalen Theorien, die gemäss Definition 2.11 zur Beschreibung eines kausalen Zusammenhangs zwischen einem Literal W und den für W kausal relevanten Literalen herangezogen werden: Die Ursachen bilden zusammen mit der Wirkung ein Bikonditional und es gilt, dass das Komplement der Ursachen von W Ursachen von $\bar{W}$ sind und umgekehrt. Sind beispiels-

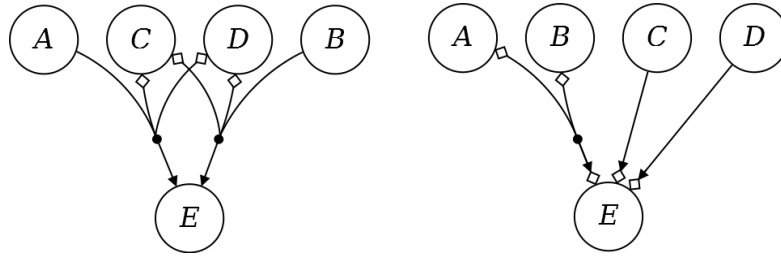


Abbildung 2.16.: Zwei wohlgeformte Hypergraphen, die kausale Abhängigkeiten darstellen können.

weise A und B zwei Alternativursachen von W , ist $\bar{A}\bar{B}$ eine Ursache von $\bar{W}$.²⁵ Je nachdem ob der Fokus auf die An- oder die Abwesenheit eines Faktors gerichtet ist, sind in Abbildung 2.16 zwei von mehreren möglichen wohlgeformten Hypergraphen dargestellt, die unter anderem obige Kausalaussagen wiedergeben, wonach A und B kausal relevant für E und C sowie D kausal relevant für $\bar{E}$ sind.

Genügt ein kausal wohlgeformter Hypergraph zusätzlich dem Prinzip der kausalen Eindeutigkeit, kann er nicht nur eine kausale Struktur repräsentieren, sondern darüber hinaus eine allumfassende Kausalstruktur. Eine Kausalstruktur ist genau dann allumfassend, wenn für jede darin verortete Wirkung alle kausalen Pfade hin zu dieser Wirkung durch Literale repräsentiert werden oder anders ausgedrückt, wenn die von der Struktur generierte Konfigurationstabelle komplett ist. So wird etwa die aus lediglich zwei Literalen A und C bestehende Kausalstruktur, gemäss welcher A eine Ursache von C ist, nicht als allumfassend betrachtet. Vielmehr stellt sie nur einen Teil einer komplexeren Kausalstruktur dar, die nicht in ihrer Ganzheit erfasst ist. Würde man davon ausgehen, dass diese Struktur allumfassend und demnach die dazugehörige, aus zwei Konfigurationen bestehende Konfigurationstabelle $\mathcal{K}_{2.17a}$ aus Tabelle 2.17 komplett ist, wäre nicht klar, in welche Richtung die entsprechende einfache Minimale Theorie $A \leftrightarrow C$ kausal zu interpretieren ist. In diesem Fall ist nicht nur A in einer einfachen Minimalen Theorie von C , sondern C auch in einer einfachen Minimalen Theorie von A . Da A gemäss der angenommenen allumfassenden Kausalstruktur das einzige kausal relevante Literal von C ist, wird auch bei jeder geeigneten Erweiterung der Faktormenge $\mathbf{J} = \{A, C\}$ diese einfache Minimale Theorie folgen und es lässt sich insgesamt sagen, dass sowohl A permanenter Teil einer einfachen Minimalen Theorie von C als auch C permanenter Teil einer einfachen Minimalen Theorie von A ist. Die daraus gefolgerte Kausalstruktur entspricht somit einer Wechselwirkung zwischen A

²⁵In der QCA-Literatur wird eine derartig direkte Transformierbarkeit einer einfachen Minimalen Theorie von W in eine einfache Minimale Theorie von $\bar{W}$ im Normalfall als unzulässig betrachtet. Das liegt daran, dass QCA jeweils von (möglicherweise mangelhaften) Realdaten ausgeht, sodass nur approximative Regularitäten abgeleitet werden können. Bei solchen Daten verlangt QCA für die Wirkungen W und $\bar{W}$ zwei gesonderte Analysen (vgl. bspw. Schneider & Wagemann (2012) S. 81 f.). Im vorliegenden Fall gehen wir jedoch davon aus, dass alle empirisch möglichen Konfigurationen vorliegen, weshalb diese Transformation unproblematisch ist.

		$\mathcal{K}_{2.17b}$	A	B	C		$\mathcal{K}_{2.17c}$	A	B	C
		c_1	1	1	1		c_1	1	1	1
		c_2	1	0	1		c_2	1	0	0
		c_3	0	1	1		c_3	0	1	0
		c_4	0	0	0		c_4	0	0	0
$\mathcal{K}_{2.17a}$	A	C				(a)				(b)
	c_1	1	1					c_1	1	1
	c_2	0	0					c_2	1	0
								c_3	0	1
								c_4	0	0
							(c)			

Tabelle 2.17.: Konfigurationstabellen $\mathcal{K}_{2.17a}$, $\mathcal{K}_{2.17b}$ und $\mathcal{K}_{2.17c}$.

und C . Tatsächlich gibt es aber zwei weitere Strukturen, die genau mit derselben permanenten einfachen Minimalen Theorie einhergehen: einerseits jene, wonach A eine Ursache von C ist und andererseits jene, wonach C eine Ursache von A ist, aber jeweils nicht umgekehrt. Aufgrund der Symmetrie der funktionalen Abhängigkeit ist somit unklar, was die Richtung der Verursachung ist und es lässt sich basierend auf den Konfigurationen nicht feststellen, ob nun A die Ursache von C ist, C die Ursache von A oder eine Wechselwirkung zwischen A und C besteht. Wie jeder andere Kausalzusammenhang hat jedoch auch die zwischen A und C bestehende kausale Abhängigkeit eine klare Richtung, nämlich von der Ursache zur Wirkung, das heisst im vorliegenden Fall von A nach C . Daraus lässt sich weder ableiten, dass C kausal relevant für A ist, noch dass $\bar{C}$ kausal relevant für $\bar{A}$ ist – kurz: Die Kausalbeziehung ist nicht symmetrisch (vgl. u. a. Baumgartner (2006), S. 111 f.).²⁶

Unter der Voraussetzung, dass die lediglich aus A und C bestehende Kausalstruktur allumfassend und $\mathcal{K}_{2.17a}$ komplett ist, ist das Prinzip der kausalen Eindeutigkeit verletzt. Da wir jedoch die Gültigkeit dieses Prinzip voraussetzen, folgt, dass jede allumfassende Kausalstruktur einen gewissen minimalen Grad an kausaler Komplexität aufweist. Aufgrund der unbestimmbaren Richtung der Relation bei einer perfekten Korrelation zweier Literale muss für eine allumfassende Kausalstruktur als Mindestanforderung gelten, dass jede Wirkung mindestens zwei Alternativursachen oder eine komplexe Ursache besitzt. Besitzt C neben A eine weitere Alternativursache B , ergibt sich aus der entsprechenden kompletten Konfigurationstabelle $\mathcal{K}_{2.17b}$ aus Tabelle 2.17 die einfache Minimale Theorie $A + B \leftrightarrow C$. Diese einfache Minimale Theorie lässt sich nur in eine Richtung kausal interpretieren: A und B sind Ursachen von C . Eine Interpretation in entgegengesetzter Richtung, wonach C eine Ursache von $A + B$ ist, ist vor dem Hintergrund des Determinismusprinzips nicht möglich, da C weder für A (vgl. c_3 aus $\mathcal{K}_{2.17b}$) noch für B (vgl. c_2 aus $\mathcal{K}_{2.17b}$) hinreichend ist. Ist die Ursache von C komplex und besteht aus der Konjunktio-

²⁶Eine Relation R ist asymmetrisch, wenn für alle Elemente einer Menge $\mathbf{M} = \{x, y, \dots\}$ gilt: Aus xRy folgt $\neg(yRx)$.

Eine Relation R ist antisymmetrisch, wenn für alle Elemente einer Menge $\mathbf{M} = \{x, y, \dots\}$ gilt: Aus xRy und yRx folgt $x = y$. Wird R als „ist kausal relevant für“ interpretiert zeigt sich bspw. bei der Betrachtung einer Wechselwirkung, dass Kausalrelationen diese beiden Eigenschaften nicht besitzen.

on AB , ergibt sich aus der entsprechenden kompletten Konfigurationstabelle $\mathcal{K}_{2.17c}$ aus Tabelle 2.17 die einfache Minimale Theorie $AB \leftrightarrow C$. Diese aus einer kompletten Konfigurationstabelle abgeleitete einfache Minimale Theorie lässt sich ebenfalls nur in eine Richtung kausal interpretieren: AB ist eine Ursache von C . Eine Interpretation in entgegengesetzter Richtung, wonach C eine Ursache von AB ist, ist vor dem Hintergrund des Kausalitätsprinzips nicht möglich, da A beziehungsweise B zwar in c_1 mit ihrer Ursache C auftreten, nicht aber in c_2 beziehungsweise c_3 .

2.11. Komplexe Minimale Theorien und komplexe Kausalstrukturen

Einfache Minimale Theorien können isoliert betrachtet Kausalstrukturen widerspiegeln, die eine einzige Wirkung beinhalten. Um aus Daten über eine Faktormenge auch Kausalstrukturen mit mehr als einer Wirkung ableiten zu können, wie beispielsweise die Common-Cause-Struktur aus Abbildung 2.6, werden entsprechend auch mehrere einfache Minimale Theorien herangezogen. Während also die beiden einfachen Minimalen Theorien $A + E \leftrightarrow B$ beziehungsweise $A + D \leftrightarrow C$ aus der Minimalen Theorie $(A + E \leftrightarrow B) \cdot (A + D \leftrightarrow C)$ von $\mathcal{K}_{2.6}$ jeweils die Grundlage zur Erfassung der Ursachen von B beziehungsweise der Ursachen von C darstellen, bildet deren Konjunktion die Grundlage zur Erfassung der Common-Cause-Struktur als Ganzes. Solche Konjunktionen von einfachen Minimalen Theorien, die zusammenhängende Strukturen mit mehreren Wirkungen wiedergeben, werden als *komplexe* Minimale Theorien bezeichnet (vgl. u. a. Baumgartner & Graßhoff (2004), S. 308; Baumgartner (2009), S. 76).

Nicht jede beliebige Konjunktion von einfachen Minimalen Theorien, die in einer Minimalen Theorie von $\mathcal{K}_F$ enthalten ist, ist eine komplexe Minimale Theorie. Neben der trivialen Voraussetzung, dass die Konjunktion aus mindestens zwei einfachen Minimalen Theorien bestehen muss, müssen in diesen einfachen Minimalen Theorien darüber hinaus Literale derselben Faktoren enthalten sein oder kurz: Sie müssen Faktorüberschneidungen aufweisen.²⁷ Überschneiden sich keine zwei einfache Minimale Theorien eines solchen Konjunktionsterms in mindestens einem Faktor, können aus der Konjunktion ableitbare kausale Abhängigkeiten auch keine zusammenhängende komplexe Kausalstruktur mit mehreren Wirkungen beschreiben, sondern lediglich voneinander unabhängige Teilstrukturen mit jeweils einer Wirkung. Im konkreten Fall der komplexen Minimalen Theorie $(A + E \leftrightarrow B) \cdot (A + D \leftrightarrow C)$ ist diese Überschneidung in Form

²⁷Bei der Überschneidung muss es sich nicht zwingend um das gleiche Literal handeln, sondern lediglich um den gleichen Faktor, von dem Literale in den einfachen Minimalen Theorien enthalten sind. So überschneiden sich beispielsweise die beiden einfachen Minimalen Theorien aus (2.10) in den Faktoren A und C (vgl. S. 41).

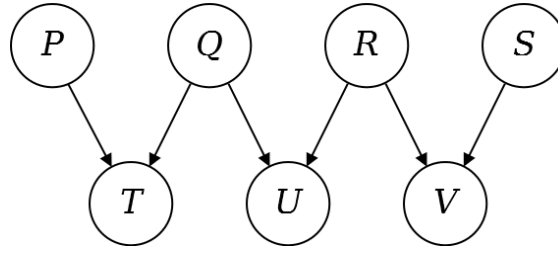


Abbildung 2.18.: Common-Cause-Struktur mit drei Wirkungen.

des Faktors A gegeben. Ist A permanenter Teil beider einfacher Minimaler Theorien, handelt es sich um eine gemeinsame Ursache von B und C und entspricht einem Literal, das die beiden kausalen Teilstrukturen von B und C zu einer zusammenhängenden komplexen Kausalstruktur mit zwei Wirkungen verbindet.

Die geforderte Bedingung der Überschneidung ist nicht etwa äquivalent zur Forderung, dass zwei beliebige einfache Minimale Theorien einer komplexen Minimalen Theorie jeweils Faktorüberschneidungen aufweisen müssen. Das Beispiel aus Abbildung 2.18 soll deutlich machen, dass zwei einfache Minimale Theorien einer komplexen Minimalen Theorie auch indirekt über weitere einfache Minimale Theorien zusammenhängen können.

Die folgende, aus den drei einfachen Minimalen Theorien β_1 , β_2 und β_3 bestehende komplexe Minimale Theorie lässt sich im Sinne der Kausalstruktur aus Abbildung 2.18 interpretieren:

$$\underbrace{(P + Q \leftrightarrow T)}_{\beta_1} \cdot \underbrace{(Q + R \leftrightarrow U)}_{\beta_2} \cdot \underbrace{(R + S \leftrightarrow V)}_{\beta_3} \quad (2.19)$$

Sowohl für die Paarung β_1 und β_2 als auch für die Paarung β_2 und β_3 gilt, dass sie Faktorüberschneidungen haben. Die erste Paarung weist das gemeinsame Literal Q auf und die beiden einfachen Minimalen Theorien der zweiten Paarung enthalten beide R . Die Forderung gilt jedoch nicht für die Paarung β_1 und β_3 . Trotzdem möchte man im Hinblick auf eine Erfassung der gesamten Common-Cause-Struktur diese beiden einfachen Minimalen Theorien als Teil einer komplexen Minimalen Theorie (2.19) betrachten. Die einfachen Minimalen Theorien β_1 und β_3 sind dabei sozusagen indirekt über β_2 miteinander verknüpft. Konkret gibt es für β_1 und β_3 mit $(\beta_1, \beta_2, \beta_3)$ eine Folge von einfachen Minimalen Theorien, bei der jede Paarung von benachbarten einfachen Minimalen Theorien Faktorüberschneidungen aufweisen.

Literale, die Faktorüberschneidungen erzeugen, müssen ebenfalls nicht jeweils Teil der RDN-Bedingung von einfachen Minimalen Theorien sein. Wie sich später im Zusammenhang mit Kausalketten noch zeigen wird, aber bereits an dieser Stelle ohne Probleme nachvollziehbar sein dürfte, können zwei einfache Minimale Theorien β_i und β_j einer komplexen Minimalen

Theorie auch derart eine Faktorüberschneidung in einem Faktor X aufweisen, dass β_i eine einfache Minimale Theorie von X und β_j eine einfache Minimale Theorie von $Y \neq X$ ist, die X in ihrer RDN-Bedingung enthält. Ist $(A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E)$ eine Minimale Theorie, haben die beiden darin enthaltenen einfachen Minimalen Theorien C gemeinsam – diese komplexe Minimale Theorie deutet auf eine Kausalkette hin (mehr dazu vgl. Abschnitt 2.11.2). Zwei einfache Minimale Theorien β_i und β_j einer komplexen Minimalen Theorie können aber nicht derart Faktorüberschneidungen aufweisen, dass es sich sowohl bei β_i als auch bei β_j um eine einfache Minimale Theorie von X handelt. Eine solche Überschneidung ist unzulässig, zumal in diesem Fall eine Konjunktion von β_i und β_j Bedingung (c) der Definition 2.8 der Minimalen Theorie verletzt und es sich damit auch nicht um eine komplexe Minimale Theorie handeln kann.

Schliesslich ist es wichtig zu verdeutlichen, dass eine komplexe Minimale Theorie einer Konfigurationstabelle $\mathcal{K}_F$ über F ein echter oder ein unechter Teil einer Minimalen Theorie von $\mathcal{K}_F$ über F sein kann. Bei den bisherigen Beispielen sind die komplexen Minimalen Theorien jeweils äquivalent zu den Minimalen Theorien der entsprechenden Konfigurationstabellen. Das liegt nicht daran, dass eine solche Äquivalenz grundsätzlich gilt, sondern daran, dass den betrachteten Konfigurationstabellen eine einzige zusammenhängende komplexe Kausalstruktur zugrunde liegt. Dementsprechend stellen auch die Minimalen Theorien dieser Konfigurationen korrekterweise als Ganzes komplexe Minimale Theorien dar. Tatsächlich können einer Konfigurationstabelle aber auch Kausalstrukturen zugrunde liegen, die aus mindestens zwei voneinander unabhängigen Teilstrukturen bestehen. Es spricht beispielsweise nichts dagegen, dass ein Ausdruck der Form $(A + D \leftrightarrow C) \cdot (A + E \leftrightarrow B) \cdot (X + Y \leftrightarrow Z)$ eine Minimale Theorie einer Konfigurationstabelle darstellt. Mit $(A + D \leftrightarrow C) \cdot (A + E \leftrightarrow B)$ wäre darin eine komplexe Minimale Theorie enthalten, die ein echter Teil der Minimalen Theorie darstellt.

Basierend auf obigen Überlegungen lässt sich der Begriff der komplexen Minimalen Theorie insgesamt folgendermassen definieren:

Definition 2.16 (Komplexe Minimale Theorie)

Sei Ψ' eine Konjunktion von einfachen Minimalen Theorien $\beta_1 \cdot \dots \cdot \beta_n$, $2 \leq n$, die in einer Minimalen Theorie Ψ von $\mathcal{K}_F$ über F enthalten ist. Ψ' ist genau dann eine komplexe Minimale Theorie relativ zu $\mathcal{K}_F$ über F , wenn für zwei beliebige β_i und β_j aus Ψ' gilt: Es gibt (mindestens) eine endliche Folge von einfachen Minimalen Theorien $(\beta_i, \dots, \beta_r, \dots, \beta_j)$ aus Ψ' , $i < r \leq j$, sodass es für jede einfache Minimale Theorie β_s aus dieser Folge und deren Folglied β_{s+1} mindestens einen Faktor aus F gibt, von dem Literale sowohl in β_s als auch in β_{s+1} enthalten sind.

Komplexe Minimale Theorien bilden die Grundlage zur Erfassung von komplexen und zusammenhängenden Kausalstrukturen mit mehreren Wirkungen. Konkret werden damit Common-

Cause-Strukturen, Kausalketten und Feedbackstrukturen erfasst. Dieser Reihenfolge entsprechend werden in den nächsten drei Abschnitten die dazugehörigen Definitionen entwickelt.

2.11.1. Common-Cause-Strukturen

Ein wesentliches Merkmal von Common-Cause-Strukturen ist bekanntlich jenes, dass sie eine gemeinsame Ursache (Common-Cause) von mindestens zwei verschiedenen Wirkungen aufweisen. Wie unter anderem bereits anhand des Beispiels aus Abbildung 2.18 und der dazugehörigen Minimalen Theorie aufgezeigt, werden diese kausalen Abhängigkeiten von komplexen Minimalen Theorien eingefangen, bei denen wenigstens ein und dasselbe Literal Teil von mindestens zwei einfachen Minimalen Theorien von Literalen unterschiedlicher Faktoren ist. Damit die in der komplexen Minimalen Theorie enthaltenen Regelmässigkeiten schliesslich auch kausal interpretiert werden können, muss gemäss der Definition 2.11 ferner die Erweiterungsklausel gelten. Beide Bedingungen zusammen liefern die Definition für eine *gemeinsame Ursache* mehrerer unterschiedlicher Wirkungen:

Definition 2.17 (Gemeinsame Ursache)

Sei ϵ ein Konjunktionsterm von Literalen $Z_1 \cdots Z_n$ mit $n \geq 1$. ϵ ist genau dann eine gemeinsame Ursache von mehreren Literalen $W_1, \dots, W_n$, $n \geq 2$, wenn es eine die Faktoren $Z_1, \dots, Z_n, W_1, \dots, W_n$ enthaltende Faktormenge $\mathbf{F}$ derart gibt, dass

- (a) es eine komplexe Minimale Theorie Ψ relativ zu allen empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}$ gibt, die von jedem Literal W_i eine einfache Minimale Theorie enthält, von der ϵ Teil ist und
- (b) es für jede geeignete Erweiterung $\mathbf{F}' \supset \mathbf{F}$ von $\mathbf{F}$ eine komplexe Minimale Theorie gibt, die relativ zu $\mathbf{F}'$ (a) erfüllt.

Neben dem Merkmal der gemeinsamen Ursache von mindestens zwei verschiedenen Wirkungen W_i und W_j ist das zweite charakteristische Kennzeichen einer Common-Cause-Struktur jenes, dass die Wirkungen nicht kausal voneinander abhängig sind. Bei einer Kausalkette zwischen drei Literalen A , B und C , die in dieser Reihenfolge kausal miteinander verknüpft sind, kann A im Grunde auch als gemeinsame Ursache von B und C betrachtet werden. Trotzdem möchte man die Kausalstruktur als Ganzes nicht als Common-Cause-Struktur bezeichnen, weil auch die beiden Wirkungen B und C kausal voneinander abhängig sind. Unter Berücksichtigung der Definition 2.11 wird dieses zweite Merkmal einer Common-Cause-Struktur so ausgedrückt, dass für beliebige, eine gemeinsame Ursache ϵ aufweisende Wirkungen W_i und W_j gefordert wird,

dass W_i kein permanenter Teil einer einfachen Minimalen Theorie von W_j ist. Darauf basierend soll der Begriff der Common-Cause-Struktur folgendermassen definiert werden:

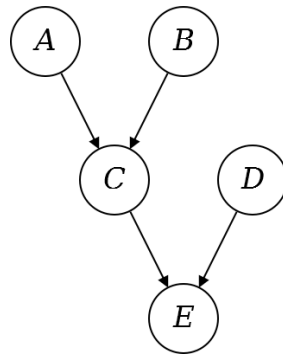
Definition 2.18 (Common-Cause-Struktur)

Sei ϵ ein Konjunktionsterm von Literalen $Z_1 \cdots Z_n$ mit $n \geq 1$. ϵ bildet zusammen mit mehreren Literalen $W_1, W_2, \dots, W_n$, $n \geq 2$, genau dann eine Common-Cause-Struktur, wenn ϵ gemeinsame Ursache von $W_1, W_2, \dots, W_n$ ist und für zwei beliebige paarweise verschiedene W_i und W_j gilt, dass W_i kein permanenter Teil einer einfachen Minimalen Theorie von W_j ist.

2.11.2. Kausalketten

Die Definition von Kausalketten kann als letztes wesentliches Puzzleteil betrachtet werden, um die gesamte, in Abschnitt 2.1.2 dargestellte Komplexität von Kausalstrukturen einzufangen. Auf Basis dieser Definition wird es möglich sein, kausale Relevanz in direkte und in indirekte kausale Relevanz aufzuspalten. Ferner bietet sie die Grundlage für die Definition von Feedbackstrukturen, die in gewisser Hinsicht ein Spezialfall von Kausalketten darstellen, da sie geschlossene kausale Verkettungen sind.

Zur Veranschaulichung soll die Kausalkette aus Abbildung 2.19 über die Faktormenge $\mathbf{K} = \{A, B, C, D, E\}$ betrachtet werden, die zwei indirekte Kausalbeziehungen aufweist: die indirekte, über C führende kausale Relevanz von A und B für E .



$\mathcal{K}_{2.19}$	A	B	C	D	E
c_1	1	1	1	1	1
c_2	1	1	1	0	1
c_3	1	0	1	1	1
c_4	1	0	1	0	1
c_5	0	1	1	1	1
c_6	0	1	1	0	1
c_7	0	0	0	1	1
c_8	0	0	0	0	0

Abbildung 2.19.: Kausalkette zusammen mit den dazugehörigen Konfigurationen $\mathcal{K}_{2.19}$ über $\mathbf{K} = \{A, B, C, D, E\}$.

Insgesamt gibt es zwei Minimale Theorien von $\mathcal{K}_{2.19}$:

$$(A + B \leftrightarrow C) \cdot (A + B + D \leftrightarrow E) \quad (2.20)$$

$$(A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E) \quad (2.21)$$

Als erstes kann festgestellt werden, dass die zur Kette gehörige Konfigurationstabelle $\mathcal{K}_{2.19}$ Ambiguitäten impliziert. Die komplexe Minimale Theorie (2.20) auf der einen Seite besitzt die wesentlichen strukturellen Eigenschaften einer Common-Cause-Struktur. Tatsächlich sind die Konfigurationen der Common-Cause-Struktur aus Abbildung 2.20, die (2.20) entspricht, genau dieselben wie jene, die in der Konfigurationstabelle aus Abbildung 2.19 enthalten sind.

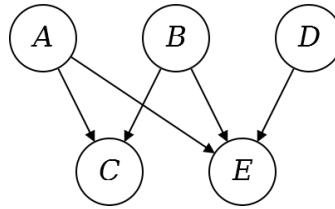


Abbildung 2.20.: Kausalgraph der Common-Cause-Struktur, deren Konfigurationen dieselben sind wie jene der Kausalkette aus Abbildung 2.19.

Die komplexe Minimale Theorie (2.21) auf der anderen Seite weist auf die tatsächliche, den Konfigurationen zugrunde liegende Kausalkette aus Abbildung 2.19 hin. C tritt dabei als Zwischenglied der Kette auf.

Im Zusammenhang mit der Erweiterungsklausel wurde bereits aufgezeigt, dass Kausalstrukturen aufgrund von unvollständigen Konfigurationstabellen über eine Faktormenge teils nicht eindeutig identifizierbar sind. Wie gleich ersichtlich wird, ist die Besonderheit im vorliegenden Fall, dass sich die Ambiguitäten auch bei der einzelnen beziehungsweise isolierten Betrachtung aller empirisch möglichen Konfigurationen über eine erweiterte Faktormenge $\mathbf{L} \supset \mathbf{K}$ nicht auflösen lassen und es unabhängig vom Grad der Komplexität einer Kausalkette auch immer jeweils Common-Cause-Strukturen gibt, deren Konfigurationen relativ zu $\mathbf{L}$ dieselben sind. Dieses Problem ist bekannt als das sogenannte *Kettenproblem* (vgl. Baumgartner (2008a)). Bevor näher auf dieses Kettenproblem eingegangen wird, soll der Fokus jedoch auf den Begriff der Kausalkette und dessen angemessener Definition gelegt werden. Mit (2.21) konnte dabei ein erster konkreter Eindruck darüber gewonnen werden, wie die komplexen Minimalen Theorien von kausalen Kettenstrukturen aussehen. Es stellt sich nun die Frage, welche Bedingungen genau zu erfüllen sind, damit eine solche komplexe Minimale Theorie auch tatsächlich als Kausalkette interpretiert werden kann. Dazu gilt es eine weitere zentrale Eigenschaft zu berücksichtigen, die der Verursa-

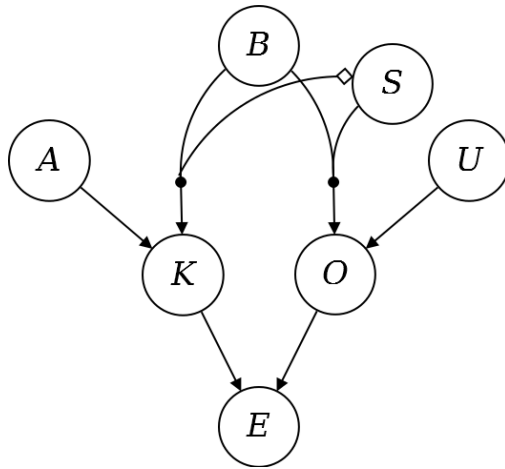
chungsbegriff aus Definition 2.11 hinsichtlich kausaler Verkettungen aufweist: die *Intransitivität* (vgl. Baumgartner (2013b), S. 97).

Intransitivität

Eine intransitive Relation R auf einer Menge $\mathbf{M}$ hat die Eigenschaft, dass für mindestens drei Elemente x , y und z aus $\mathbf{M}$ sowohl xRy als auch yRz gilt, nicht aber xRz . Das dazugehörige Gegenstück, wonach bei einer Relation R aus xRy und yRz immer xRz folgt, wird transitive Relation genannt.²⁸ Übertragen auf die Kausalbeziehung besagt die Intransitivität, dass es kausale Abhängigkeiten gibt, für die gilt, dass A kausal relevant für B , B kausal relevant für C , A aber nicht (indirekt) kausal relevant für C ist.

Innerhalb des kausaltheoretischen Forschungsfeldes ist die Frage nach der Transitivität der Kausalität nicht abschliessend geklärt. Einerseits ist angesichts der unzähligen Beispiele, wonach neben der kausalen Relevanz von A für B sowie B für C auch die indirekte kausale Relevanz von A für C intuitiv als gültig betrachtet wird, eine Tendenz hin zur Annahme der Transitivität erkennbar. Etwa Lewis (1973) betrachtet die Kausalität als transitiv und verbaut diese Eigenschaft fest in seine kontrafaktische Theorie der Kausalität, indem er kausale Relevanz über den transitiven Abschluss von kontrafaktischer Abhängigkeit definiert (vgl. S. 560 ff.). Seiner Theorie zufolge ist ein Ereignis a genau dann kausal relevant für ein anderes Ereignis b , wenn es eine von ersterem zu letzterem führende Kette $(a, x_1, x_2, \dots, x_n, b)$ von Ereignissen x_i , $1 \leq i \leq n$, gibt, sodass alle zwei benachbarten Ereignisse dieser Kette durch kontrafaktische Abhängigkeit verbunden sind. Ein Ereignis ist dabei genau dann kontrafaktisch von einem anderen Ereignis abhängig, wenn gilt, dass, wenn letzteres nicht stattgefunden hätte, auch ersteres nicht stattgefunden hätte und wenn letzteres stattgefunden hätte, auch ersteres stattgefunden hätte. Solche kontrafaktischen Konditionale wird man intuitiv oftmals anführen, um die Transitivität der kausalen Relation zwischen den eine Sequenz bildenden Literalen A , B und C zu stützen, indem gesagt wird, dass C ohne eine Instanz von A nicht aufgetreten wäre. Dennoch lassen einige einfache Beispiele Zweifel darüber aufkommen, ob dies für jede Abfolge von Literalen gilt, für die ein Literal jeweils kausal relevant für dessen Folglied ist. Ein bekanntes Beispiel hierzu ist etwa (A) der Verlust eines Fingers, der (B) durch einen chirurgischen Eingriff wieder angebracht wird und dadurch (C) seine volle Funktionsfähigkeit wiedererlangt. Die kausale Relevanz von A für B sowie jene von B für C würde man kaum bestreiten, dass jedoch A als kausal relevant

²⁸ Teilweise wird der oben definierte Begriff der Intransitivität auch als Non-Transitivität bezeichnet und die intransitive Relation derart definiert, dass für beliebige Elemente $x, y, z \in \mathbf{M}$ mit xRy und yRz nie xRz gilt (vgl. bspw. Lemmon (1998), S. 183). In der vorliegenden Arbeit bedeutet Intransitivität jedoch lediglich, dass aus xRy und yRz nicht immer xRz folgt.



$\mathcal{K}_{2,21}$	B	S	A	U	K	O	E
c_1	1	1	1	1	1	1	1
c_2	1	1	1	0	1	1	1
c_3	1	1	0	1	0	1	1
c_4	1	1	0	0	0	1	1
c_5	1	0	1	1	1	1	1
c_6	1	0	1	0	1	0	1
c_7	1	0	0	1	1	1	1
c_8	1	0	0	0	1	0	1
c_9	0	1	1	1	1	1	1
c_{10}	0	1	1	0	1	0	1
c_{11}	0	1	0	1	0	1	1
c_{12}	0	1	0	0	0	0	0
c_{13}	0	0	1	1	1	1	1
c_{14}	0	0	1	0	1	0	1
c_{15}	0	0	0	1	0	1	1
c_{16}	0	0	0	0	0	0	0

Abbildung 2.21.: Graphische Darstellung des hypothetischen kausalen Prozesses zur Senkung der betrieblichen Kosten zusammen mit der dazugehörigen Konfigurationstabelle $\mathcal{K}_{2,21}$ über $\mathbf{N} = \{B, S, A, U, K, O, E\}$.

für C betrachtet werden soll, scheint man intuitiv nicht akzeptieren zu wollen (vgl. z. B. Kvart (2001), S. 398 f.; Hitchcock (2001)).

Einem kausalen Pluralismus folgend stehen sich bei dieser Transitivitätsdebatte zwei inkonsistente vor-theoretische Kausalintuitionen gegenüber, von denen nicht eine wahr und die andere folglich falsch ist. Ob die kausale Relation transitiv ist oder nicht, wird typischerweise durch die Theorie der Kausalität zum Ausdruck gebracht, deren Angemessenheit durch ihren Anwendungskontext gerechtfertigt wird (vgl. S. 13). So ist anders als bei der kontrafaktischen Theorie von Lewis der über Minimale Theorien definierte Verursachungsbegriff intransitiv. Dies lässt sich anhand des folgenden hypothetischen Beispiels aus der Betriebswirtschaft nachweisen.

Angenommen ein Land befinde sich in einer schlechten Wirtschaftslage, worauf die Regierung ein Förderprogramm initiiert, das besonders stark betroffene Unternehmen bei der Erhaltung der Arbeitsplätze unterstützen soll. Ein Unternehmen selbst, das ein Gesuch um eine entsprechende Subvention einreichen kann, unterliege dabei insgesamt dem hypothetischen kausalen Kernprozess, der zusammen mit den dazugehörigen Konfigurationen über die Faktormenge $\mathbf{N} = \{B, S, A, U, K, O, E\}$ in Abbildung 2.21 dargestellt ist. Die einzelnen Faktoren sind dabei

folgendermassen definiert:

$S :=$ „Unternehmen erhält Fördergelder“ ($\bar{S} :=$ „Unternehmen erhält keine Fördergelder“),

$B :=$ „Unbedingte Bereitschaft des Unternehmens, die Betriebskosten zu senken, ohne Kündigungen auszusprechen“,

$A :=$ „Ausfall von essenziellen (und versicherten) Produktionsmaschinen“,

$U :=$ „Verpflichtung zu einer umweltschonenden Produktion“,

$K :=$ „Arbeitnehmer verrichten Kurzarbeit“,

$O :=$ „Das Unternehmen optimiert die Arbeitsprozesse“,

$E :=$ „Das Unternehmen erzielt Einsparungen“.

Die Zielsetzung eines Unternehmens, Kosten zur Erhaltung der Arbeitsplätze einzusparen (B), ist zusammen mit der Gewährung von staatlichen Subventionen (S) auf der einen Seite kausal relevant für das Vorhandensein genügender Ressourcen, die zur Optimierung von Arbeitsprozessen eingesetzt werden können. Die Prozessoptimierungen sind schliesslich ihrerseits kausal relevant für die Erzielung von Einsparungen (E). Auf der anderen Seite ist die Bereitschaft des Unternehmens, Kosten zu senken (B) zusammen mit dem Ausbleiben von staatlichen Subventionen ($\bar{S}$) kausal relevant für die Einführung von Kurzarbeit mit entsprechend reduzierter Entlohnung, die ihrerseits ebenfalls zu Einsparungen führt (E). Um dem Prinzip der kausalen Eindeutigkeit und dem damit einhergehenden minimalen Grad an Komplexität gerecht zu werden, seien die beiden Faktoren K und O mit jeweils einer Alternativursache ausgestattet (vgl. Abschnitt 2.10.2): Als Alternative zur unbedingten Bereitschaft, Einsparungen zu erzielen, kann auch ein länger anhaltender Ausfall spezifischer und essenzieller Maschinen, die im Schadenfall versichert sind (A), dazu führen, dass Arbeitnehmer Kurzarbeit verrichten. Ferner kann ebenfalls die Verpflichtung, umweltschonend zu produzieren (U), eine Ursache dafür sein, dass Arbeitsprozesse optimiert werden. Sowohl A als auch U führen vor dem Hintergrund der schlechten Wirtschaftslage indirekt auch zu Einsparungen (E).

Ausgehend von $\mathcal{K}_{2,21}$ lassen sich die folgenden vier Minimalen Theorien ableiten:

$$(BS + U \leftrightarrow O) \cdot (B\bar{S} + A \leftrightarrow K) \cdot (A + B + O \leftrightarrow E) \quad (2.22)$$

$$(BS + U \leftrightarrow O) \cdot (B\bar{S} + A \leftrightarrow K) \cdot (A + B + U \leftrightarrow E) \quad (2.23)$$

$$(BS + U \leftrightarrow O) \cdot (B\bar{S} + A \leftrightarrow K) \cdot (B + K + U \leftrightarrow E) \quad (2.24)$$

$$(BS + U \leftrightarrow O) \cdot (B\bar{S} + A \leftrightarrow K) \cdot (K + O \leftrightarrow E) \quad (2.25)$$

Zum einen resultieren mit (2.22) bis (2.24) komplexe Minimale Theorien, die den Konfigurationen relativ zu $\mathbf{N} = \{B, S, A, U, K, O, E\}$ Strukturen zugrunde legen, die Common-Cause-Teile aufweisen. Diese Ambiguitäten werden in Kürze diskutiert und können für den Nachweis der Intransitivität vernachlässigt werden (vgl. S. 80 ff.). Denn die Transitivität ist eine konditionale Bedingung, gemäss der die Kausalbeziehung genau dann transitiv wäre, wenn gelten würde: Ist A kausal relevant für B und B kausal relevant für C , dann ist A kausal relevant für C . Bei Common-Cause-Strukturen ist der Antezedenz dieser Implikation nicht erfüllt, weshalb sich die Frage, ob die Kausalbeziehung transitiv ist oder nicht, anhand dieser Strukturen nicht prüfen lässt. Für den Moment ist bezüglich der komplexen Minimalen Theorien (2.22) bis (2.24) einzig die Feststellung wichtig, dass weder S noch $\bar{S}$ in der jeweiligen einfachen Minimalen Theorie von E enthalten sind. Ohne Einschränkung der Gültigkeit des Arguments kann der Fokus auf die komplexe Minimale Theorie (2.25) gelegt werden, die im Sinne der Kausalstruktur aus Abbildung 2.21 zu interpretieren ist: BS sowie U sind zwei Alternativursachen von O , das seinerseits eine Ursache von E ist. Weiter ist neben A die Konjunktion $B\bar{S}$ eine komplexe Ursache von K , das alternativ zu O ebenfalls kausal relevant für E ist.

Damit einem Literal gemäss Definition 2.11 kausale Relevanz für ein anderes Literal zugeschrieben werden kann, muss ersteres unter anderem Teil einer einfachen Minimalen Theorie von letzterem sein. Möchte man also die komplexe Minimale Theorie (2.25) derart als Kausalkette interpretieren, dass U , B und S (bzw. A , B und $\bar{S}$) neben deren direkten kausalen Relevanz für O (bzw. K) über ebendieses Literal O (bzw. K) auch indirekt kausal relevant für E sind, müssen sie auch Teil einer einfachen Minimalen Theorie von E sein. Für B , A und U ist dies tatsächlich der Fall, nicht aber für S und $\bar{S}$. Bei O und K noch Teil von einfachen Minimalen Theorien, sind S und $\bar{S}$ nicht Teil von einfachen Minimalen Theorien von E , obwohl E gemäss obigem Kausalgraph weiter unten in der Kausalkette angesiedelt ist. Teil einer einfachen Minimalen Theorie von E zu sein bedeutet, ein Difference-Maker von E zu sein. Weder S noch $\bar{S}$ sind solche Difference-Maker von E . Denn sofern B gegeben ist, tritt E unabhängig davon auf, ob S oder $\bar{S}$ gegeben ist. Der Einfluss von S beziehungsweise $\bar{S}$ auf E beschränkt sich auf den kausalen Pfad, der jeweils ausgehend von B bis hin zur Hervorbringung von E aktiviert wird: Ist S anwesend, wird E über O verursacht, ist S abwesend, wird E über K verursacht. S und $\bar{S}$ haben quasi die Funktion von Weichenstellern und können als *Switching-Faktorlitterale* oder kurz *Switching-Litterale* bezeichnet werden. Dieser Umstand muss auch graphisch erfasst werden, zumal der Kausalgraph aus Abbildung 2.21 aufgrund der Pfade zwischen S beziehungsweise $\bar{S}$ und E suggeriert, dass S und $\bar{S}$ indirekt kausal relevant für E sind. Zu diesem Zweck werden solche Switching-Litterale durch eine gestrichelte Kreisumrandung ausgezeichnet. Mit dieser Notation wird ausgedrückt, dass sich der kausale Einfluss des entsprechenden Literals

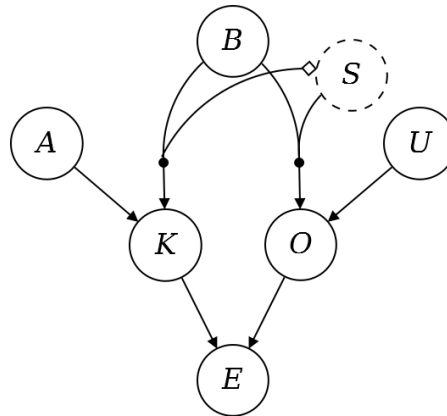


Abbildung 2.22.: Kausaler Prozess zur Senkung der betrieblichen Kosten aus Abbildung 2.21 mit Switching-Literal S beziehungsweise $\bar{S}$.

nicht über alle Folgeglieder der Verkettung erstreckt. Abbildung 2.22 zeigt den Kausalgraphen aus Abbildung 2.21 mit der entsprechenden Anpassung.

Es bleibt an dieser Stelle zu erwähnen, dass diese Notation in Bezug auf ein Switching-Literal im Allgemeinen nicht alle Informationen wiedergibt, die auch in den dazugehörigen Minimalen Theorien enthalten sind. Konkret zeigt diese graphische Notation, im Gegensatz zu den entsprechenden Minimalen Theorien, nicht an, bis zu welchem Folgeglied der Verkettung sich die kausale Relevanz des Switching-Literals fortpflanzt. Anders gesagt: Es ist nicht zwangsläufig so, dass ein Switching-Literal jeweils nur ein Difference-Maker für den direkten Nachfolger in der Kette ist und es diese Qualität für alle darauffolgenden Glieder verliert. Die Difference-Maker-Qualität eines Switching-Literals erstreckt sich zwar jeweils nicht über alle Glieder einer Kette, kann je nach Kausalstruktur aber über mehrere Kettenglieder erhalten bleiben. Ferner kann es vorkommen, dass das positive Literal eines Faktors ein Switching-Literal ist, das negative aber nicht (oder umgekehrt). Demnach würde durch die obige Notation nicht ersichtlich, welches der beiden Literale ein Switching-Literal ist und welches nicht. Indem solche Spezialfälle in der vorliegenden Arbeit nicht betrachtet werden, wird es jedoch als ausreichend erachtet, ein Switching-Literal alleine durch die gestrichelte Kreismrandung zu kennzeichnen.

Ist S Teil einer einfachen Minimalen Theorie des Literals O , das seinerseits Teil einer einfachen Minimalen Theorie eines weiteren Literals E ist, gilt nicht immer, dass letzteres auch eine einfache Minimale Theorie besitzt, die S enthält. Bei Annahme der Erweiterungsklausel existieren folglich kausale Abhängigkeiten derart, dass S kausal relevant für O , O kausal relevant für E , nicht aber S (indirekt) kausal relevant für E ist. Die über Minimale Theorien definierte Kausalbeziehung ist also intransitiv.

Die Intransitivität ist nicht nur an den hier verwendeten Verursachungsbegriff gebunden. Sie betrifft jede Theorie der Kausalität, deren Grundlage unter anderem ein Difference-Making-Prinzip darstellt, wie jenes aus der vorliegenden Arbeit, gemäss dem jede Ursache ein Difference-Maker ist (vgl. S. 61 bzw. Def. 2.10, S. 47). Dazu gehören beispielsweise probabilistische Theorien oder Interventionstheorien. Denn bezogen auf das vereinfachte Beispiel aus Abbildung 2.21 ändert weder der Erhalt der Subventionen gegenüber dem Ausbleiben der Gelder (und umgekehrt) die Wahrscheinlichkeit von Einsparungen, noch bringt eine Intervention des Staates in Form von Fördergeldvergaben im Vergleich zum Ausbleiben solcher Subventionen eine Änderung bezüglich der Einsparungen des Unternehmens mit sich.²⁹

Definition des Begriffs der Kausalkette unter Berücksichtigung der Intransitivität

Die Intransitivität hat nicht zur Folge, dass auf dem Difference-Making-Prinzip basierende Theorien der Kausalität für die Aufdeckung von indirekten Kausalzusammenhängen ungeeignet sind. Zur Identifikation indirekter Kausalbeziehungen ist jedoch die Verwendung des Kettenschlusses unzulässig. Geht man beim konkreten Beispiel der Kausalkette aus Abbildung 2.19 davon aus, dass A permanenter Teil einer einfachen Minimalen Theorie von C und C permanenter Teil einer einfachen Minimalen Theorie von E ist, kann daraus nicht direkt gefolgert werden, dass A auch (indirekt) kausal relevant für E ist. Um die kausale Relevanz von A für E zu etablieren, muss über die komplexe Minimale Theorie (2.21) von Seite 71 hinaus zusätzlich eine einfache Minimale Theorie von E existieren, von der A Teil ist. Im Fall der Kausalkette aus Abbildung 2.19 liegt mit der einfachen Minimalen Theorie $A + B + D \leftrightarrow E$ eine solche vor und zusammen mit der komplexen Minimalen Theorie $(A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E)$ kann den Konfigurationen die in Abbildung 2.19 dargestellte Kausalkette zugeordnet werden. Diese zusätzliche einfache Minimale Theorie $A + B + D \leftrightarrow E$ wird dabei selbst nicht als weiteres Konjunkt in die komplexe Minimale Theorie (2.21) integriert und lediglich zur Bestätigung dafür herangezogen, dass A auch ein Difference-Maker von E ist. Würde man diese einfache Minimale Theorie in die komplexe Minimale Theorie integrieren, würde eine Konjunktion von RDN-Bikonditionalen resultieren, die zwei RDN-Bikonditionale desselben Literals beinhaltet. Eine solche Konjunktion ist keine Minimale Theorie, da sie Bedingung (c) der Definition 2.16 verletzt.

Obwohl aus komplexen Minimalen Theorien aufgrund der Intransitivität alleine noch keine Kausalketten erschlossen werden können, übernehmen sie eine wichtige Rolle bei der Aufdeckung von Ketten. Die komplexen Minimalen Theorien tragen unter anderem einen wesentlichen Teil zur Bestimmung des Typs der Kausalstruktur bei. Zwischen den beiden komplexen Minimalen

²⁹Mehr zur probabilistischen Theorie findet man z. B. bei Suppes (1970). Zur Interventionstheorie vgl. u. a. Woodward (2003).

Theorien $(A + B \leftrightarrow C) \cdot (A + B + D \leftrightarrow E)$ und $(A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E)$ derselben Konfigurationstabelle gibt es einen strukturellen Unterschied: Erstere beschreibt Regelmässigkeiten im Auftrittsverhalten von Faktoren, die auf eine Common-Cause-Struktur zurückgehen und letztere deutet, sofern sich die indirekten Zusammenhänge bestätigen lassen, auf eine Kausalkette hin. Die komplexen Minimalen Theorien besitzen somit eine strukturgebende Funktion und die einfachen Minimalen Theorien, die darüber hinaus aufgrund der Intransitivität herangezogen werden, dienen lediglich zur Bestätigung von indirekten Kausalbeziehungen.

Die Anzahl und Form der einfachen Minimalen Theorien, die bei der Aufdeckung von Kausalketten neben der entsprechenden komplexen Minimalen Theorie zur Bestätigung der indirekten kausalen Relevanz von Literalen herangezogen werden, sind von der Faktormenge sowie der Struktur der aus den Daten abgeleiteten kausalen Beziehungen abhängig. Betrachtet man beispielsweise einen etwas komplexeren Fall einer Kausalkette zwischen vier Literalen Z_1, Z_2, Z_3 und Z_4 , die in dieser Reihenfolge in einer Kausalstruktur enthalten sind, muss für den Schluss auf diese kausale Verkettung notwendigerweise gelten, dass es (a1) eine komplexe Minimale Theorie gibt, die von jedem $Z_i, 2 \leq i \leq 4$, eine einfache Minimale Theorie enthält, die Z_{i-1} aufweist. Darüber hinaus muss aufgrund der Intransitivität (a2) Z_1 Teil einer einfachen Minimalen Theorie von Z_3 sowie Z_4 und Z_2 Teil einer einfachen Minimalen Theorie von Z_4 sein. Für die aus vier Literalen bestehende Kausalkette sind somit mehr als eine einfache Minimale Theorie involviert, welche über die komplexe Minimale Theorie hinaus zur Bestätigung der indirekten Kausalzusammenhänge herangezogen werden müssen.

Schliesslich hat gemäss Definition 2.11 der kausalen Relevanz (b) die Erweiterungsklausel zu gelten, wonach die Regelmässigkeiten aus (a1) und (a2) bei geeigneten Erweiterungen der Faktormenge bestehen bleiben müssen. Zu beachten gilt es hierbei, dass sich via Erweiterung $\mathbf{F}'$ einer Faktormenge $\mathbf{F}$ Literale hinzugenommener Faktoren aus $\mathbf{F}' \setminus \mathbf{F}$ möglicherweise als zusätzliche Zwischenglieder in der relativ zu $\mathbf{F}$ bestehenden Kette positionieren lassen. Demnach kann es vorkommen, dass sich im Rahmen einer Analyse relativ zu $\mathbf{F}$ eine Kettenstruktur ableiten lässt, die sich via Erweiterungen dahingehend verändert, dass sie durch zusätzliche Literale ergänzt wird, die sich zwischen die Kettenglieder der relativ zu $\mathbf{F}$ bereits abgeleiteten Kettenstruktur einbinden lassen. Sofern sich die neuen, nach Erweiterungen formulierten Kausalhypothesen nur durch die Integration weiterer Faktoren von jener unterscheiden, die unter $\mathbf{F}$ formuliert wurde, ist letztere immer noch wahr und unterscheidet sich lediglich im Detaillierungsgrad von ersterer. Um dieser möglichen Veränderung von Kausalhypothesen über Kettenstrukturen, die sich durch Faktorerweiterungen ergeben kann, in der Definition von Kausalketten Rechnung zu tragen, muss die Erweiterungsklausel etwas angepasst werden. So wird nicht gefordert, dass eine die obigen beiden Bedingungen (a1) und (a2) erfüllende Folge $(Z_1, \dots, Z_m)$ von Literalen über alle Erweiterungen hinweg genau so bestehen bleibt, sondern lediglich verlangt, dass $(Z_1, \dots, Z_m)$ je-

weils als Teilfolge einer die Bedingungen (a1) und (a2) erfüllenden Folge auftritt. Dabei ist eine Teilfolge einer Folge ebenfalls eine Folge, bei der die gleichen oder weniger Elemente berücksichtigt werden. Dieser Umstand soll nochmals konkret anhand der Kausalkette (Z_1, Z_2, Z_3, Z_4) betrachtet werden. Relativ zu einer Faktormenge $\mathbf{F}$, die die Faktoren Z_1, Z_2 und Z_4 enthält, nicht aber Z_3 , resultiert aus allen dazugehörigen empirisch möglichen Konfigurationen über $\mathbf{F}$ eine komplexe Minimale Theorie Ψ_1 , die eine die obigen beiden Bedingungen (a1) und (a2) erfüllende Folge von Literalen (Z_1, Z_2, Z_4) aufweist. Wird die gleiche Kausalkette relativ zur erweiterten Faktormenge $\mathbf{F}' = \mathbf{F} \cup \{Z_3\}$ betrachtet, lässt sich aus den entsprechenden Konfigurationen eine komplexe Minimale Theorie Ψ_2 bilden, die eine einfache Minimale Theorie mehr enthält und die obigen beiden Bedingungen (a1) und (a2) erfüllende Folge (Z_1, Z_2, Z_3, Z_4) aufweist. Alleine bezogen auf diese komplexe Minimale Theorie Ψ_2 erfüllt die Folge (Z_1, Z_2, Z_4) die obige Bedingung (a1) nicht, zumal Ψ_2 keine einfache Minimale Theorie von Z_4 aufweist, die Z_2 enthält. Die Folge (Z_1, Z_2, Z_4) ist aber eine Teilfolge von (Z_1, Z_2, Z_3, Z_4) , bei der zwar Z_3 nicht berücksichtigt wird, die Reihenfolge der restlichen Glieder ansonsten aber identisch ist.

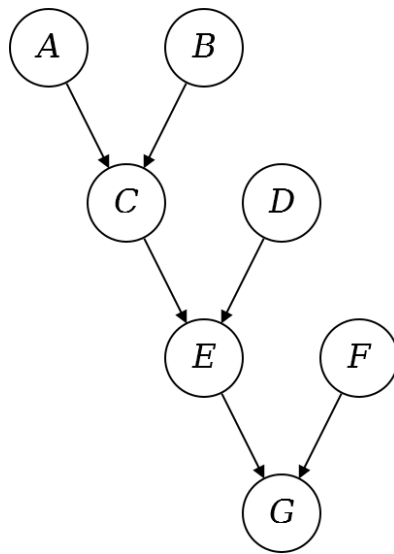
Die drei Bedingungen (a1), (a2) und (b) können so verallgemeinert formuliert werden, dass sich der Begriff der Kausalkette mit einer beliebigen Anzahl von Zwischengliedern definieren lässt:

Definition 2.19 (Kausalkette)

Eine Folge von Literalen $(Z_1, \dots, Z_m)$, $m \geq 3$, bildet in dieser Reihenfolge genau dann eine kausale Verkettung, wenn es eine die Faktoren $Z_1, \dots, Z_m$ enthaltende Faktormenge $\mathbf{F}$ derart gibt, dass

- (a) relativ zu allen empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}$ eine aus Literalen bestehende Folge $(Z_1, \dots, Z_m)$ mit folgenden Eigenschaften existiert:
 - (a1) Es gibt eine komplexe Minimale Theorie Ψ , die von jedem Z_i der Folge, $i \geq 2$, eine einfache Minimale Theorie enthält, von der Z_{i-1} Teil ist,
 - (a2) über Ψ hinaus gilt zusätzlich für jedes j und jedes k , $1 \leq j < j+1 < k \leq m$, dass Z_j Teil einer einfachen Minimalen Theorie von Z_k ist und
- (b) für jede geeignete Erweiterung $\mathbf{F}' \supset \mathbf{F}$ der Faktormenge $\mathbf{F}$ eine aus Literalen von Faktoren aus $\mathbf{F}'$ bestehende Folge existiert, die relativ zu $\mathbf{F}'$ (a) erfüllt und $(Z_1, \dots, Z_m)$ als Teilfolge enthält.

Eine Kausalstruktur, die mindestens eine kausale Verkettung aufweist, ist eine Kausalkette.



$\mathcal{K}_{2,23}$	A	B	C	D	E	F	G
c_1	1	1	1	1	1	1	1
c_2	1	1	1	1	1	0	1
c_3	1	1	1	0	1	1	1
c_4	1	1	1	0	1	0	1
c_5	1	0	1	1	1	1	1
c_6	1	0	1	1	1	0	1
c_7	1	0	1	0	1	1	1
c_8	1	0	1	0	1	0	1
c_9	0	1	1	1	1	1	1
c_{10}	0	1	1	1	1	0	1
c_{11}	0	1	1	0	1	1	1
c_{12}	0	1	1	0	1	0	1
c_{13}	0	0	0	1	1	1	1
c_{14}	0	0	0	1	1	0	1
c_{15}	0	0	0	0	0	1	1
c_{16}	0	0	0	0	0	0	0

Abbildung 2.23.: Kausalkette mit dazugehöriger Konfigurationstabelle $\mathcal{K}_{2,23}$ über $\mathbf{O} = \{A, B, C, D, E, F, G\}$.

Die Erweiterungsklausel (b), die bereits bei der Definition 2.11 der kausalen Relevanz motiviert wurde, zeigt ihre Tragweite im Zusammenhang mit Kausalketten nochmals auf besondere Weise. Ihre Sonderrolle begründet sich im mittlerweile schon mehrfach erwähnten Kettenproblem, das im Folgenden nun diskutiert werden soll.

Das Kettenproblem

Das Beispiel aus Abbildung 2.19 hat gezeigt, dass die zu dieser Kausalkette gehörigen Konfigurationen $\mathcal{K}_{2,19}$ über $\mathbf{K} = \{A, B, C, D, E\}$ genau dieselben sind wie jene, die von der Common-Cause-Struktur aus Abbildung 2.20 generiert werden (vgl. S. 70 f.). Auch die im Vergleich zur Kausalkette aus Abbildung 2.19 um ein Kettenglied verlängerte Kette aus Abbildung 2.23 ist nicht die einzige kausal wohlgeformte Struktur, welche die entsprechenden Konfigurationen $\mathcal{K}_{2,23}$ über $\mathbf{O} = \{A, B, C, D, E, F, G\}$ generiert.

Insgesamt gibt es mit den Strukturen aus Abbildung 2.24 in diesem Fall sogar fünf weitere kausal wohlgeformte Strukturen, die $\mathcal{K}_{2,23}$ ebenfalls generieren. Auch hier unterscheiden sich die

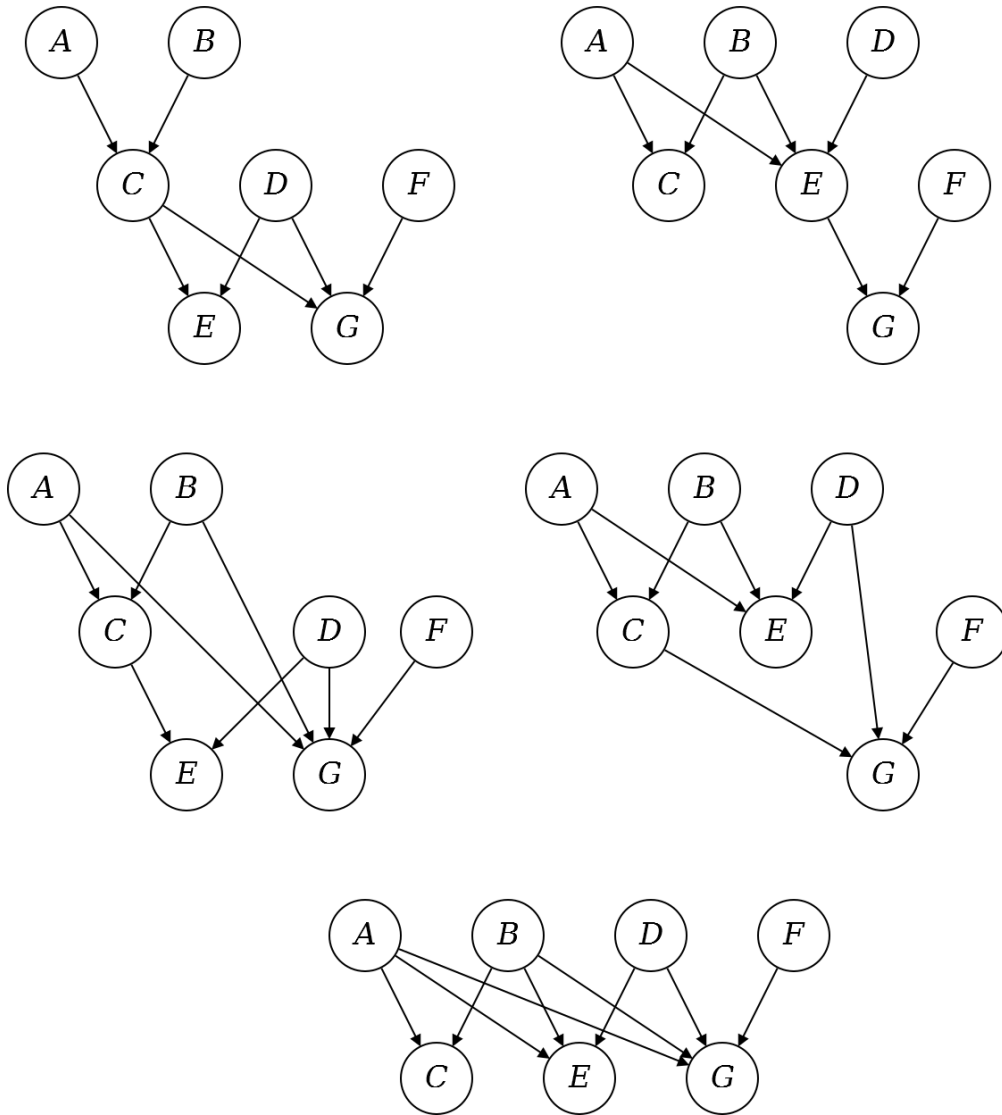


Abbildung 2.24.: Fünf kausal wohlgeformte Strukturen mit Common-Causes, die $\mathcal{K}_{2,23}$ ebenfalls generieren.

Strukturen aus Abbildung 2.24 derart von der Kette aus Abbildung 2.23, dass erstere Common-Cause-Strukturen aufweisen. Für jede dieser zur Kausalkette empirisch äquivalenten Strukturen gilt, dass mindestens drei in der Kette direkt oder indirekt kausal aufeinanderfolgende Literale $(..., Z_i, ..., Z_j, ..., Z_k, ...)$, $i < j < k$, derart als Common-Cause-Struktur angeordnet sind, dass Z_i eine gemeinsame (direkte oder indirekte) Ursache von Z_j und Z_k ist, wobei Z_j und Z_k untereinander nicht kausal verknüpft sind. Die kausale Unabhängigkeit zwischen Z_j und Z_k führt dabei deshalb nicht zu Konfigurationen, die sich von jenen aus $\mathcal{K}_{2.23}$ unterscheiden, weil Z_j und Z_k relativ zur Kette nicht andere Ursachen aufweisen, sondern lediglich direkte durch indirekte Ursachen substituiert werden.

In kausal wohlgeformten Strukturen kann das gleiche Instantiierungsverhalten einer Wirkung durch direkte oder indirekte Ursachen ausgedrückt werden. Das führt allgemein dazu, dass immer, wenn Konfigurationen von Faktoren $Z_1, Z_2, ..., Z_m$, $m \geq 3$, als Verkettungen $(Z_1, Z_2, ..., Z_m)$ modelliert, sie auch als Strukturen modelliert werden können, die Common-Cause-Strukturen aufweisen. Das ist das Kettenproblem, das kurz und kompakt ausgedrückt besagt: Zu jeder Kausalkette gibt es eine oder mehrere empirisch äquivalente Common-Cause-Strukturen.

Das Kettenproblem wird auch von den Definitionen 2.11 bis 2.13 sowie 2.17 bis 2.19 eingefangen. Diese Definitionen folgen dem allgemeinen Schema, dass es eine notwendige Bedingung (a) gibt, die zusammen mit der Erweiterungsklausel (b) hinreichend ist, um Regelmässigkeiten im An- und Abwesenheitsverhalten von Faktoren kausal zu interpretieren.³⁰ Während die Bedingung (a) jeweils Kriterien bestimmt, die es relativ zu einer Faktormenge zu erfüllen gilt, fordert die Bedingung (b), dass diese Kriterien auch für jede geeignete Erweiterung dieser Faktormenge gelten. Der Bedingung (a) zufolge müssen die potentiellen Ursachen Teil von einfachen Minimalen Theorien der Wirkungen sein. Vergleicht man diese Bedingung (a) der Definition 2.19 für Kausalketten mit den Bedingungen aus Definition 2.18, die relativ zu einer Faktormenge notwendig für den Schluss auf Common-Cause-Strukturen sind, sieht man, dass immer wenn erstere erfüllt ist, auch letztere erfüllt sind. Man betrachte dazu nochmals das konkrete Beispiel der kausalen Verkettung (A, C, E) aus Abbildung 2.19 (vgl. S. 70). Relativ zur Faktormenge $\mathbf{K} = \{A, B, C, D, E\}$ muss es gemäss Definition 2.19

- (a1) eine komplexe Minimale Theorie Ψ geben, die eine A enthaltende einfache Minimale Theorie von C sowie eine C enthaltende einfache Minimale Theorie von E aufweist und
- (a2) über Ψ hinaus zusätzlich eine einfache Minimale Theorie von E geben, von der A Teil ist.

³⁰Bedingung (c) der Definition 2.18 von Common-Cause-Strukturen bildet hier insofern eine Ausnahme, als diese Bedingung im Gegensatz zu den Bedingungen (a) und (b) dem Zweck dient, kausale Abhängigkeiten zwischen Literalen zu verhindern.

Konkret handelt es sich um folgende Minimale Theorien (vgl. S. 71):

$$(a1) (A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E),$$

$$(a2) A + B + D \leftrightarrow E.$$

Sind die beiden Bedingungen (a1) und (a2) erfüllt, gehen aus den Konfigurationen also die folgenden drei einfachen Minimalen Theorien hervor, wobei bei jener, die in Anbetracht der dahinterliegenden Verkettung aufgrund der Intransitivität lediglich zur Bestätigung dafür herangezogen wird, dass A auch ein Difference-Maker von E ist, der Doppelpfeil „ $\leftrightarrow$ “ mit einem für „indirekt“ stehenden „ $\overset{i}{\leftrightarrow}$ “ versehen wird:

$$A + B \leftrightarrow C \tag{2.26}$$

$$C + D \leftrightarrow E \tag{2.27}$$

$$A + B + D \overset{i}{\leftrightarrow} E \tag{2.28}$$

Welche einfache Minimale Theorien im Allgemeinen zusätzlich mit einem „ $\overset{i}{\leftrightarrow}$ “ versehen werden, wird jeweils relativ zu einer auf eine kausale Verkettung hindeutenden Minimalen Theorie bestimmt. Seien $\beta_1, \dots, \beta_n$, $n \geq 3$, einfache Minimale Theorien relativ zu einer Konfigurationstabelle $\mathcal{K}_F$ und Ψ eine Minimale Theorie von $\mathcal{K}_F$, die Bedingung (a1) der Definition 2.19 erfüllt. Relativ zu Ψ wird jede in Ψ nicht enthaltene einfache Minimale Theorie β_i , $1 \leq i \leq n$, mit einem „ $\overset{i}{\leftrightarrow}$ “ ausgestattet, die der Bedingung (a2) der Definition 2.19 genügt.

Aus funktionaler Sicht gibt es zwischen „ $\leftrightarrow$ “ und „ $\overset{i}{\leftrightarrow}$ “ keinen Unterschied und das Bikonditional (2.28) ist genauso eine einfache Minimale Theorie von E wie auch (2.27) eine einfache Minimale Theorie von E ist. Indem (2.28) ebenfalls eine einfache Minimale Theorie ist, muss sie gemäss Definition 2.9 Teil einer Minimalen Theorie der zur Kette gehörigen Konfigurationen sein. Da weiter gemäss Bedingung (c) der Definition einer Minimalen Theorie keine zwei einfache Minimale Theorien von Literalen desselben Faktors in einer Minimalen Theorie enthalten sein können sowie (2.27) und (2.28) beides einfache Minimale Theorien von E sind, muss es eine weitere Minimale Theorie geben, die sich von der auf die tatsächliche Kette hindeutenden Minimalen Theorie unterscheidet. Diese Minimale Theorie weist neben (2.28) auch (2.26) auf. Aus der konjunktiven Verknüpfung dieser beiden einfachen Minimalen Theorien resultiert die Minimale Theorie (2.20), die logisch äquivalent zu (2.21) ist. Aufgrund der logischen Äquivalenz wird die entsprechende Struktur, die aus der kausalen Interpretation von (2.20) resultiert, genau dieselben Konfigurationen generieren wie jene, die aus der kausalen Interpretation von (2.21) folgt, das heisst, die beiden Strukturen sind empirisch äquivalent. Im Unterschied zur komplexen Minimalen Theorie (2.21), die sich aus (2.26) und (2.27) zusammensetzt, wird man bei (2.20) die Regularitäten aber auf die Common-Cause-Struktur aus Abbildung 2.20 zurück-

führen und nicht auf die Kausalkette aus Abbildung 2.19. Mit (2.27) bleibt zwar eine einfache Minimale Theorie von E vorhanden, die C enthält und damit Bedingung (c) der Definition 2.18 nicht stützt. Diese Bedingung ist dadurch jedoch auch nicht verletzt, zumal sie ausschliesst, dass C permanenter Teil einer einfachen Minimalen Theorie von E ist, aber durchaus erlaubt, dass C relativ zu einer begrenzten Anzahl geeigneter Faktormengen Teil einer einfachen Minimalen Theorie von E sein kann. Relativ zu einer Faktormenge kann nicht beurteilt werden, ob die Bedingung (c) erfüllt ist oder nicht und sowohl einfache Minimale Theorien, die die Bedingung stützen, als auch solche, die dies nicht tun, sind möglich.

Es zeigt sich also insgesamt, dass die Definitionen das konkrete, anhand des Beispiels aus Abbildung 2.19 veranschaulichte Kettenproblem abbilden. Dies gilt auch für die Verallgemeinerung des Kettenproblems auf kausale Verkettungen beliebiger Länge. Sei $(Z_1, Z_2, \dots, Z_m)$, $m \geq 3$, eine Folge von Literalen von Faktoren einer Menge $\mathbf{F}$, die in dieser Reihenfolge kausal miteinander verkettet sind. Um folgern zu können, dass den Konfigurationen eine Kausalstruktur mit der kausalen Verkettung $(Z_1, Z_2, \dots, Z_m)$ zugrunde liegt, muss es relativ zu $\mathbf{F}$ Minimale Theorien geben, die der Bedingung (a) aus Definition 2.19 für Kausalketten genügen. Gibt es diese Minimalen Theorien, gehen diese gezwungenermassen aus folgendem System $\Delta_{\mathbf{F}}$ von einfachen Minimalen Theorien hervor, das aus den Konfigurationen gefolgert wird:

$$\begin{aligned}
 &Z_1 \dots \leftrightarrow Z_2 \\
 &Z_2 \dots \leftrightarrow Z_3 \\
 &Z_1 \dots \overset{i}{\leftrightarrow} Z_3 \\
 &\vdots \\
 &Z_{m-1} \dots \leftrightarrow Z_m \\
 &Z_1 \dots \overset{i}{\leftrightarrow} Z_m \\
 &Z_2 \dots \overset{i}{\leftrightarrow} Z_m \\
 &\vdots \\
 &Z_{m-2} \dots \overset{i}{\leftrightarrow} Z_m
 \end{aligned}$$

Zur besseren Veranschaulichung wurde darauf verzichtet, jede der einfachen Minimalen Theorien der Vollständigkeit halber mit Ko-Literalen X_i , $i \in \mathbb{N}$, und Alternativursachen $\mathcal{Y}_j$, $j \in \mathbb{N}$, zu ergänzen. Diese wurden stattdessen durch Punkte ‘...’ ersetzt. Ferner wurden wiederum jene einfachen Minimale Theorien mit einem ‘ i ’ versehen, die in Anbetracht der tatsächlichen kausalen Verkettung $(Z_1, Z_2, \dots, Z_m)$ aufgrund der Intransitivität zur Erfüllung der Bedingung (a2) herangezogen werden, selbst aber nicht Teil der entsprechenden komplexen Minimalen Theorie

aus Bedingung (a1) sind. Vom k -ten Literal der Verkettung $(Z_1, Z_2, \dots, Z_k, \dots, Z_m)$ sind $k - 1$ einfache Minimale Theorien in Δ_F enthalten. Es können nicht weniger sein, da dies bedeuten würde, dass mindestens zwei in der Verkettung vor Z_k positionierte Literale Z_i und Z_j nur in derselben einfachen Minimalen Theorie von Z_k auftreten oder mindestens ein vor Z_k positioniertes Literal in keiner einfachen Minimalen Theorie von Z_k enthalten ist. Letzteres würde heissen, dass dieses Literal kein Difference-Maker von Z_k ist, was angesichts der dahinterliegenden Verkettung und der Definition 2.11 der kausalen Relevanz einen Widerspruch zur Folge hätte. Dasselbe gilt für ersteres: In Anbetracht der kausalen Verkettung $(Z_1, Z_2, \dots, Z_i, \dots, Z_j, \dots, Z_k, \dots, Z_m)$ wäre eines der beiden in der Struktur übereinanderliegenden Literale Z_i und Z_j redundant. Denn eine Variation von Z_i , die zu einer Variation von Z_k führt, führt aufgrund der Verkettung ebenfalls zu einer Variation des Zwischengliedes Z_j , womit es kein Difference-Making-Paar von Z_i (bzw. Z_j) bezüglich Z_k geben kann. Daraus wiederum folgt, dass es keine Z_i und Z_j enthaltende einfache Minimale Theorie von Z_k gibt.

Aus diesem für den Schluss auf die Verkettung $(Z_1, Z_2, \dots, Z_m)$ notwendigen System Δ_F von einfachen Minimalen Theorien lassen sich relativ zu F erneut komplexe Minimale Theorien bilden, die darauf schliessen lassen, dass den Konfigurationen Common-Cause-Strukturen zugrunde liegen können. Im Extremfall lässt sich eine komplexe Minimale Theorie bilden, die auf überhaupt keine Verkettung $(Z_1, Z_2, \dots, Z_m)$ oder Teile davon hindeutet. Dazu werden von jedem Literal aus der Verkettung jene einfachen Minimalen Theorien aus Δ_F konjunktiv miteinander verknüpft, die das erste Literal Z_1 aufweisen:

$$(Z_1 \dots \leftrightarrow Z_2) \cdot (Z_1 \dots \leftrightarrow Z_3) \cdot (Z_1 \dots \leftrightarrow Z_4) \cdots (Z_1 \dots \leftrightarrow Z_m) \quad (2.29)$$

Wiederum stützen die restlichen einfachen Minimalen Theorien des Systems Δ_F zwar die Bedingung (c) der Definition 2.18 von Common-Cause-Strukturen nicht, sie widerlegen sie aber aufgrund der Relativierung auf eine Faktormenge auch nicht. (2.29) ist empirisch äquivalent zu jener komplexen Minimalen Theorie, die auf die Verkettung $(Z_1, Z_2, \dots, Z_m)$ hindeutet, da (2.29) jedes Z_i , $1 \leq i \leq m$, der Folge $(Z_1, Z_2, \dots, Z_m)$ aufweist und eine einfache Minimale Theorie von jedem Z_j , $j \geq 2$, enthält. Die mit der komplexen Minimalen Theorie (2.29) einhergehende, auf Regelmässigkeiten basierende Erklärung der An- und Abwesenheiten der Faktoren umfasst somit denselben Teil der Konfigurationen, der von der auf eine Kausalkette hindeutenden Minimalen Theorie erfasst wird.

Neben diesem Extremfall kann den Konfigurationen relativ zu F eine Vielzahl an Mischformen zugrunde gelegt werden, die sowohl Common-Cause-Strukturen als auch Verkettungen aufweisen. Eine entsprechende komplexe Minimale Theorie wäre etwa jene, die im Unterschied zur komplexen Minimalen Theorie (2.29), anstelle der einfachen Minimalen Theorie $Z_1 \dots \leftrightarrow Z_3$, die

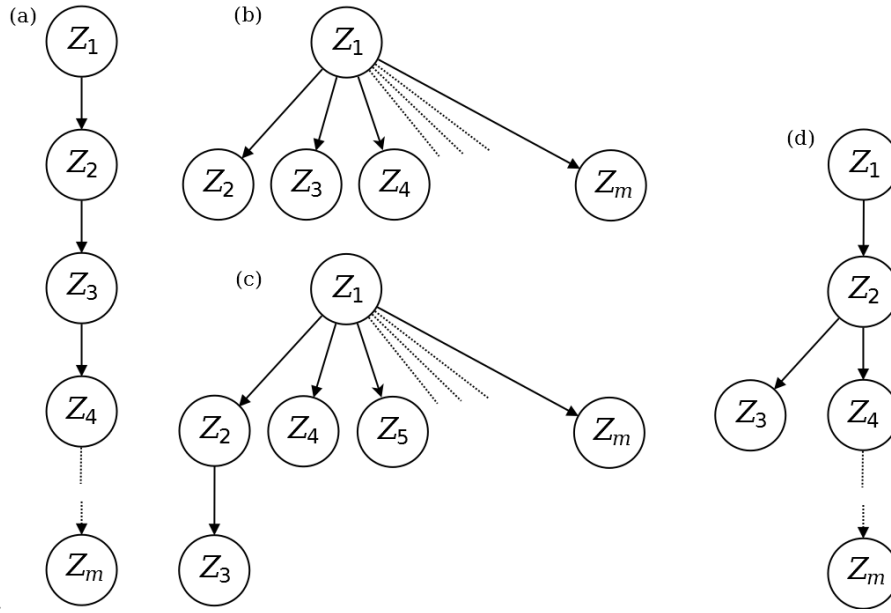


Abbildung 2.25.: Beispiele von Strukturen, die gemäss dem allgemeinen Kettenproblem relativ zu einer Faktormenge $\{Z_1, Z_2, \dots, Z_m\}$ dieselben Konfigurationen aufweisen wie die Kausalkette (a). Die gestrichelten Linien der Strukturen (b) und (c) deuten die direkte kausale Relevanz von Z_1 für die restlichen, im Kausalgraphen nicht dargestellten Literale aus $\{Z_1, Z_2, \dots, Z_m\}$ dar. Bei den Graphen (a) und (d) repräsentiert die gestrichelte und unterbrochene Linie zwischen Z_4 und Z_m eine Zusammenfassung der Verkettung $(Z_4, Z_5, Z_6, \dots, Z_m)$.

einfache Minimale Theorie $Z_2 \dots \leftrightarrow Z_3$ aufweist und $Z_1 \dots \leftrightarrow Z_3$ aufgrund der Intransitivität zur Bestätigung dafür herangezogen wird, dass Z_1 auch ein Difference-Maker von Z_3 ist:

$$(Z_1 \dots \leftrightarrow Z_2) \cdot (Z_2 \dots \leftrightarrow Z_3) \cdot (Z_1 \dots \leftrightarrow Z_4) \cdot (Z_1 \dots \leftrightarrow Z_5) \cdots (Z_1 \dots \leftrightarrow Z_m) \quad (2.30)$$

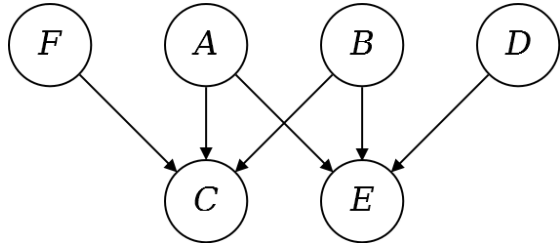
Zur Bildung einer komplexen Minimalen Theorie aus Δ_F , die das Auftrettsverhalten der Wirkungen aus der tatsächlichen Kausalkette einfängt, kann für jedes Z_i eine von insgesamt $i - 1$ einfachen Minimalen Theorien mit einfachen Minimalen Theorien der restlichen Literale konjunktiv verknüpft werden. Allgemein können somit bei einer Verkettung mit n aufeinanderfolgenden Wirkungen aus den entsprechenden Konfigurationen mindestens $n!$ komplexe Minimale Theorien abgeleitet werden. Diese lassen neben einer Verkettung $(Z_1, Z_2, \dots, Z_m)$ auf $m! - 1$ weitere mögliche Kausalstrukturen schliessen, die relativ zur Faktormenge den Konfigurationen zugrunde liegen können. Abbildung 2.25 zeigt neben der kausalen Verkettung (a) sowie den beiden Strukturen (b) und (c), die über (2.29) und über (2.30) mit $Z_1 \dots \overset{i}{\leftrightarrow} Z_3$ folgen, eine weitere dieser $m! - 1$ möglichen Strukturen (d).

Das Kettenproblem und die Erweiterungsklausel

Während relativ zu einer Faktormenge $\mathbf{F}$ aus einer Konfigurationstabelle über $\mathbf{F}$ nicht eindeutig auf die Kette geschlossen werden kann und mit der Kette auch empirisch äquivalente Common-Cause-Strukturen einhergehen, löst sich dieses Kettenproblem nach Erhärtung der Erweiterungsklausel auf. Denn aus einem System von einfachen Minimalen Theorien wie obigem $\Delta_{\mathbf{F}}$, das der Bedingung (a) der Definition 2.19 genügt und für das zusätzlich gilt, dass es diese Bedingung auch für jede Erweiterung der Faktormenge erfüllt, kann eindeutig auf die kausale Verkettung $(Z_1, Z_2, \dots, Z_m)$ geschlossen werden. Gelten die einfachen Minimalen Theorien aus $\Delta_{\mathbf{F}}$ auch für jede Erweiterung der Faktormenge, gilt gemäss Definition 2.11 für jede darin enthaltene einfache Minimale Theorie $Z_i \dots \leftrightarrow Z_j$, $i < j \leq m$, dass Z_i kausal relevant für Z_j ist. Die Gesamtheit dieser kausalen Abhängigkeiten ist mit einer kausalen Verkettung $(Z_1, Z_2, \dots, Z_m)$ vereinbar, nicht aber mit einer Common-Cause-Struktur, zumal sich letztere unter anderem dadurch auszeichnet, dass die darin enthaltenen Wirkungen nicht kausal voneinander abhängig sind.

Dies soll nochmals konkret anhand des Beispiels aus Abbildung 2.19 (vgl. S. 70) veranschaulicht werden, wobei der Einfachheit halber davon ausgegangen wird, dass die den Konfigurationen zugrundeliegende Kette allumfassend ist. Letzteres bedeutet, dass auch bei jeder Erweiterung der Faktormenge die Konfigurationen der Faktoren aus $\mathbf{K}$ dieselben bleiben. Dementsprechend resultieren auch bei jeder Erweiterung die drei einfachen Minimalen Theorien (2.26), (2.27) und (2.28), woraus sich ein permanentes, mit dem Index „Kette“ gekennzeichnetes System $[(2.26), (2.27), (2.28)]_{\text{Kette}}$ ergibt. Die Regularität zwischen C und E aus (2.27) kann in diesem Fall gemäss Definition 2.11 nicht mehr kausal uninterpretiert bleiben und aufgrund des Umstandes, dass C permanenter Teil einer einfachen Minimalen Theorie von E ist, wird C kausale Relevanz für E zugesprochen. Das permanente System $[(2.26), (2.27), (2.28)]_{\text{Kette}}$ kann damit nur auf eine Kausalkette zurückgeführt werden.

Würde es sich stattdessen tatsächlich um eine Common-Cause-Struktur handeln, wäre C nicht kausal relevant für E . Dementsprechend würde sich über alle Erweiterungen hinweg nicht das System $[(2.26), (2.27), (2.28)]_{\text{Kette}}$, sondern ein System herauskristallisieren, das lediglich eine einfache Minimale Theorie von C sowie eine einfache Minimale Theorie von E , von der C nicht Teil ist, enthält. Ein entsprechendes Szenario für diesen zweiten Fall könnte beispielsweise derart aussehen, dass die tatsächliche Kausalstruktur jene aus Abbildung 2.26 ist und eine relativ zu $\mathbf{K}$ latente Alternativursache F von C existiert, die relativ zu $\mathbf{K}$ konstant abwesend war. Dementsprechend erweist sich die Konfigurationstabelle $\mathcal{K}_{2.19}$ nur als Teil aller tatsächlich empirisch möglichen Konfigurationen über $\mathbf{K} = \{A, B, C, D, E\}$, die über $\mathbf{P} = \{A, B, C, D, E, F\} \supset \mathbf{K}$ von der Gestalt aus $\mathcal{K}_{2.26}$ sind. Der Einfachheit halber soll erneut davon ausgegangen werden, dass diese um F ergänzte Common-Cause-Struktur in dem Sinne allumfassend ist, dass keine



$\mathcal{K}_{2.26}$	A	B	C	D	E	F
c_1	1	1	1	1	1	1
c_2	1	1	1	1	1	0
c_3	1	1	1	0	1	1
c_4	1	1	1	0	1	0
c_5	1	0	1	1	1	1
c_6	1	0	1	1	1	0
c_7	1	0	1	0	1	1
c_8	1	0	1	0	1	0
c_9	0	1	1	1	1	1
c_{10}	0	1	1	1	1	0
c_{11}	0	1	1	0	1	1
c_{12}	0	1	1	0	1	0
c_{13}	0	0	1	1	1	1
c_{14}	0	0	1	0	0	1
c_{15}	0	0	0	1	1	0
c_{16}	0	0	0	0	0	0

Abbildung 2.26.: Um F ergänzte Common-Cause-Struktur aus Abbildung 2.20 zusammen mit den entsprechenden Konfigurationen über $\mathbf{P} = \{A, B, C, D, E, F\}$.

weiteren Literale in einer nicht über F , A , B und D führenden kausalen Abhängigkeit mit den Wirkungsfaktoren dieser Struktur stehen.

Im Gegensatz zu $\mathcal{K}_{2.19}$ folgt aus der kompletten Konfigurationstabelle $\mathcal{K}_{2.26}$ genau eine (komplexe) Minimale Theorie $(A + B + F \leftrightarrow C) \cdot (A + B + D \leftrightarrow E)$ und damit genau die zwei folgenden einfachen Minimalen Theorien:

$$A + B + F \leftrightarrow C \quad (2.31)$$

$$A + B + D \leftrightarrow E \quad (2.32)$$

Die Hinzunahme des Faktors F bewirkt, dass Variationen im Auftrittsverhalten des Faktors C von Variationen im Auftrittsverhalten des Faktors E isoliert werden und so die in $\mathcal{K}_{2.19}$ noch vorhandene Regularität zwischen C und E aufgebrochen wird. Das daraus resultierende System von einfachen Minimalen Theorien weist keine einfache Minimale Theorie von E auf, die C enthält. Ferner ist mit der Voraussetzung, dass die Common-Cause-Struktur allumfassend ist, das aus (2.31) und (2.32) bestehende, mit dem für „Common-Cause“ stehenden Index „CC“ gekennzeichnete System $[(2.31), (2.32)]_{CC}$ invariant gegenüber Faktorerweiterungen. $[(2.31), (2.32)]_{CC}$ ist also ein permanentes System und erfüllt damit Bedingung (a), (b) sowie insbesondere auch

(c) der Definition 2.18 und kann über die Bildung der komplexen Minimalen Theorie eindeutig als Common-Cause-Struktur kausal interpretiert werden.

Übertragen auf das allgemeine Kettenproblem mit dem System Δ_F von einfachen Minimalen Theorien (vgl. S. 84) bedeutet dies, dass sich die tatsächlich diesem System zugrunde liegende Kausalstruktur über die Bestätigung der Erweiterungsklausel eindeutig identifizieren lässt. Mit der Bestätigung der Erweiterungsklausel stellt sich heraus, welche einfachen Minimalen Theorien oder Teile davon bei geeigneten Erweiterungen bestehen bleiben und damit kausal interpretiert werden. Sind es alle einfachen Minimalen Theorien aus Δ_F , handelt es sich gemäss Definition 2.19 um eine kausale Verkettung $(Z_1, Z_2, \dots, Z_i, \dots, Z_m)$. Denn in diesem Fall sind gemäss Definition 2.11 alle im System Δ_F enthaltenen einfachen Minimalen Theorien kausal zu interpretieren und es ist alleine die kausale Verkettung, die all diesen kausalen Abhängigkeiten Rechnung trägt. Eine Common-Cause-Struktur, wie etwa die Struktur (b) aus Abbildung 2.25, kann ein solches System Δ_F nicht über alle Erweiterungen aufrechterhalten. Die Common-Cause-Struktur würde der Folgerung widersprechen, dass die Wirkungen kausal relevant füreinander sind. Bleiben über alle geeigneten Erweiterungen jedoch tatsächlich nur jene einfachen Minimalen Theorien $Z_i \dots \leftrightarrow Z_j$ aus Δ_F erhalten, für die $i = 1$ und $j \neq 1$ gilt, handelt es sich gemäss Definition 2.18 um die Common-Cause-Struktur (b) aus Abbildung 2.25. Analog zu diesen beiden Extremfällen wird etwa die Kausalstruktur (d) aus Abbildung 2.25 genau dann eingefangen, wenn über alle Erweiterungen jene einfachen Minimalen Theorien $Z_i \dots \leftrightarrow Z_j$ aus Δ_F erhalten bleiben, für die $i < j$ und $i \neq 3$ gilt.

Insgesamt zeigt sich die angesprochene Sonderrolle der Erweiterungsklausel für die Definition 2.19 von Kausalketten derart, dass durch sie die mit dem Kettenproblem einhergehenden Ambiguitäten aufgelöst werden. Es muss jedoch gesagt sein, dass diese Ambiguitäten im Rahmen des kausalen Schliessens immer bestehen. Jede Kausalanalyse wird immer relativ zu einer Faktormenge durchgeführt, weshalb bei der Analyse einer Kausalkette jeweils auch Common-Cause-Strukturen abgeleitet werden, die den entsprechenden Konfigurationen zugrunde gelegt werden können. Wie bei jedem anderen Schluss auf Kausalität auch muss man sich beim Schluss auf eine Kausalkette damit begnügen, dass das Risiko eines Szenarios wie jenes aus Abbildung 2.26 durch Folgeanalysen mit erweiterten Faktormengen zwar kontinuierlich sinkt, aber nicht vollständig verschwindet.

Direkte und indirekte kausale Relevanz

Mithilfe der Definition 2.19 und des darin definierten Begriffs der kausalen Verkettung lässt sich die kausale Relevanz in die direkte und indirekte kausale Relevanz aufspalten. Direkte und

indirekte kausale Relevanz sind dabei explizit an die Faktormenge zu binden: Betrachtet man erneut die Kausalkette aus Abbildung 2.19, ist A relativ zur Faktormenge $\mathbf{K} = \{A, B, C, D, E\}$ indirekt kausal relevant für E . Für die gleiche Kausalstruktur ist A relativ zur Faktormenge $\mathbf{Q} = \{A, B, D, E\} \subset \mathbf{K}$, das heisst ohne Betrachtung des Faktors C , direkt kausal relevant für E . Direkte und indirekte kausale Relevanz werden folgendermassen definiert:

Definition 2.20 (Indirekte kausale Relevanz)

X ist relativ zu einer Faktormenge $\mathbf{F}$ genau dann indirekt kausal relevant für W , wenn es eine aus Literalen von Faktoren aus $\mathbf{F}$ bestehende Folge $(Z_1, Z_2, \dots, Z_m)$ gibt, die in dieser Reihenfolge eine kausale Verkettung bilden sowie $X = Z_k$ und $W = Z_l$, $1 \leq k < k+1 < l \leq m$, gilt.

Definition 2.21 (Direkte kausale Relevanz)

X ist relativ zu einer Faktormenge $\mathbf{F}$ genau dann direkt kausal relevant für W , wenn X kausal relevant für W ist und relativ zu $\mathbf{F}$ mindestens ein kausaler Zusammenhang zwischen X und W nicht indirekt ist.

Es zeigt sich vor allem anhand der Definition 2.21 deutlich, dass nicht ausgeschlossen ist, dass ein Literal innerhalb der gleichen Kausalstruktur relativ zu einer Faktormenge sowohl direkt als auch indirekt kausal relevant für dieselbe Wirkung sein kann.

2.11.3. Feedbackstrukturen

Mit kleinen Adaptionen der Definition 2.19 des Begriffs der Kausalkette lässt sich zum Schluss eine Definition für Wechselwirkungen sowie Zyklen formulieren. Damit ist die in Abschnitt 2.1.3 dargestellte Komplexität von Kausalstrukturen vollständig erfasst. Für Feedbackstrukturen gilt das wesentliche Charakteristikum, dass sie kausale Verkettungen beinhalten, die zyklisch sind. Jedes Literal eines kausalen Zyklus ist kausal relevant für jedes andere Literal dieser geschlossenen Verkettung. Bilden beispielsweise Z_1 , Z_2 sowie Z_3 in dieser Reihenfolge einen kausalen Zyklus, ist Z_1 direkt kausal relevant für Z_2 , Z_2 direkt kausal relevant für Z_3 und Z_3 wiederum direkt kausal relevant für Z_1 . Darüber hinaus ist jedes dieser drei Literale indirekt kausal relevant für seinen direkten Verursacher. Relativ zu einer Faktormenge $\mathbf{F}$, die die drei Faktoren Z_1 , Z_2 sowie Z_3 enthält, wird der Zyklus von einer komplexen Minimalen Theorie Ψ eingefangen, die eine einfache Minimale Theorie von Z_2 enthält, von der Z_1 Teil ist, eine einfache Minimale Theorie von Z_3 enthält, von der Z_2 Teil ist, und eine einfache Minimale Theorie von Z_1 enthält, von der Z_3 Teil ist. Aufgrund der Intransitivität müssen neben Ψ zusätzlich einfache Minimale Theorien existieren, die Z_1 als Difference-Maker von Z_3 , Z_2 als Difference-Maker

von Z_1 sowie Z_3 als Difference-Maker von Z_2 ausweisen. Der in Ψ enthaltene Zyklus zwischen Z_1 , Z_2 und Z_3 kann genau dann kausal interpretiert werden, wenn es auch für jede geeignete Erweiterung von $\mathbf{F}$ Minimale Theorien gibt, die diese Eigenschaften der zyklischen Folge aufweisen oder darin eine zyklische Folge mit analogen Eigenschaften existiert, die Z_1 , Z_2 und Z_3 in dieser Reihenfolge als echte Teilfolge enthält. Letzteres trägt wiederum dem Umstand Rechnung, dass sich durch Erweiterungen allenfalls Literale als zusätzliche Glieder in den Zyklus integrieren lassen (vgl. dazu S. 78). Da Feedbackstrukturen als geschlossene Verkettungen auch vom Kettenproblem betroffen sind, soll des Weiteren auch nochmals auf die Sonderrolle der Erweiterungsklausel aufmerksam gemacht werden (vgl. Abschnitt 2.11.2, S. 87 ff.).

Ausgehend von Definition 2.19 gilt es für die Definition des Begriffs der Feedbackstruktur zu berücksichtigen, dass eine Folge $(Z_1, \dots, Z_m)$, die der Bedingung (a) der Definition 2.19 genügt, bei Feedbackstrukturen zusätzlich zyklisch sein muss, das heißt, auf Z_m muss wieder Z_1 folgen. Bezüglich der Bedingung (a1) wird dies durch die Ergänzung der Folge um das Literal Z_{m+1} erreicht, das identisch mit Z_1 ist. Mithilfe der daraus resultierenden Folge $(Z_1, Z_2, \dots, Z_m, Z_{m+1})$ lässt sich eine komplexe Minimale Theorie Ψ beschreiben, bei der jedes Z_i der Folge, $1 \leq i \leq m$, Teil einer in Ψ enthaltenen einfachen Minimalen Theorie von Z_{i+1} ist, wobei sich darunter aufgrund von $Z_1 = Z_{m+1}$ eine einfache Minimale Theorie von Z_1 befindet, die Z_m aufweist.

Auch bezüglich der Bedingung (a2), die aufgrund der Intransitivität zur Bestätigung der indirekten Kausalzusammenhänge erfüllt sein muss, ist eine Anpassung notwendig. Bei Kausalketten sind Literale indirekt kausal relevant für weiter unten in der Kette positionierte Literale. Im Unterschied dazu gilt bei Feedbackstrukturen aufgrund der geschlossenen Verkettungen, dass die darin enthaltenen Literale, abgesehen von der direkten kausalen Relevanz für den direkten Nachfolger, für jedes andere Literal indirekt kausal relevant sind. Damit ist über die komplexe Minimale Theorie aus Bedingung (a1) hinaus zusätzlich gefordert, dass jedes Z_j , $1 \leq j \leq m$, Teil einer einfachen Minimalen Theorie jedes anderen Z_k , $1 \leq k \leq m$, der zyklischen Verkettung ist. Wie bereits bei der Definition der Kausalkette sollen dabei die Fälle $k = j + 1$ explizit ausgeschlossen werden, um die einfachen Minimalen Theorien, die Teil der komplexen Minimalen Theorie sind, klar von den einfachen Minimalen Theorien, die aufgrund der Intransitivität herangezogen werden, zu trennen. Mit Hinzunahme der Erweiterungsklausel können Feedbackstrukturen insgesamt folgendermassen definiert werden:

Definition 2.22 (Feedbackstruktur)

Eine zyklische Folge von Literalen $(Z_1, \dots, Z_m)$, $m \geq 2$, sodass auf Z_m wiederum Z_1 folgt, bildet in dieser Reihenfolge genau dann eine kausale Feedbackstruktur, wenn es eine die Faktoren $Z_1, \dots, Z_m$ enthaltende Faktormenge $\mathbf{F}$ derart gibt, dass

- (a) relativ zu allen empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}$ eine aus Literalen bestehende Folge $(Z_1, \dots, Z_m, Z_{m+1})$, $Z_{m+1} = Z_1$, mit folgenden Eigenschaften existiert:
 - (a1) Es gibt eine komplexe Minimale Theorie Ψ , die von jedem Z_i der Folge, $i \geq 2$, eine einfache Minimale Theorie enthält, die Z_{i-1} aufweist,
 - (a2) über Ψ hinaus gilt zusätzlich für jedes j , $1 \leq j \leq m$, und jedes k , $1 \leq k \leq m+1$ und $k \neq j+1$, dass Z_j in einer einfachen Minimalen Theorie von Z_k enthalten ist und
- (b) für jede geeignete Erweiterung $\mathbf{F}' \supset \mathbf{F}$ der Faktormenge $\mathbf{F}$ eine aus Literalen aus Faktoren aus $\mathbf{F}'$ bestehende Folge $(Z_1, \dots, Z_{n+1})$, $n \geq m$ und $Z_{n+1} = Z_1$, existiert, die relativ zu $\mathbf{F}'$ (a) erfüllt und $(Z_1, \dots, Z_m)$ als Teilfolge enthält.

Aufgrund der zyklischen Anordnung kann für die Folge $(Z_1, \dots, Z_m, Z_{m+1})$, auf die sich (a) und (b) der Definition 2.22 beziehen, ein beliebiges Literal der Feedbackstruktur als Z_1 beziehungsweise Z_{m+1} definiert werden, sofern die benachbarten Literale jeweils dieselben bleiben. Dies hat lediglich einen Einfluss auf die Anordnung der konjunktiv verknüpften einfachen Minimalen Theorien innerhalb der komplexen Minimalen Theorie Ψ . Da die konjunktive Verknüpfung kommutativ ist, resultiert daraus jeweils dieselbe komplexe Minimale Theorie Ψ .

Zur Veranschaulichung der Definition 2.22 wird die zyklische Teilstruktur aus Abbildung 2.27 zusammen mit der Folge (A, B, C, D, A) betrachtet, wobei $A = Z_1 = Z_5$, $B = Z_2$, $C = Z_3$ sowie $D = Z_4$ gilt.

Die dargestellten kausalen Abhängigkeiten werden genau dann von Definition 2.22 als Feedbackstruktur erfasst, wenn es relativ zu einer Faktormenge $\mathbf{F}$, $\{A, B, C, D\} \subseteq \mathbf{F}$, eine Folge

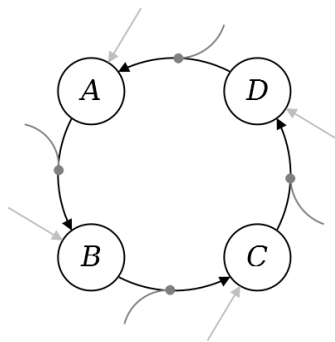


Abbildung 2.27.: Kausale Feedbackstruktur mit einem aus den vier Literalen A , B , C und D bestehenden Zyklus.

(A, B, C, D, A) gibt, sodass gilt: (a1) Es gibt eine komplexe Minimale Theorie der dazugehörigen Konfigurationen über $\mathbf{F}$, die eine einfache Minimale Theorie von $A (= Z_5)$ enthält, von der $D (= Z_4)$ Teil ist, eine einfache Minimale Theorie von $B (= Z_2)$ enthält, von der $A (= Z_1)$ Teil ist, eine einfache Minimale Theorie von $C (= Z_3)$ enthält, von der $B (= Z_2)$ Teil ist, sowie eine einfache Minimale Theorie von $D (= Z_4)$ enthält, von der $C (= Z_3)$ Teil ist. Ferner gibt es (a2) aufgrund der Intransitivität zur Bestätigung der indirekten Zusammenhänge bei der Betrachtung derselben Faktormenge

- eine einfache Minimale Theorie von $A (= Z_1 = Z_5)$, von der $C (= Z_3)$ Teil ist, und eine von A , von der $D (= Z_4)$ Teil ist,
- eine einfache Minimale Theorie von $B (= Z_2)$, von der $D (= Z_4)$ Teil ist, und eine von B , von der $A (= Z_1 = Z_5)$ Teil ist,
- eine einfache Minimale Theorie von $C (= Z_3)$, von der $A (= Z_1 = Z_5)$ Teil ist, und eine von C , von der $B (= Z_2)$ Teil ist, sowie
- eine einfache Minimale Theorie von $D (= Z_4)$, von der $B (= Z_2)$ Teil ist, und eine von D , von der $C (= Z_3)$ Teil ist.

Schliesslich gelten (b) die in (a1) und (a2) erfassten Regularitäten auch bei jeder geeigneten Erweiterung $\mathbf{F}'$ der Faktormenge $\mathbf{F}$. Das heisst, dass es für jede Erweiterung $\mathbf{F}'$ von $\mathbf{F}$ eine zyklische Folge gibt, für die relativ zu den Faktoren in $\mathbf{F}'$ die Bedingungen (a1) und (a2) gelten und (A, B, C, D) als Teilfolge aufweist.

Ein Spezialfall von Feedbackstrukturen bilden die Wechselwirkungen, die sich dadurch auszeichnen, dass A kausal relevant für B ist und B wiederum kausal relevant für A ist. Auch diese Kausalstrukturen werden durch den obigen Begriff der Feedbackstrukturen definiert, wobei für m aus der entsprechenden Definition 2.22 $m = 2$ gilt.

Mit der Gesamtheit der Definitionen dieses und des vorherigen Abschnitts können alle in Abbildung 2.1 dargestellten Typen von Kausalstrukturen über die dazugehörigen Konfigurationen eingefangen werden. Aus diesen Typen lassen sich durch modulare Kombination Kausalstrukturen mit den unterschiedlichsten Graden an Komplexität bilden.

2.12. In Kürze

Zur Entwicklung eines auf Konfigurationen basierenden Verfahrens zur Aufdeckung von kausalen Abhängigkeiten muss in einem ersten Schritt definiert werden, was genau gesucht wird. Mit dem vorliegenden Kapitel zum kausaltheorietischen Rahmen wurde die entsprechende Antwort geliefert. Es hat sich gezeigt, dass nur jene in Konfigurationen auftretenden, auf der Suffizienz und Notwendigkeit basierenden Regelmässigkeiten kausal interpretiert werden können, die von Minimalen Theorien beschrieben werden. Mittels Minimaler Theorien können nicht nur komplexe und alternative Ursachen einer Wirkung erfasst werden, sondern auch Kausalstrukturen mit mehreren Wirkungen, wie Common-Cause-Strukturen, Kausalketten oder Feedbackstrukturen. Die wesentliche Eigenschaft der Minimalen Theorien ist dabei die strikte Redundanzfreiheit, die dazu führt, dass jeder Teil der Minimalen Theorien unverzichtbar für die Beschreibung der Regelmässigkeiten ist. Konkret beschreibt eine Minimale Theorie eine strukturell minimale Konjunktion von RDN-Bikonditionalen. Mit der Forderung der strukturellen Minimalität wird ein Typ von Redundanz ausgeschlossen, der bisher unberücksichtigt blieb.

Eine Minimale Theorie wird immer relativ zu einer Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ über eine Faktormenge $\mathbf{F}$ bestimmt. Aus diesem Grund ist es zwar notwendig, dass Literale Teil einer einfachen Minimalen Theorie einer Wirkung W sind, um als Ursachen von W interpretiert werden zu können, nicht aber hinreichend. Literale können erst dann als kausal relevant für W interpretiert werden, wenn gilt, dass sie relativ zu allen empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}$ Teil einer einfachen Minimalen Theorie von W sind und relativ zu allen empirisch möglichen Konfigurationen der Faktoren jeder geeigneten Erweiterung $\mathbf{F}' \supset \mathbf{F}$ Teil einer einfachen Minimalen Theorie bleiben.

Die über Minimale Theorien definierten kausalen Abhängigkeiten sind intransitiv und genügen fünf fundamentalen kausalen Prinzipien. Letztere können als Geltungsbereich aufgefasst werden, innerhalb dessen konfigurationale Methoden zur Anwendung kommen und mit dem unter anderem ein minimaler Grad an Komplexität der Kausalstrukturen einhergeht. Für das konfigurationale kausale Modellieren gilt es ferner zu berücksichtigen, dass einerseits aufgrund der Intransitivität des über Minimale Theorien definierten Verursachungsbegriffs der Kettenschluss nicht anwendbar ist. Andererseits gibt es gemäss dem Kettenproblem relativ zu einer Faktormenge zu jeder Kausalkette empirisch äquivalente Common-Cause-Strukturen. Dieses Problem löst sich jedoch über die Bestätigung der Erweiterungsklausel auf.

3. Empirische Grundlage

Der kausaltheoretische Rahmen aus Kapitel 2 liefert die Antwort auf die Frage, wie der Verursachungsbegriff zu definieren ist, um basierend auf Regelmässigkeiten im An- und Abwesenheitsverhalten von Faktoren kausale Abhängigkeiten beliebiger Komplexität begrifflich abzubilden. In einem nächsten Schritt gilt es sich mit den empirischen Daten auseinanderzusetzen, die für die Aufdeckung solcher Abhängigkeiten generiert werden können. Der Fokus verschiebt sich somit vom begrifflichen Unterfangen der Definition der Kausalbeziehung hin zum epistemologischen Unterfangen des kausalen Schliessens, dessen Ausgangspunkt empirische Daten sind. Mit der Erzeugung von Vergleichssituationen zwischen Konfigurationen ist das wesentliche Merkmal solcher Daten bereits bekannt. Im vorliegenden Kapitel wird diese Datengrundlage weiter spezifiziert und ihre für das kausale Schliessen wesentlichen Eigenschaften werden dargestellt.

3.1. Von den Rohdaten zur Konfigurationstabelle

Konfigurationale Methoden decken kausale Abhängigkeiten zwischen Faktorliteralen auf, die durch singuläre Ereignisse oder Zustände instantiiert werden. Solche in Raum und Zeit verortete Ereignisse und Zustände stellen das empirische Fundament dar, aus dem sich Konfigurationstabellen bilden lassen. Untersucht man beispielsweise die kausale Relevanz verschiedener Literale für eine hohe politische Repräsentation von Frauen (R), muss ermittelt werden, in welcher Weise ihre Instanzen in einzelnen Fällen koinzidieren. Als Grundlage dienen dazu Rohdaten wie etwa jene aus Tabelle 3.1, die aus Krook (2010) entnommen wurden (vgl. S. 893). Dabei werden folgende Faktoren untersucht, wobei teilweise in Klammern zusätzlich angegeben wird, was in den Rohdaten konkret gemessen wird:

P := „proportionales Wahlsystem“ (Wahlsystem),

Q := „Frauenquote“,

S := „hoher gesellschaftlicher Status der Frau“ (Modell des Wohlfahrtsstaates),

Land	P	Q	S	B	L	R
Schweden	Proportional	Ja	Sozialdemokratisch	Keine	5	47.3
Finnland	Proportional	Nein	Sozialdemokratisch	Keine	5	42
Norwegen	Proportional	Ja	Sozialdemokratisch	Autonom	14	37.9
Dänemark	Proportional	Nein	Sozialdemokratisch	Autonom	7	36.9
Niederlande	Proportional	Ja	Konservativ	Autonom	7	36.7
Spanien	Proportional	Ja	Konservativ	Autonom	0	36
Belgien	Proportional	Ja	Konservativ	Autonom	13	34.7
Österreich	Proportional	Ja	Konservativ	Keine	8	32.2
Neuseeland	Mischform	Nein	Liberal	Autonom	9	32.2
Island	Proportional	Ja	Sozialdemokratisch	Autonom	9	31.7
Deutschland	Mischform	Ja	Konservativ	Autonom	7	31.6
Schweiz	Proportional	Ja	Konservativ	Keine	4	25
Australien	Mehrheit	Ja	Liberal	Autonom	0	24.7
Luxemburg	Proportional	Nein	Konservativ	Keine	7	23.3
Portugal	Proportional	Ja	Konservativ	Keine	0	21.3
Kanada	Mehrheit	Nein	Liberal	Autonom	0	20.8
Grossbritannien	Mehrheit	Ja	Liberal	Autonom	0	19.7
Frankreich	Mehrheit	Ja	Konservativ	Autonom	0	18.5
Italien	Mischform	Ja	Konservativ	Keine	3	17.3
USA	Mehrheit	Nein	Liberal	Autonom	0	16.3
Griechenland	Proportional	Ja	Konservativ	Keine	0	13
Irland	Mehrheit	Ja	Liberal	Autonom	2	13.3

Tabelle 3.1.: Beispiel von Rohdaten für die Untersuchung der hohen politischen Repräsentation von Frauen (R).

B := „Frauenbewegung“,

L := „starke Linksparteien“ (prozentualer Anteil Sitze im nationalen Parlament),

R := „hohe politische Repräsentation von Frauen“ (prozentualer Anteil von Frauen im nationalen Parlament).

Die Erhebung solcher Rohdaten ist eng mit der Aufstellung eines angemessenen Forschungsdesigns verknüpft. So muss etwa bestimmt werden, welche Fälle untersucht werden. Weiter müssen die untersuchten Faktoren gewählt und definiert sowie bestimmt werden, wie ihre Ausprägung in den einzelnen Fällen gemessen wird. Diese Vorarbeit ist massgeblich von substanziellem und theoretischem Wissen abhängig (vgl. u. a. Berg-Schlosser & Meur (2009)). Gut zu erkennen ist dies beispielsweise bei der Ermittlung der An- und Abwesenheit des Faktors S . Um

qualitative Unterschiede bezüglich des gesellschaftlichen Status der Frau in westlichen Ländern einzufangen, wählt Krook Modelle von Wohlfahrtsstaaten als Indikator. Als Begründung verweist sie auf eine über zwanzigjährige Forschungsarbeit, die aufzeigt, dass die drei erfassten Modelle eine ausschlaggebende Rolle bei der Gestaltung des Status der Frau einnehmen. Während demnach sozialdemokratische Wohlfahrtsstaaten die Vereinbarkeit von Mutterschaft und beruflicher Tätigkeit stärken, fördern konservative Wohlfahrtsstaaten die traditionelle Familienrolle der Frau. Liberale Wohlfahrtsstaaten zeichnen sich durch möglichst wenige Interventionen aus, wodurch sich Ungleichheiten oftmals verstärken (vgl. Krook (2010), S. 893 f.).

Sind die Rohdaten erhoben, werden sie in einem nächsten Schritt *kalibriert* (vgl. u. a. Ragin (2008), Part 2).¹ Mit der Kalibrierung wird die Frage beantwortet, wann eine Instanz die von Faktoren beschriebenen Eigenschaften in welchem Masse erfüllt. So muss im Zusammenhang mit obigen Rohdaten unter anderem definiert werden, wann der Status der Frau als hoch oder wann Linksparteien als stark einzustufen sind. Die Rohdaten selbst geben darüber oftmals noch nicht direkt Auskunft und liefern vorerst nur den unmittelbar erhobenen qualitativen oder quantitativen Wert, auf dem diese Einstufung vorgenommen wird. Der Status der Frau etwa wird über die verschiedenen Modelle von Wohlfahrtsstaaten eingefangen. Welches Modell aber mit einem hohen Status einhergeht und welches nicht, muss noch ermittelt werden. Wie bei der Erhebung ist dabei auch die Kalibrierung stark von substanziellem und theoretischem Wissen abhängig. So beruft sich Krook bei der Kalibrierung des Status der Frau auf Untersuchungen, die sozialdemokratische Länder im Vergleich zu den anderen Modellen von Wohlfahrtsstaaten als „frauenfreundlich“ nachweisen. Krook zufolge dürften Frauen damit besonders in sozialdemokratischen Ländern einen hohen sozialen sowie wirtschaftlichen Status genießen. Sie ermittelt dementsprechend genau für jene Fälle einen hohen Status, bei denen es sich um solche sozialdemokratische Länder handelt. Weiter definiert Krook, gestützt auf theoretischem Wissen sowie der Verteilung der einzelnen Werte in den Rohdaten, beispielsweise für starke Linksparteien einen Schwellenwert von sieben Prozent, sodass genau jene Linksparteien als stark eingestuft werden, die im Parlament einen Sitzanteil von sieben Prozent oder mehr aufweisen. Auf ähnliche Weise wird die Repräsentation von Frauen in nationalen Parlamenten genau dann als hoch eingestuft, wenn der prozentuale Anteil mindestens 30 Prozent erreicht (vgl. Krook (2010), S. 394 f.). Mit einer transparenten Kalibrierung für die restlichen Faktoren ergibt sich daraus schließlich die kalibrierte Datentabelle 3.2a, wobei eine 0 beziehungsweise eine 1 wie gewohnt zum Ausdruck bringt, dass der Faktor in der Spaltenüberschrift im entsprechenden Fall anwesend beziehungsweise abwesend ist.

¹Für die Kalibrierung quantitativer Rohdaten wurden bei QCA u. a. eine direkte und eine indirekte Methode des Kalibrierens entwickelt (vgl. z. B. Ragin (2008), Kapitel 5). Für die Kalibrierung qualitativer Rohdaten vgl. bspw. Basurto & Speer (2012).

Land	<i>P</i>	<i>Q</i>	<i>S</i>	<i>B</i>	<i>L</i>	<i>R</i>							
Schweden	1	1	1	0	0	1							
Finnland	1	0	1	0	0	1							
Norwegen	1	1	1	1	1	1							
Dänemark	1	0	1	1	1	1							
Niederlande	1	1	0	1	1	1							
Spanien	1	1	0	1	0	1							
Belgien	1	1	0	1	1	1							
Österreich	1	1	0	0	1	1	<i>P</i>	<i>Q</i>	<i>S</i>	<i>B</i>	<i>L</i>	<i>R</i>	Fallzahl
Neuseeland	0	0	0	1	1	1	1	1	1	1	1	1	2
Island	1	1	1	1	1	1	1	1	1	0	0	1	1
Deutschland	0	1	0	1	1	1	1	1	0	1	1	1	2
Schweiz	1	1	0	0	0	0	1	1	0	1	0	1	1
Australien	0	1	0	1	0	0	1	1	0	0	1	1	1
Luxemburg	1	0	0	0	1	0	1	1	0	0	0	0	3
Portugal	1	1	0	0	0	0	1	0	1	1	1	1	1
Kanada	0	0	0	1	0	0	1	0	1	0	0	1	1
Grossbritannien	0	1	0	1	0	0	1	0	0	0	1	0	1
Frankreich	0	1	0	1	0	0	0	1	0	1	1	1	1
Italien	0	1	0	0	0	0	0	1	0	1	0	0	4
USA	0	0	0	1	0	0	0	1	0	0	0	0	1
Griechenland	1	1	0	0	0	0	0	0	0	1	1	1	1
Irland	0	1	0	1	0	0	0	0	0	1	0	0	2

(a)

(b)

Tabelle 3.2.: Tabelle 3.2a zeigt die aus den Rohdaten aus Tabelle 3.1 erzeugten kalibrierten Daten. Tabelle 3.2a widerspiegelt dieselben Informationen in Form einer Konfigurationstabelle mit Fallzahlen.

Jede der aus Nullen und Einsen bestehenden Zeilen aus Tabelle 3.2a gibt einen kalibrierten Fall wieder. Aus dieser kalibrierten Datentabelle wird schliesslich die Konfigurationstabelle 3.2b gebildet. Dazu werden jene kalibrierten Fälle, bei denen die Faktoren das gleiche Instantiierungsverhalten aufweisen, zu einer *Konfiguration* zusammengefasst und letztere ihrerseits um die Anzahl der zusammengefassten Fälle, der sogenannten Fallzahl, ergänzt. Eine Konfiguration ist also ein bestimmter Falltyp, der ein Instantiierungsverhalten von Faktoren in geeigneter räumlicher und zeitlicher Nähe beschreibt, das von Fällen instantiiert wird.² Die aus lauter Einsen bestehende Konfiguration beispielsweise besagt, dass alle sechs Faktoren in der Spaltenüberschrift zusammen anwesend sind. Das heisst, ein proportionales Wahlsystem, Frauenquoten, ein hoher

²Zur räumlichen und zeitlichen Nähe vgl. S. 24 f.

Status der Frau, Frauenbewegungen, starke Linksparteien und eine hohe politische Repräsentation der Frau treten gemeinsam auf. Die entsprechenden zwei Einzelfälle dieser Konfiguration sind Norwegen und Island. Wie dieses konkrete Beispiel andeutet, stehen bei den Konfigurationen nicht mehr die Einzelfälle im Vordergrund, sondern die allgemeinere Information, ob ein bestimmtes Instantiierungsverhalten der Faktoren realisierbar ist. Schliesslich ist es in erster Linie auch diese von Einzelfällen abstrahierte Information, die für die Suche nach Kausalstrukturen mittels konfiguratoraler Methoden wesentlich ist.

Es gilt an dieser Stelle zu bemerken, dass Konfigurationen hier im Sinne von CNA verstanden werden. Konkret bedeutet dies, dass mit einer gegebenen Konfiguration von Faktoren aus einer Menge F lediglich eine Aussage darüber getroffen wird, ob und wie die Faktoren aus F gemeinsam auftreten. Welches Literal dabei eine Wirkung ist oder als solche betrachtet werden sollte, ist (noch) nicht festgelegt. Solche Konfigurationen werden auch als *Koinzidenzen* bezeichnet, woraus sich der Name „Coincidence Analysis“, kurz CNA, ergibt. Im Gegensatz dazu ist bei QCA eine Konfiguration jeweils eine Verlinkung von als „Bedingungen“ bezeichneten potenziellen Ursachen und einer als „Outcome“ bezeichneten Wirkung (vgl. u. a. Rihoux & Ragin (2009), S. xxiv). Bei QCA steht also bereits von Beginn weg explizit eine Wirkung im Fokus, deren Verhalten in Abhängigkeit von potenziellen kausal relevanten Literalen untersucht wird. Eine solche Konfiguration im Sinne von QCA kann verstanden werden als Konfiguration im Sinne von CNA, die zusätzlich mit der Information ausgestattet ist, welche Literale Wirkungen sind oder als solche betrachtet werden. Diese zusätzliche Information ergibt sich typischerweise aus der Forschungsfrage und hängt vom theoretischen und substanziellen Wissen ab, das bei der Entwicklung des Forschungsdesigns eingesetzt wird (vgl. u. a. Berg-Schlosser & Meur (2009), S. 20). So gesehen können Konfigurationen im Sinne von CNA als Verallgemeinerungen von Konfigurationen im Sinne von QCA verstanden werden, indem erstere mit den Informationen angereichert werden, welche Literale Wirkungen sind. Angesichts der konkreten Forschungsfrage kann dies im obigen Beispiel umgesetzt werden, indem man oberhalb der Faktoren eine weitere Zeile hinzufügt, die anzeigt, dass P , Q , S , B und L Bedingungen sind und R der Outcome. Dabei interessiert man sich typischerweise dafür, welche Literale dazu führen, dass R auftritt, das heisst, man interessiert sich für die Ursachen des positiven Literals R .

In Tabelle 3.2b sind alle Konfigurationen dichotomisiert. Demnach sind die qualitativen und quantitativen Primärdaten aus Tabelle 3.1 auf die Werte 0 und 1 codiert. Eine aus lediglich diesen beiden Werten bestehende Konfigurationstabelle wird auch als *crisp-set Konfigurationstabelle* bezeichnet. Entgegen einer solchen crisp-set Kalibrierung kommt es oftmals vor, dass man ein Ereignis (oder Zustand) nicht entweder als eine Instanz oder keine Instanz eines Faktors betrachten möchte, da die vom Faktor repräsentierten Eigenschaften nur teilweise gegeben sind. Die Zweiwertigkeit greift oft zu kurz und verzerrt die beobachteten Daten. Zwar mag

beispielsweise Norwegen mit 14 Prozent Sitzanteil starke Linksparteien aufweisen und etwa Spanien mit 0 Prozent nicht. Aber Neuseeland zum Beispiel liegt mit einem Sitzanteil von 9 Prozent irgendwo dazwischen, sodass man die Linksparteien für diesen Fall nicht unbedingt als „stark“, sondern eher als „ziemlich stark“ oder „eher stark als nicht stark“ bewerten möchte. Zu diesem Zweck hat QCA die zweiwertige Analysevariante, die sogenannte *crisp-set QCA* (csQCA), um mehrwertige Analysevarianten erweitert, von denen sich die von Lotfi Zadeh Zadeh (1965) entwickelten *Fuzzy-Mengen* nutzende Analysemethode, die sogenannte *fuzzy-set QCA* (fsQCA), besonders hervorhebt (vgl. bspw. Ragin (2008); Schneider & Wagemann (2012)). Diese Analysevariante erlaubt es beispielsweise im Fall von Neuseeland, dem Faktor *L* anstelle des Wertes 1 einen (Fuzzy-)Wert wie 0.6 zuzuweisen. Die vorliegende Arbeit beschränkt sich jedoch auf die Entwicklung eines Verfahrens für die Analyse von crisp-set Konfigurationstabellen, auf die auch das kausaltheoretische Fundament des vorhergehenden Kapitels ausgelegt ist. Die wesentlich komplexere fuzzy-set Variante, die als Verallgemeinerung der crisp-set Variante betrachtet werden kann, muss in einer weiterführenden Arbeit behandelt werden.

Die aus der Kalibrierung von erhobenen Rohdaten resultierenden Konfigurationstabellen bilden den Ausgangspunkt des analytischen Moments, das in der vorliegenden Arbeit im Zentrum steht und jenem Teilprozess des konfiguralen kausalen Modellierens entspricht, den man als logische Datenverarbeitung bezeichnen könnte. Die Kernaufgabe des analytischen Moments ist es, ausgehend von Konfigurationstabellen, algorithmisch Minimale Theorien zu suchen. Bevor näher auf dieses analytische Moment eingegangen wird, werden neben der crisp-set Kalibrierung der Daten zwei zusätzliche Annahmen getroffen. Diese Annahmen bezwecken zusammen mit weiteren Bedingungen, die zu einem späteren Zeitpunkt in diesem Kapitel ausgearbeitet werden, die Idealisierung der untersuchten Daten. Diese Idealisierungsmaßnahmen sind notwendig, um die Korrektheit eines Aufdeckungsalgorithmus zu prüfen. Die Daten werden soweit idealisiert, dass allfällige unerwünschte oder fehlerhafte Folgerungen des Verfahrens eindeutig auf ein Fehlverhalten des Algorithmus zurückgeführt werden können und nicht etwa auf Datenmängel. Diese Idealisierungen führen zwar zu praxisfernen Daten, doch erst wenn die Korrektheit des Algorithmus unter idealen Bedingungen sichergestellt ist, ist es sinnvoll, auch Kausalanalysen zu betrachten, die der Praxis näher kommen. Ist die Korrektheit des Algorithmus verifiziert, können die Bedingungen nach und nach abgebaut und das Verfahren kann entsprechend weiterentwickelt werden.

Zum einen wird in der Folge angenommen, dass die Datenerhebung frei von praktischen Hürden wie mangelhaften Kalibrierungen oder Messfehlern ist. Solche Schwierigkeiten können zwar zu Fehlschlüssen führen, gehen aber nicht auf ein Fehlverhalten des Algorithmus zurück. So kann etwa fälschlicherweise die Suffizienz von *A* für *W* gefolgert werden, weil das entsprechende Instantiierungsverhalten der beiden Faktoren in den beobachteten Fällen nicht korrekt gemes-

$\mathcal{K}_{3,3}$	A	B	C
c_1	1	1	1
c_2	1	0	1
c_3	0	1	1
c_4	0	0	0

Tabelle 3.3.: Beispiel einer hypothetischen crisp-set Konfigurationstabelle $\mathcal{K}_{3,3}$ über die geprüfte Faktormenge $\mathbf{A} = \{A, B, C\}$, wie sie künftig typischerweise als empirische Analysegrundlage dargestellt wird.

sen wurde. Der Ausschluss solcher (praktischen) Probleme wird künftig unter dem Begriff der *angemessenen Güte* der Daten subsumiert. Wird angenommen, dass eine untersuchte Konfigurationstabelle eine angemessene Güte hat, bedeutet dies, dass alle in der Tabelle enthaltenen Konfigurationen tatsächlich auch in dieser Form empirisch möglich sind.

Zum anderen werden die Fallzahlen vernachlässigt, indem davon ausgegangen wird, dass sie für jede Konfiguration identisch sind. Die Fallzahlen kommen vor allem dann zum Tragen, wenn die Daten durch Störgrößen beeinflusst werden. Beispielsweise bieten sie ein Mittel für den Umgang mit ungewöhnlichen Ausreißern. Derartige Szenarien, die vor allem in der praktischen Umsetzung auftreten, werden hier nicht betrachtet. Ist die Korrektheit des Algorithmus in idealer Umgebung nachgewiesen, können solche Einflussgrößen in einem nächsten Entwicklungsschritt eingebaut werden.

Die bis hierhin formulierten Voraussetzungen werden im weiteren Verlauf dieser Arbeit üblicherweise nicht mehr angegeben. Sofern nicht explizit anders erwähnt, werden sie standardmäßig als gegeben angenommen. Insgesamt werden somit künftig crisp-set Konfigurationstabellen betrachtet, die die Form haben wie beispielsweise jene aus Tabelle 3.3 und in dieser Gestalt bereits in Kapitel 2 ausgiebig genutzt wurden. Für jede in einer Konfigurationstabelle enthaltene Konfiguration gilt, dass sie einem empirisch möglichen Instantiierungsverhalten der Faktoren in der Spaltenüberschrift entspricht. Aus $\mathcal{K}_{3,3}$ lässt sich somit unter anderem gemäss c_3 entnehmen, dass B zusammen mit C auftritt, während A abwesend ist. Konkret wurden also Fälle beobachtet oder gemessen, bei denen B und C zusammen instantiiert sind und A nicht gegeben ist. Fehlerquellen wie Mess- oder Kalibrierungsfehler können dabei ausgeschlossen werden.

Bei solchen Konfigurationstabellen wie $\mathcal{K}_{3,3}$ setzt das analytische Moment an, das im folgenden Abschnitt etwas ausführlicher skizziert werden soll. Zumal der Begriff einer Konfigurationstabelle im Sinne von CNA verstanden wird, orientiert sich dabei auch das analytische Moment insgesamt an CNA. Zweck dieser Skizzierung ist in erster Linie die Vermittlung der groben Systematik des Aufdeckungsverfahrens, die für das Verständnis der noch folgenden Charakterisierungen der empirischen Grundlage und ihrer Einflüsse auf das kausale Schliessen von

$\mathcal{K}_{3,4}$	A	B	C	D	E
c_1	1	1	1	1	1
c_2	1	1	1	0	1
c_3	1	0	1	1	1
c_4	1	0	1	0	1
c_5	0	1	1	1	1
c_6	0	1	1	0	1
c_7	0	0	0	1	1
c_8	0	0	0	0	0

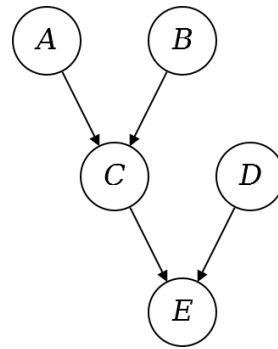


Abbildung 3.4.: Einer Kausalanalyse unterzogene Konfigurationen $\mathcal{K}_{3,4}$ der Faktoren aus $\{A, B, C, D, E\}$, zusammen mit einem Kausalgraphen, der ein Kausalmodell von $\mathcal{K}_{3,4}$ darstellt.

Nutzen sein wird. Dazu werden einzelne ausgewählte Prozessschritte ausführlicher beschrieben als andere, sofern sie für dieses Verständnis besonders förderlich sind. Eine detaillierte logische Zergliederung aller Prozessschritte des analytischen Moments wird jedoch erst in Kapitel 5 vorgenommen.

3.2. Skizzierung des analytischen Moments

Innerhalb des analytischen Moments werden die Minimalen Theorien einer Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ über eine Faktormenge $\mathbf{F}$ gesucht und daraus schliesslich sogenannte *Kausalmodelle* gebildet, die das Endprodukt des analytischen Moments darstellen. Ein zu einer beliebigen Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ gehöriges Kausalmodell kann als mögliche, den Konfigurationen $\mathcal{K}_{\mathbf{F}}$ zugrunde liegende Kausalstruktur interpretiert werden und lässt sich durch kausal wohlgeformte Hypergraphen darstellen. Betrachtet man beispielsweise die bereits bekannten Konfigurationen $\mathcal{K}_{3,4}$ aus Abbildung 3.4, wäre unter anderem die in derselben Abbildung dargestellte Kausalkette ein Kausalmodell von $\mathcal{K}_{3,4}$.

Als mögliche, den Konfigurationen $\mathcal{K}_{3,4}$ zugrunde liegende Kausalstruktur wird von diesem Kausalmodell gesagt, dass es mit $\mathcal{K}_{3,4}$ *verträglich* ist. Allgemein wird ein mit beliebigen Konfigurationen $\mathcal{K}_{\mathbf{F}}$ verträgliches Kausalmodell folgendermassen definiert:

Definition 3.1 (Kausalmodell)

Seien $\mathcal{K}_{\mathbf{F}}$ Konfigurationen der Faktoren einer Faktormenge $\mathbf{F}$. Ein mit $\mathcal{K}_{\mathbf{F}}$ verträgliches Kausalmodell ω ist eine Faktoren aus $\mathbf{F}' \supseteq \mathbf{F}$ beinhaltende, als kausal wohlgeformter Hypergraph

darstellbare Struktur oder als kausal wohlgeformter Hypergraph darstellbare Teilstruktur einer Struktur, die die untersuchten Konfigurationen $\mathcal{K}_F$ generieren kann.

Unter anderem ist die den Konfigurationen $\mathcal{K}_F$ tatsächlich zugrunde liegende Kausalstruktur trivialerweise ein mit $\mathcal{K}_F$ verträgliches Kausalmodell. Ferner gilt, dass, wenn es ein mit $\mathcal{K}_F$ verträgliches Kausalmodell gibt, grundsätzlich auch alle kausal wohlgeformten Teilstrukturen dieses Kausalmodells mit $\mathcal{K}_F$ verträglich sind. So ist bezüglich des Beispiels aus Abbildung 3.4 neben der Kausalkette auch jene Struktur ein mit $\mathcal{K}_{3,4}$ verträgliches Kausalmodell, die nur aus den Alternativursachen A sowie B und deren Wirkung C besteht. Umgekehrt kann ein Kausalmodell auch Faktoren beinhalten, die nicht Teil der untersuchten Faktormenge sind. Ein derartiges Kausalmodell kann beispielsweise dann gebildet werden, wenn durch die Integration solcher zusätzlichen Faktoren ein relevanter Beitrag zur kausalen Erklärung des in den untersuchten Konfigurationen enthaltenen Instantiierungsverhaltens geleistet wird und die Integration darüber hinaus basierend auf substanziellem Wissen gerechtfertigt ist. Üblicherweise interessiert man sich jedoch nur für jene Kausalmodelle, die eine kausale Erklärung für den grösstmöglichen Teil der Konfigurationstabelle über eine Faktormenge F liefern und auch nur aus Faktoren aus F bestehen. Gibt es also einerseits zwei nur aus Faktoren aus F bestehende sowie mit $\mathcal{K}_F$ verträgliche Kausalmodelle ω_1 und ω_2 , sodass ω_1 eine Teilstruktur von ω_2 beschreibt, wird üblicherweise nur ω_2 betrachtet. Liegt andererseits ein mit $\mathcal{K}_F$ verträgliches Kausalmodell ω_3 vor, welches ω_2 als Teilstruktur aufweist und zusätzlich die Faktoren der Faktormenge $\{X_1, \dots, X_k\}$ beinhaltet, wobei $\{X_1, \dots, X_k\} \cap F = \emptyset$ und $k \geq 1$ gilt, betrachtet man typischerweise ebenfalls ω_2 und nicht ω_3 . Das Modell ω_3 wird allenfalls relativ zur Faktormenge $F' = \{X_1, \dots, X_k\} \cup F$ gegenüber ω_2 bevorzugt, sofern es nach wie vor mit den entsprechenden Konfigurationen $\mathcal{K}_{F'}$ verträglich ist.

Wird ein mit Konfigurationen $\mathcal{K}_F$ verträgliches Kausalmodell gebildet, wird damit noch nicht behauptet, dass dieses Kausalmodell die tatsächlichen Kausalzusammenhänge wiedergibt, sondern etwas abgeschwächt gesagt, dass die im Modell dargestellten Zusammenhänge relativ zu den vorhandenen Daten aus $\mathcal{K}_F$ als kausale Abhängigkeiten infrage kommen. Mit dieser abgeschwächten Interpretation trägt man unter anderem dem induktiven Risiko Rechnung, das gemäss dem hier verwendeten Verursachungsbegriff erst nach Erhärtung der Erweiterungsklausel als tragbar betrachtet werden kann. Ferner wird mit dieser schwächeren Deutung auch berücksichtigt, dass Konfigurationstabellen mehrdeutig sein können. Angesichts der Definition 2.11 der kausalen Relevanz können die mit Konfigurationen $\mathcal{K}_F$ verträglichen Kausalmodelle beim konfigurationsalen kausalen Modellieren demnach mit den Minimalen Theorien der untersuchten Konfigurationstabellen gleichgesetzt werden, die, etwa ausgedrückt mittels Kausalgraphen, im Rahmen einer Kausalhypothese kausal interpretiert werden.

Das analytische Moment, dessen Ziel die Bildung solcher Kausalmodelle ist, kann grob in drei aufeinanderfolgende Verfahrensschritte gegliedert werden, die sich ihrerseits in weitere Teilschritte unterteilen lassen. Die drei Hauptschritte bestehen aus (I) der Bestimmung aller RDN-Bikonditionale, (II) der Bildung von Minimalen Theorien aus den RDN-Bikonditionalen und (III) der Bildung von Kausalmodellen aus den Minimalen Theorien. Was genau unter diesen Schritten zu verstehen ist, soll anhand einer Analyse der Konfigurationen aus $\mathcal{K}_{3,4}$ beschrieben werden. Dabei soll davon ausgegangen werden, dass bis auf die untersuchten Konfigurationen keine weiteren Informationen vorliegen, die in irgendeiner Form bei der Folgerung eines Kausalmodells behilflich sein könnten. Dieser Annahme zufolge kann neben der Konfigurationstabelle $\mathcal{K}_{3,4}$ auf kein zusätzliches Vorwissen für eine Analyse zurückgegriffen werden. In der Praxis liegt eine solche Ausgangslage typischerweise dann vor, wenn man versucht herauszufinden, ob und wie untersuchte Faktoren kausal miteinander verknüpft sind, ohne dass substanzielles Wissen oder andere Indizien, wie beispielsweise die zeitliche Abfolge der Faktorinstantiierungen, herangezogen werden können. Im Allgemeinen vereinfacht das Vorhandensein von Vorwissen das Aufdeckungsverfahren, indem beispielsweise gewisse Literale von Faktoren der geprüften Faktormenge von Beginn weg als mögliche Ursachen, nicht aber als Wirkungen von Literalen anderer Faktoren aus der geprüften Faktormenge infrage kommen. Mit diesem kausalen Vorwissen kann je nachdem von Anfang an auf eine Suche nach einfachen Minimalen Theorien eines solchen Literals verzichtet werden, womit sich das Aufdeckungsverfahren um einige Analyseschritte reduzieren lässt. So gesehen kann ein Aufdeckungsverfahren ohne Annahme von Vorwissen als die allgemeine Variante betrachtet werden, die das analytische Moment in seiner gesamten logischen Zergliederung aufzeigt und bei jeder Analyse in dieser oder in einer vereinfachten Form davon zum Einsatz kommt. Sofern nicht explizit anders erwähnt, wird deshalb künftig standardmässig davon ausgegangen, dass bis auf die Konfigurationstabelle keine weiteren Informationen vorliegen.

(I) Bestimmung aller RDN-Bikonditionale

Die wesentlichen Bausteine zur Bildung eines Kausalmodells sind die Minimalen Theorien. Die Minimalen Theorien bestehen aus RDN-Bikonditionalen, die als Teil der Minimalen Theorien einfache Minimale Theorien darstellen. Ist man auf der Suche nach Kausalmodellen, die mit einer untersuchten Konfigurationstabelle $\mathcal{K}_F$ verträglich sind, ist es also sinnvoll, zuerst die RDN-Bikonditionale zu ermitteln. Diese RDN-Bikonditionale lassen sich allenfalls konjunktiv zu Minimalen Theorien von $\mathcal{K}_F$ verknüpfen und im Rahmen eines Kausalmodells kausal interpretieren.

Liegen bis auf eine beliebige Konfigurationstabelle $\mathcal{K}_F$ keine weiteren Informationen vor, bedeutet dies unter anderem, dass nicht bekannt ist, von welchen Literalen von Faktoren aus der geprüften Faktormenge F RDN-Bikonditionale existieren. Dementsprechend muss man vorerst davon ausgehen, dass es von jedem Literal RDN-Bedingungen gibt. Bei einigen Faktoren aus F kann jedoch von Beginn weg auf eine Suche nach RDN-Bikonditionalen ihrer Literale verzichtet und so der Aufdeckungsprozess etwas effizienter gestaltet werden. Erfüllt ein Faktor Z aus F eine der folgenden beiden Bedingungen, lässt sich aus den entsprechenden Konfigurationen kein RDN-Bikonditional eines Literals dieses Faktors ableiten:

- (WD1) Der Faktor Z ist in der geprüften Konfigurationstabelle $\mathcal{K}_F$ (a) in allen Konfigurationen anwesend oder (b) in allen Konfigurationen abwesend.
- (WD2) In der geprüften Konfigurationstabelle $\mathcal{K}_F$ sind zwei Konfigurationen enthalten, die sich nur darin unterscheiden, dass Z in der einen Konfiguration anwesend ist und in der anderen nicht.

Von einem Literal eines Faktors Z , der Bedingung (WD1a) erfüllt, wird sich keine minimal hinreichende Bedingung finden lassen. Es können zwar hinreichende Bedingungen von Z abgeleitet werden, die können aber nicht minimal hinreichend für Z sein. Letzteres gilt deshalb, weil bei einer solchen Datenlage für jede in $\mathcal{K}_F$ gegebene Konjunktion von Literalen von Faktoren aus $F \setminus \{Z\}$ gilt, dass sowohl diese als auch jeder echte Teil davon bis hin zum 0-stelligen Konjunktionsterm hinreichend für Z sind. Sei konkret $X_1 X_2$ eine solche hinreichende Bedingung von Z . Ist Z in allen Konfigurationen anwesend, gilt auch für die echten Teile X_1 und X_2 von $X_1 X_2$, dass immer, wenn diese gegeben sind, Z gegeben ist. Demnach ist $X_1 X_2$ zwar hinreichend, aber nicht minimal hinreichend für Z . Doch auch die hinreichenden Literale X_1 und X_2 sind nicht minimal hinreichend für Z , da aufgrund des immer auftretenden Faktors Z ebenfalls der aus keinem Literal bestehende 0-stellige Konjunktionsterm hinreichend für Z ist. Der 0-stellige Konjunktionsterm ist ein echter Teil von X_1 und X_2 , kann seinerseits aber keine minimal hinreichende Bedingung von Z darstellen, da eine solche per Definition aus mindestens einem Literal bestehen muss (vgl. S. 33).³ Auf analoge Weise kann für jede weitere in $\mathcal{K}_F$ enthaltene hinreichende Bedingung von Z argumentiert werden, womit insgesamt gefolgert werden kann, dass aus $\mathcal{K}_F$ keine minimal hinreichende Bedingung von Z und damit kein RDN-Bikonditional von Z abgeleitet werden können.

³Aus rein formaler Sicht mag der 0-stellige Konjunktionsterm in diesem Fall zwar minimal hinreichend für Z sein, da er hinreichend ist für Z und es trivialerweise keinen echten Teil davon gibt, der hinreichend für Z ist. Der Begriff der minimalen Suffizienz wird in der vorliegenden Arbeit jedoch im Hinblick auf das kausale Modellieren definiert, weshalb der 0-stellige Konjunktionsterm als minimal hinreichende Bedingung ausgeschlossen und gefordert wird, dass jede minimal hinreichende Bedingung mindestens ein Literal enthält. Andernfalls wäre man gezwungen zu sagen, dass Z durch nichts verursacht wird, was dem Kausalitätsprinzip widersprechen würde.

Bei (WD1b) sind aufgrund des jeweils ausbleibenden Faktors Z keine Bedingungen auszumachen, für die gilt, dass, wenn sie gegeben sind, auch Z gegeben ist. Die einzige Ausnahme, für die Letzteres gilt, sind Kontradiktionen wie die Konjunktion $X\bar{X}$ eines Faktors X aus $\mathbf{F}$. Eine solche Kontradiktion wird innerhalb des konfiguralen kausalen Modellierens jedoch deshalb nicht als hinreichende Bedingung betrachtet, weil gleiche Faktoren in einer Konfiguration entweder anwesend oder abwesend sind, aber nicht beides. Die Kontradiktion ist also empirisch nicht möglich und kann deshalb auch nie eine Ursache sein. Bei (WD1b) lassen sich also mit dem Ausschluss solcher Kontradiktionen keine hinreichenden Bedingungen von Z finden und damit schliesslich wiederum auch kein RDN-Bikonditional von Z .

Bei einer Datenlage gemäss Bedingung (WD2) lässt sich aus $\mathcal{K}_{\mathbf{F}}$ keine aus minimal hinreichenden Bedingungen bestehende notwendige Disjunktion von Z ableiten, weil es in $\mathcal{K}_{\mathbf{F}}$ mindestens eine Konfiguration mit auftretendem Z gibt, bei der keine hinreichende Bedingung von Z gegeben ist. Sei $X_1 \cdots X_n$, $n \geq 1$, die Konjunktion von Literalen der Faktoren aus $\mathbf{F} \setminus \{Z\}$, die in einer Konfiguration mit Z und in einer Konfiguration ohne Z gegeben ist. Da Z einerseits zusammen mit $X_1 \cdots X_n$ gegeben ist und andererseits nicht obige Bedingung (WD1) erfüllt, muss zur Bildung eines RDN-Bikonditionals von Z der Konjunktionsterm $X_1 \cdots X_n$ oder echte Teile davon minimal hinreichend für Z sein. Dies ist jedoch deshalb nicht möglich, weil derselbe Konjunktionsterm $X_1 \cdots X_n$ auch zusammen mit $\bar{Z}$ gegeben ist und demnach weder $X_1 \cdots X_n$ noch echte Teile davon hinreichend für Z sein können. Folglich kann aus $\mathcal{K}_{\mathbf{F}}$ für Z keine aus minimal hinreichenden Bedingungen bestehende notwendige Disjunktion und damit kein RDN-Bikonditional von Z abgeleitet werden.

Bei obiger Argumentation wird von einer strengen Auslegung der Definition von Minimalen Theorien und damit der Definition von RDN-Bikonditionalen ausgegangen. Eine Minimale Theorie ist demnach genau dann wahr, wenn die Bedingungen aus Definition 2.8 von allen Fällen gestützt werden. Eine derart strikte Anwendung ergibt dann Sinn, wenn die Daten in idealisierter Form vorliegen und beispielsweise zu Ausreissern führende Störgrössen ausgeschlossen sind. Zumal das Ziel dieses Kapitels unter anderem die Ausarbeitung einer idealen Datenlage ist, ist diese strenge Auslegung angebracht.

Wird die Konfigurationstabelle $\mathcal{K}_{3,4}$ auf Faktoren überprüft, die obige Bedingungen (WD1) oder (WD2) erfüllen, bleiben die beiden Faktoren C und E zurück. Alle anderen Faktoren verletzen die Bedingung (WD2) und es können demnach basierend auf $\mathcal{K}_{3,4}$ keine RDN-Bikonditionale von Literalen dieser Faktoren gefunden werden. Kein Literal dieser Faktoren kann so mit Literalen von anderen Faktoren der geprüften Faktormenge in Beziehung gebracht werden, dass es einzeln auf einer Seite eines Bikonditionals steht und sich auf der anderen Seite dieses Bikonditionals eine für das Literal minimal notwendige Disjunktion minimal hinreichender Be-

dingungen befindet. Dies soll nochmals konkret anhand von A veranschaulicht werden. Gibt es ein RDN-Bikonditional von A , das sich aus $\mathcal{K}_{3,4}$ ableiten lässt, muss aus $\mathcal{K}_{3,4}$ eine minimal notwendige Disjunktion von minimal hinreichenden Bedingungen von A gefolgert werden können. Daraus folgt unter anderem, dass $\mathcal{K}_{3,4}$ ein solches RDN-Bikonditional von A nur dann impliziert, wenn immer, wenn A gegeben ist, auch mindestens eine minimal hinreichende Bedingung von A aus dem RDN-Bikonditional gegeben ist. In c_1 ist zwar A anwesend, aber es ist keine minimal hinreichende Bedingung von A gegeben. Wäre eine solche gegeben, müsste es sich um die Konjunktion $BCDE$ oder einen echten Teil davon handeln. Wie anhand der Konfiguration c_5 ersichtlich ist, ist diese Konjunktion jedoch nicht hinreichend für A und damit weder $BCDE$ noch ein Teil davon minimal hinreichend für A . Folglich kann aus $\mathcal{K}_{3,4}$ kein RDN-Bikonditional und damit schliesslich auch keine einfache Minimale Theorie von A abgeleitet werden.

Nach dem Ausschluss von Faktoren aus $\mathbf{F}$ mit Zuhilfenahme von (WD1) und (WD2) wird nach den RDN-Bikonditionalen von Literalen der verbleibenden Faktoren gesucht. Für C und E können aus der Konfigurationstabelle $\mathcal{K}_{3,4}$ bekanntlich folgende RDN-Bikonditionale abgeleitet werden (vgl. Abschnitt 2.11.2, S. 70 ff.):

$$A + B \leftrightarrow C \quad (3.1)$$

$$A + B + D \leftrightarrow E \quad (3.2)$$

$$C + D \leftrightarrow E \quad (3.3)$$

(II) Bildung von Minimalen Theorien

Die Gesamtheit aller RDN-Bikonditionale aus Schritt (I), die aus der untersuchten Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ folgen, erzeugen zusammen ein System Δ von RDN-Bikonditionalen. Gibt es Minimale Theorien von $\mathcal{K}_{\mathbf{F}}$, bestehen diese notwendigerweise aus einer konjunktiven Verknüpfung von RDN-Bikonditionalen aus Δ .

Für Minimale Theorien gilt gemäss Bedingung (c) der Definition 2.8, dass die Faktoren, von deren Literalen die Minimale Theorie RDN-Bikonditionale enthält, paarweise verschieden sind. Folglich kann eine Minimale Theorie von $\mathcal{K}_{\mathbf{F}}$ von jedem Literal mit RDN-Bedingung maximal ein RDN-Bikonditional aus dem System Δ von RDN-Bikonditionalen enthalten. In Betracht dessen werden in einem ersten Teilschritt alle möglichen Konjunktionen von RDN-Bikonditionalen aus Δ gebildet, sodass jede Konjunktion von jedem Literal genau ein RDN-Bikonditional enthält. Aus dem aus $\mathcal{K}_{3,4}$ gefolgerten System lassen sich insgesamt zwei solche Konjunktionen bilden:

$$(A + B \leftrightarrow C) \cdot (A + B + D \leftrightarrow E) \quad (3.4)$$

$$(A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E) \quad (3.5)$$

Damit liegen bereits die Minimalen Theorien von $\mathcal{K}_{3,4}$ vor. Neben der Bedingung (c) aus Definition 2.8 gilt für beide Ausdrücke (3.4) und (3.5), dass sie (a) äquivalent zur Konjunktion $(A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E) \cdot (A + B + D \leftrightarrow E)$ aller aus $\mathcal{K}_{3,4}$ gefolgerten RDN-Bikonditionale sind und dies (b) für keinen echten Teil von (3.4) beziehungsweise (3.5) gilt.

(III) Bildung von Kausalmodellen

Mit der Ermittlung der Minimalen Theorien und den damit einhergehenden einfachen Minimalen Theorien liegen alle redundanzfreien Regularitäten vor, aus denen die mit $\mathcal{K}_F$ verträglichen Kausalmodelle $\omega_1, \dots, \omega_i, \dots, \omega_n$, $1 \leq i \leq n$, gebildet werden können. In einem letzten Schritt werden die Minimalen Theorien kausal interpretiert und damit auch ihre strukturellen Eigenheiten ausgearbeitet. Bezüglich letzteren wird etwa untersucht, ob in den Minimalen Theorien komplexe Minimale Theorien enthalten sind und falls ja, um welche Strukturtypen es sich bei den kausalen Interpretationen dieser komplexen Minimalen Theorien handelt. Befinden sich unter den kausal interpretierten komplexen Minimalen Theorien auch Strukturen, die auf kausale Verkettungen hindeuten, muss schliesslich geprüft werden, ob sich die damit einhergehenden indirekten Zusammenhänge angesichts der Intransitivität tatsächlich bestätigen lassen.

Im Fall des Beispiels der Kausalkette aus Abbildung 3.4 handelt es sich bei beiden Minimalen Theorien (3.4) und (3.5) um komplexe Minimale Theorien. Beide einfachen Minimalen Theorien der jeweiligen Minimalen Theorie von $\mathcal{K}_{3,4}$ weisen Faktorüberschneidungen auf: Die einfachen Minimalen Theorien aus (3.4) enthalten beide A sowie B und die einfachen Minimalen Theorien aus (3.5) C .

In einem zweiten Teilschritt wird schliesslich für jede komplexe Minimale Theorie der Minimalen Theorien geprüft, ob darin Folgen $(Z_1, Z_2, \dots)$ von Literalen enthalten sind, die auf eine kausale Verkettung oder einen Zyklus hindeuten und damit unter anderem die Bedingungen (a1) der Definition 2.19 einer Kausalkette oder der Definition 2.22 einer Feedbackstruktur erfüllen (vgl. S. 79 und S. 91). Gibt es solche Folgen, muss geprüft werden, ob auch die Bedingung (a2) der jeweiligen Definition erfüllt ist. Es muss also geprüft werden, ob für jede auf einen indirekten Zusammenhang hindeutende Beziehung zwischen einem Z_i und seiner möglichen Wirkung Z_j aus diesen Folgen gilt, dass Z_i ein Difference-Maker von Z_j ist. Die einfachen Minimalen Theorien, die diese indirekten Kausalzusammenhänge bestätigen, werden zur entsprechenden Minimalen Theorie hinzugefügt und bilden mit letzterer ein System von Minimalen Theorien. Sie sind jedoch von den Minimalen Theorien aus Schritt (II), zu denen sie ergänzend hinzugefügt werden und mit denen sie ein System von Minimalen Theorien bilden, zu unterscheiden. Wie

im Zusammenhang mit dem Kettenproblem bereits verwendet, werden dazu einfache Minimale Theorien, die in diesem letzten Teilschritt zur Bestätigung indirekter Zusammenhänge in einem System aufgenommen werden, mit einem für „indirekt“ stehenden „i“ ausgestattet, das oberhalb des Doppelpfeiles positioniert ist. Das daraus hervorgehende System von Minimalen Theorien, das schliesslich kausal interpretiert wird, entspricht einem mit $\mathcal{K}_F$ verträglichen Kausalmodell ω_i , das als kausal wohlgeformter Hypergraph dargestellt werden kann.

Im Fall des Beispiels mit der Konfigurationstabelle $\mathcal{K}_{3,4}$ beinhaltet die komplexe Minimale Theorie (3.4) keine Folgen von Literalen, die auf indirekte Kausalzusammenhänge hindeuten. A und B können als gemeinsame Ursachen von C und E interpretiert werden. Anders verhält es sich bei der komplexen Minimalen Theorie (3.5), die mit (A, C, E) und (B, C, E) zwei Folgen von Literalen aufweist, die eine kausale Verkettung nahelegen. Der indirekte Zusammenhang zwischen A und E sowie B und E wird durch die einfache Minimale Theorie (3.1), $A + B + D \leftrightarrow E$, bestätigt, nach der A und B auch Difference-Maker von E sind. Diese einfache Minimale Theorie wird mit einem „i“ versehen und zu (3.5) hinzugefügt, woraus sich ein System von Minimalen Theorien ergibt. Insgesamt resultieren somit zwei Systeme von Minimalen Theorien, die kausal interpretiert zwei mit $\mathcal{K}_{3,4}$ verträgliche Kausalmodelle ω_1 und ω_2 darstellen. Das erste System ω_1 gibt eine Common-Cause-Struktur mit A und B als gemeinsame Ursachen von C und E wieder:

$$(A + B \leftrightarrow C) \cdot (A + B + D \leftrightarrow E) \quad (3.4)$$

Das zweite System ω_2 entspricht einer Kausalkette mit C als Zwischenglied:

$$(A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E) \quad (3.5)$$

$$A + B + D \overset{i}{\leftrightarrow} E \quad (3.2)$$

In Abbildung 3.5 sind die beiden resultierenden Kausalmodelle ω_1 und ω_2 als Kausalgraphen dargestellt.

3.3. Kausaler Fehlschluss

Welche Kausalmodelle aus einer Konfigurationstabelle über das analytische Moment hervorgehen, ist massgeblich von der gewählten Faktormenge, den beobachteten Einzelfällen sowie der Güte der erhobenen Daten abhängig. Bezüglich der Güte wurde bereits vorausgesetzt, dass diese in angemessener Form vorliegt (vgl. S. 100). Demzufolge sind die in einer zu analysierenden crisp-set Konfigurationstabelle enthaltenen Kombinationen von an- und abwesenden Faktoren tatsächlich empirisch möglich. Diese Annahme ist notwendig, aber nicht hinreichend dafür,

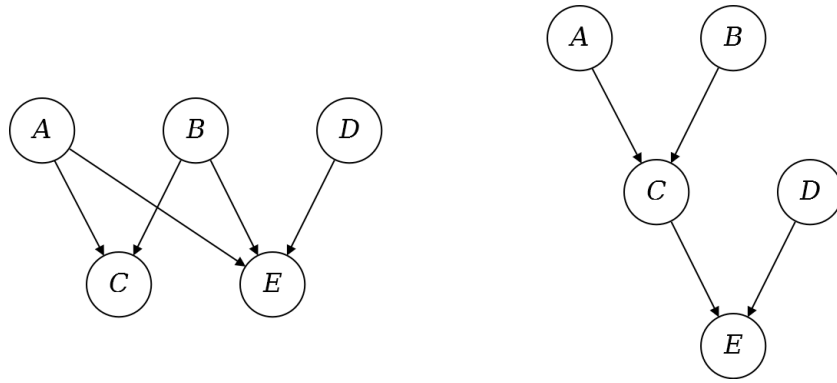


Abbildung 3.5.: Die beiden aus den Konfigurationen $\mathcal{K}_{3,4}$ gefolgerten Kausalmodelle ω_1 und ω_2 .

dass aus den Konfigurationstabellen Minimale Theorien hervorgehen, die die dahinterliegenden kausalen Abhängigkeiten erfassen. Da jede Kausalanalyse relativ zu einer Faktormenge durchgeführt wird, hängt der Erfolg einer Untersuchung auch wesentlich von den gewählten Faktoren ab. Des Weiteren ist ebenfalls entscheidend, welche Konfigurationen dieser Faktoren vorliegen. Denn mit der Annahme der angemessenen Güte der Daten ist unter anderem noch nicht garantiert, dass für jede empirisch mögliche Konfiguration auch tatsächlich Einzelfälle beobachtet und die entsprechenden Konfigurationen somit in die zu analysierende Konfigurationstabelle aufgenommen werden. Je nach Zusammenspiel dieser Charakteristika geben die Konfigurationstabellen im Idealfall alle Informationen wieder, welche die Bildung von Minimalen Theorien derart ermöglichen, dass die erfassten Regelmässigkeiten relativ zur geprüften Faktormenge eindeutig auf die tatsächliche Kausalstruktur zurückgeführt werden können. Im schlechtesten Fall handelt es sich um Informationen, die kausale Fehlschlüsse zur Folge haben. Bei einem kausalen Fehlschluss wird man dazu verleitet, den untersuchten Konfigurationen ein Kausalmodell zugrunde zu legen, das vermeintliche Kausalbeziehungen aufweist, die tatsächlich keine kausalen Abhängigkeiten sind. Ebenfalls einem kausalen Fehlschluss entspricht die Folgerung einer den untersuchten Konfigurationen zugrunde gelegten Struktur, die kausal relevante Literale zu vermeintlichen komplexen Ursachen verbindet, obwohl diese alternativen Ursachen angehören und umgekehrt.

Für die Folgerung von Kausalmodellen, die zuverlässige Aussagen über die Kausalzusammenhänge zwischen Literalen der untersuchten Faktoren zulassen, ist man mindestens darauf angewiesen, dass die analysierten Konfigurationen so miteinander vergleichbar sind, dass kausale Fehlschlüsse ausgeschlossen werden können. Liegt ein kausaler Fehlschluss vor, ist kein aus den untersuchten Konfigurationen $\mathcal{K}_F$ gefolgertes Kausalmodell ω_i , $i \geq 1$, identisch mit der tatsächlichen Kausalstruktur Ω oder einer Teilstruktur von Ω . Anders formuliert bedeutet dies:

Definition 3.2 (Kausaler Fehlschluss)

Seien $\mathcal{K}_F$ Konfigurationen der Faktoren einer Faktormenge F . Sei weiter Ω jene Kausalstruktur, die $\mathcal{K}_F$ tatsächlich zugrunde liegt und $\omega_1, \dots, \omega_n$, $n \geq 1$, die aus $\mathcal{K}_F$ gefolgerten Kausalmodelle.

Ein kausaler Fehlschluss liegt genau dann vor, wenn für jedes ω_i , $1 \leq i \leq n$, mindestens eine der beiden folgenden Bedingungen gilt:

- (a) Das Kausalmodell ω_i weist mindestens ein Literal eines Faktors $A \in F$ als kausal relevant für ein anderes Literal eines Faktors $W \in F$ aus, wobei A gemäss Ω nicht kausal relevant für W ist.
- (b) Das Kausalmodell ω_i verknüpft für ein Literal eines Faktors $W \in F$ kausal relevante Literale von Faktoren aus F disjunktiv beziehungsweise konjunktiv, die gemäss Ω nicht Alternativursachen von W darstellen beziehungsweise keine komplexe Ursache von W bilden.

Wie aus der Definition 3.2 hervorgeht, liegt ein kausaler Fehlschluss erst dann vor, wenn jedes aus einer Konfigurationstabelle gefolgerte Kausalmodell Bedingung (a) oder (b) der Definition 3.2 erfüllt. Das heisst unter anderem, dass aus Konfigurationen, die keine kausalen Fehlschlüsse provozieren, nicht zwangsläufig eindeutig auf die datengenerierende Kausalstruktur geschlossen werden kann. Ambiguitäten sind, wie unter anderem bereits im Zusammenhang mit dem Kettenproblem gesehen, keine Seltenheit und werden akzeptiert, solange eine Kausalanalyse mindestens ein Modell hervorbringt, für das weder (a) noch (b) der Definition 3.2 gilt. Diese Korrektheitsanforderung kann beim kausalen Modellieren durchaus als Standard betrachtet werden (vgl. bspw. Spirtes et al. (2000), S. 81; Kalisch et al. (2012), S. 7). Gehen aus einer Konfigurationstabelle mehrere Kausalmodelle $\omega_1, \dots, \omega_n$, $n \geq 2$, hervor, ist jede der mit den Modellen dargestellte Struktur eine mögliche, die Daten generierende Kausalstruktur. Gemäss dem Prinzip der kausalen Eindeutigkeit wird sich jedoch mit Erhärtung der Erweiterungsklausel durch Folgeanalysen ein Modell als die tatsächliche Struktur herauskristallisieren. Entsprechend wird die aus einer Analyse gefolgerte Kausalhypothese von der Form sein, dass ω_1 oder ω_2 oder ... oder ω_n die datengenerierende Kausalstruktur ist. Eine solche Aussage, die aus einer disjunktiven Verknüpfung von Teilaussagen ω_i , $1 \leq i \leq n$, besteht, ist genau dann wahr, wenn mindestens eine der Teilaussagen ω_i wahr ist. Des Weiteren ist bei der Analyse einer Konfigurationstabelle, die keinen kausalen Fehlschluss generiert, auch nicht ausgeschlossen, dass keine Kausalmodelle abgeleitet werden können. Insgesamt können also Bedingungen, die dazu führen, dass Konfigurationstabellen keine kausalen Fehlschlüsse provozieren, als Minimalanforderung betrachtet werden, die notwendig ist, um kausal schliessen zu können.

3.3.1. Entstehung eines kausalen Fehlschlusses

Kausale Fehlschlüsse können abhängig vom Instantiierungsverhalten von Faktoren auf unterschiedliche Weise entstehen. Um dies zu zeigen, werden zuerst drei Beispiele von Szenarien betrachtet, die Kausalanalysen darstellen und kausale Fehlschlüsse zur Folge haben. Die dabei aufgezeigten Eigenheiten dieser Szenarien werden später zur Illustration verschiedener Punkte benötigt, die für das kausale Schliessen wesentlich sind.

Szenario 1

Stehe die Kausalstruktur aus Abbildung 3.6 im Fokus, wobei D, E, F, G und H die fünf Faktoren der geprüften Faktormenge $\mathbf{P} = \{D, E, F, G, H\}$ sind. Die Faktoren der beiden für H ebenfalls kausal relevanten Literale X und Y seien nicht Teil der geprüften Faktormenge $\mathbf{P}$, was im Kausalgraph mithilfe von gepunkteten Linien erkennbar gemacht werden soll. Gemäss der datengenerierenden Struktur Ω_1 bilden einerseits D und X sowie andererseits E und Y zwei komplexe Ursachen der Wirkung H , die mit den beiden weiteren Ursachen F und G insgesamt vier Alternativursachen besitzt. Die in Abbildung 3.6 dargestellte Konfigurationstabelle $\mathcal{K}_{3,6}$ enthalte jene Konfigurationen, die aus beobachteten Fällen aufbereitet wurden und die für die Untersuchung zur Verfügung stehen.

Relativ zur Faktormenge $\mathbf{P}$ gibt es mit H einen Faktor aus der Konfigurationstabelle $\mathcal{K}_{3,6}$, der (WD1) sowie (WD2) nicht erfüllt und für den das eindeutige RDN-Bikonditional $DE + F + G \leftrightarrow H$ abgeleitet werden kann. Dieses RDN-Bikonditional ist genau für jene Belegungen wahr, die einer Konfiguration aus $\mathcal{K}_{3,6}$ entsprechen. Es stellt somit die einzige (einfache) Minimale Theorie von $\mathcal{K}_{3,6}$ dar. In Abbildung 3.7 ist das entsprechende Kausalmodell ω_1 abgebildet.

Im Kausalmodell ω_1 bilden D und E eine komplexe Ursache von H , obwohl D und E gemäss Ω_1 tatsächlich Teil verschiedener Alternativursachen von H sind. Das Kausalmodell erfüllt somit Bedingung (b) der Definition 3.2. Der Grund für diesen kausalen Fehlschluss sind die beiden unberücksichtigten Faktoren X und Y , die bei gewissen Konfigurationen anwesend sind und bei gewissen nicht. Ein mögliches Instantiierungsverhalten der Faktoren X und Y , das vor dem Hintergrund der tatsächlichen Kausalstruktur zur untersuchten Konfigurationstabelle $\mathcal{K}_{3,6}$ führt, ist in Tabelle 3.8 dargestellt. Der grosse Spaltenabstand zwischen H und X soll dabei kennzeichnen, dass die Faktoren links von dieser Lücke in der geprüften Faktormenge sind und jene rechts davon nicht. Das heisst, die letzten beiden Spalten sind nicht Teil von $\mathcal{K}_{3,6}$. Die Konfigurationen aus Tabelle 3.8 zeigen, dass X nur genau bei der Konfiguration c_4 anwesend ist und sonst nicht. Das führt dazu, dass, wenn entweder nur D (vgl. c_8) oder nur E (c_{12}) gegeben ist, H nicht

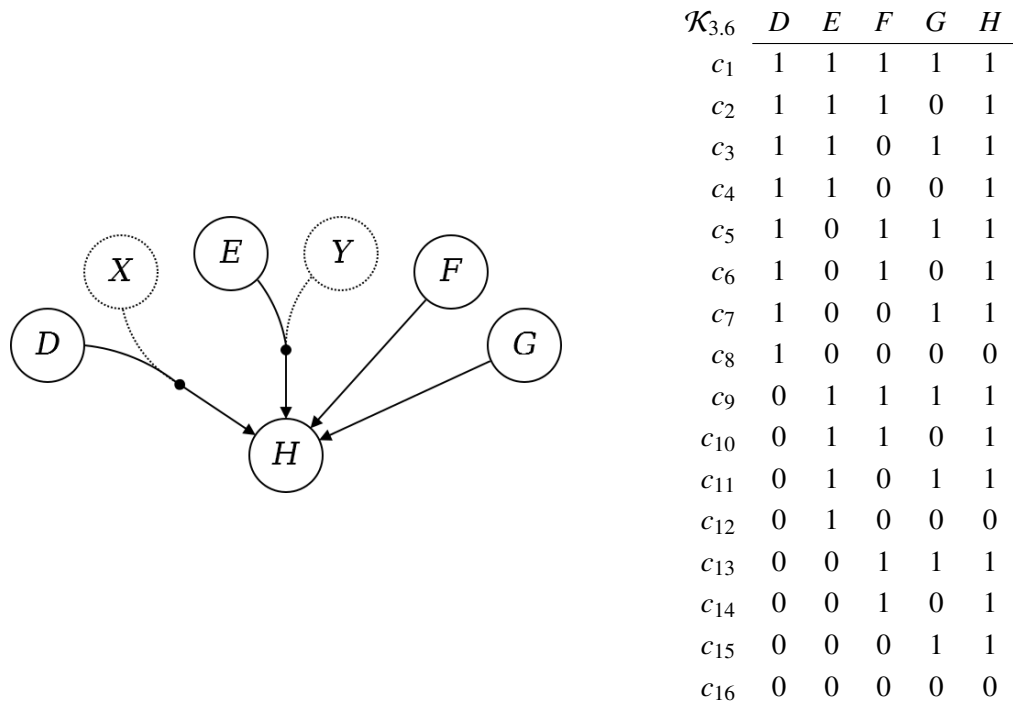


Abbildung 3.6.: Tatsächliche Kausalstruktur Ω_1 zusammen mit den für eine Kausalanalyse zur Verfügung stehenden Konfigurationen $\mathcal{K}_{3,6}$ der Faktoren der Menge $\mathbf{P} = \{D, E, F, G, H\}$.

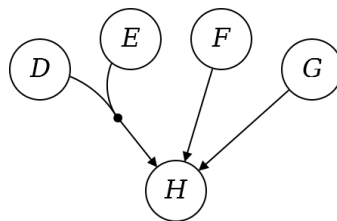


Abbildung 3.7.: Aus den Konfigurationen $\mathcal{K}_{3,6}$ mittels Bestimmung der einfachen Minimalen Theorie $DE + F + G \leftrightarrow H$ abgeleitetes Modell ω_1 .

gegeben ist, und wenn sowohl D als auch E anwesend sind, auch H anwesend ist (c_1 bis c_4). Folglich sind D und E isoliert gemäss $\mathcal{K}_{3,6}$ nicht hinreichend und die Konjunktion DE minimal hinreichend für H .

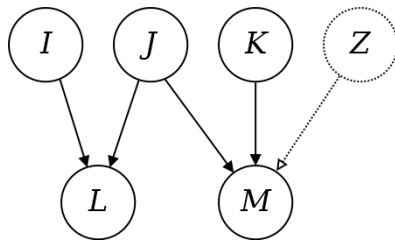
	D	E	F	G	H	X	Y
c_1	1	1	1	1	1	0	0
c_2	1	1	1	0	1	0	0
c_3	1	1	0	1	1	0	0
c_4	1	1	0	0	1	1	0
c_5	1	0	1	1	1	0	0
c_6	1	0	1	0	1	0	0
c_7	1	0	0	1	1	0	0
c_8	1	0	0	0	0	0	0
c_9	0	1	1	1	1	0	0
c_{10}	0	1	1	0	1	0	0
c_{11}	0	1	0	1	1	0	0
c_{12}	0	1	0	0	0	0	0
c_{13}	0	0	1	1	1	0	0
c_{14}	0	0	1	0	1	0	0
c_{15}	0	0	0	1	1	0	0
c_{16}	0	0	0	0	0	0	0

Tabelle 3.8.: Konfigurationen der Faktoren aus $\mathbf{P} = \{D, E, F, G, H\}$ zusammen mit den in $\mathbf{P}$ nicht enthaltenen Faktoren X und Y .

Szenario 2

Seien I, J, K, L und M die Faktoren der geprüften Faktormenge $\mathbf{N} = \{I, J, K, L, M\}$, die zusammen tatsächlich eine Common-Cause-Struktur Ω_2 mit J als gemeinsamer Ursache von L und M bilden. L weist mit I eine und M mit K und Z zwei weitere Alternativursachen auf, wobei der Faktor Z nicht Teil der geprüften Faktormenge ist. Letzteres ist im entsprechenden Kausalgraphen aus Abbildung 3.9 wiederum durch die Verwendung gestrichelter Linien gekennzeichnet. Untersucht werden die ebenfalls in Abbildung 3.9 dargestellten Konfigurationen $\mathcal{K}_{3,9}$.

Ersetzt man I durch A , J durch B , K durch D , L durch C und M durch E stellt sich heraus, dass die Konfigurationstabelle $\mathcal{K}_{3,9}$ identisch mit $\mathcal{K}_{3,4}$ ist (vgl. S. 102). Dementsprechend können



$\mathcal{K}_{3,9}$	I	J	K	L	M
c_1	1	1	1	1	1
c_2	1	1	0	1	1
c_3	1	0	1	1	1
c_4	1	0	0	1	1
c_5	0	1	1	1	1
c_6	0	1	0	1	1
c_7	0	0	1	0	1
c_8	0	0	0	0	0

Abbildung 3.9.: Tatsächliche Kausalstruktur Ω_2 zusammen mit der untersuchten Konfigurationstabelle $\mathcal{K}_{3,9}$ über die Faktormenge $\mathbf{N} = \{I, J, K, L, M\}$.

unter Berücksichtigung dieser Substitutionen ohne Umschweife die beiden Kausalmodelle $\omega_{2,1}$ und $\omega_{2,2}$ aus Abbildung 3.10 gefolgert werden.

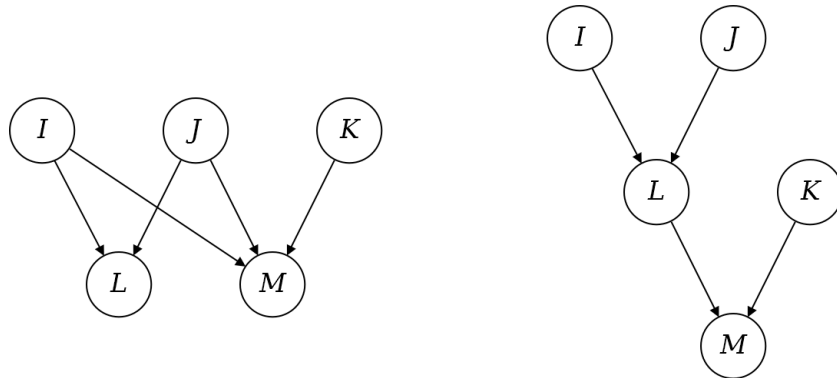


Abbildung 3.10.: Aus den Konfigurationen $\mathcal{K}_{3,9}$ abgeleitete Kausalmodelle.

Es folgen also ausschliesslich Kausalmodelle, die Bedingung (a) der Definition 3.2 erfüllen. In beiden Modellen wird I (direkte oder indirekte) kausale Relevanz für M zugesprochen, obwohl I gemäss der datengenerierenden Kausalstruktur Ω_2 tatsächlich nicht kausal relevant für M ist.

Wie in Szenario 1 entsteht auch hier der kausale Fehlschluss aufgrund des Instantiierungsverhaltens unberücksichtigter Faktoren. Angesichts der tatsächlichen Kausalstruktur kommen die Konfigurationen $\mathcal{K}_{3,9}$ etwa dann zustande, wenn Z ausschliesslich bei der Konfiguration c_4 anwesend ist. In Tabelle 3.11 sind entsprechende, um Z ergänzte Konfigurationen nochmals dargestellt.

	<i>I</i>	<i>J</i>	<i>K</i>	<i>L</i>	<i>M</i>	<i>Z</i>
c_1	1	1	1	1	1	0
c_2	1	1	0	1	1	0
c_3	1	0	1	1	1	0
c_4	1	0	0	1	1	1
c_5	0	1	1	1	1	0
c_6	0	1	0	1	1	0
c_7	0	0	1	0	1	0
c_8	0	0	0	0	0	0

Tabelle 3.11.: Konfigurationen der Faktoren aus $\mathbf{N} = \{I, J, K, L, M\}$ zusammen mit dem in $\mathbf{N}$ nicht enthaltenen Faktor Z .

Szenario 3

Seien A , B und C die drei Faktoren der geprüften Faktormenge $\mathbf{O} = \{A, B, C\}$, für deren Literale gilt, dass sie in keinem kausalen Zusammenhang zueinander stehen. Der Kausalgraph Ω_3 aus Abbildung 3.12 zeigt diese kausale Unabhängigkeit mittels drei Teilstrukturen, die untereinander an keiner Stelle kausal verknüpft sind. Der Einfachheit halber werden dabei die Ursachen von A , B und C jeweils durch ein einzelnes Literal repräsentiert, wobei der entsprechende Faktor nicht Teil der geprüften Faktormenge ist. Für eine Analyse sollen die Konfigurationen $\mathcal{K}_{3,12}$ aus Abbildung 3.12 zur Verfügung stehen.

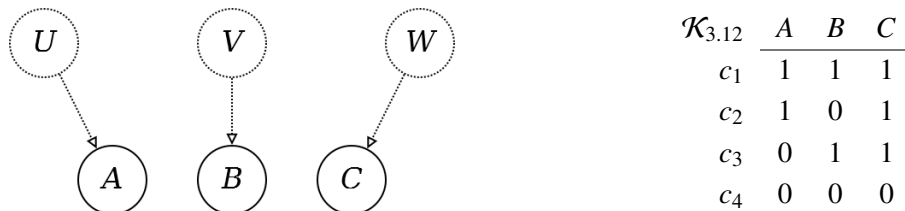


Abbildung 3.12.: Tatsächliche Kausalstruktur Ω_3 zusammen mit der für die Analyse zur Verfügung stehenden Konfigurationstabelle $\mathcal{K}_{3,12}$ über die Faktormenge $\mathbf{O} = \{A, B, C\}$.

Basierend auf $\mathcal{K}_{3,12}$ gibt es relativ zur Faktormenge $\mathbf{O}$ mit C einen Faktor, der weder (WD1) noch (WD2) erfüllt und für den sich mit $A + B \leftrightarrow C$ genau eine einfache Minimale Theorie ableiten lässt. Abbildung 3.13 zeigt das entsprechende, aus $\mathcal{K}_{3,12}$ gefolgerte Kausalmodell ω_3 .

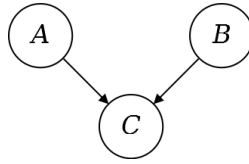


Abbildung 3.13.: Aus den Konfigurationen $\mathcal{K}_{3,12}$ mittels Bestimmung der einfachen Minimalen Theorie $A + B \leftrightarrow C$ abgeleitetes Modell ω_3 .

Weil ω_3 Bedingung (a) der Definition 3.2 erfüllt, liegt ein kausaler Fehlschluss vor. Im Modell werden A und B als Alternativursachen von C ausgewiesen, obwohl die drei Literale tatsächlich in keiner Kausalbeziehung zueinander stehen. Begründen lässt sich dieser Fehlschluss durch die unberücksichtigten Ursachen U , V und W , die zufällige Korrelationen zwischen A , B und C erzeugen, welche die Bedingungen einer einfachen Minimalen Theorie erfüllen. Tabelle 3.14 zeigt die entsprechenden, um die Faktoren U , V und W ergänzten Konfigurationen.

	A	B	C	U	V	W
c_1	1	1	1	1	1	1
c_2	1	0	1	1	0	1
c_3	0	1	1	0	1	1
c_4	0	0	0	0	0	0

Tabelle 3.14.: Konfigurationen der Faktoren aus $\mathbf{O} = \{A, B, C\}$ zusammen mit den Faktoren U , V und W , die nicht in $\mathbf{O}$ enthalten sind.

3.3.2. Störursachen

Die obigen Szenarien zeigen, dass kausale Fehlschlüsse durch unberücksichtigte Ursachen entstehen können, die unbemerkt zwischen verschiedenen Konfigurationen variieren. Der Begriff der Variation, der im Folgenden noch eine wichtige Rolle spielen wird, kann dabei allgemein folgendermassen definiert werden:

Definition 3.3 (Variation)

Ein Literal Z oder ein Konjunktionsterm $Z_1 Z_2 \cdots Z_n$, $n \geq 2$, von Literalen variiert genau dann, wenn Z beziehungsweise $Z_1 Z_2 \cdots Z_n$ bei einer Konfiguration gegeben und bei einer Konfiguration nicht gegeben ist. Dabei gilt, dass $Z_1 Z_2 \cdots Z_n$ genau dann gegeben ist, wenn alle Literale der Konjunktion gegeben sind.

Bei unberücksichtigten Ursachen, die derart unbemerkt variieren, dass ein kausaler Fehlschluss entsteht, kann es sich gemäss dem ersten Szenario einerseits um Ko-Literale von Literalen der

geprüften Faktoren aus $\mathbf{F}$ handeln, die konjunktiv verknüpft eine komplexe Ursache von Literalen der Faktoren aus $\mathbf{F}$ bilden. Andererseits kann es sich, wie im zweiten und dritten Szenario, um kausal relevante Literale handeln, die ohne Literale der geprüften Faktoren aus $\mathbf{F}$ eine Ursache von Literalen der Faktoren aus $\mathbf{F}$ darstellen. Unbemerkt bleibt eine Variation einer solchen unberücksichtigten Ursache eines Literals W eines Faktors aus $\mathbf{F}$ dann, wenn sie nicht aus Literalen $S_1, \dots, S_n$, $n \geq 1$, besteht, die (a) nur indirekt über ein Literal eines Faktors aus $\mathbf{F}$ kausal relevant für W sind oder (b) deren Ursachen Literale von Faktoren aus $\mathbf{F}$ enthalten. Beide Fälle (a) und (b) lassen sich beispielhaft anhand der Kausalkette aus Abbildung 3.4 veranschaulichen (vgl. S. 102). Für den Fall (a) können etwa relativ zur geprüften Faktormenge $\{C, D, E\}$ A und B als unberücksichtigte Ursachen von E betrachtet werden, deren Variationen nicht unbemerkt zu einer Variation von E führen: Eine von ihnen verursachte Variation von E würde sich jeweils im Auftrittsverhalten des berücksichtigten Faktors C manifestieren. Für den Fall (b) wäre beispielsweise C relativ zur Faktormenge $\{A, B, D, E\}$ eine unberücksichtigte Ursache von E , die nicht unbemerkt eine Variation von E erzeugen kann. In diesem Fall liesse sich eine von C verursachte Variation von E im Instantiierungsverhalten der berücksichtigten Faktoren A und B erkennen. A und B sind Ursachen von C , weshalb eine solche Variation von E überhaupt erst durch sie ermöglicht wird.

Weiter handelt es sich bei Fehlschluss generierenden Ursachen nur um solche jener Literale der geprüften Faktoren aus $\mathbf{F}$, die *potenzielle Wirkungen* sind. Potenzielle Wirkungen sind die Literale, von denen relativ zu $\mathbf{F}$ einfache Minimale Theorien abgeleitet werden können und die nicht durch kausales Vorwissen als Wirkungen ausgeschlossen sind. Für solche Literale werden aus den Analysen Kausalmodelle hervorgehen, in denen sie als Wirkungen enthalten sind und über die somit bezüglich ihrer Ursachen falsch geschlossen werden kann. Die Menge aller in einer geprüften Faktormenge $\mathbf{F}$ enthaltenen Faktoren, deren Literale potenzielle Wirkungen sind, wird in der Folge speziell ausgezeichnet, indem die entsprechende Faktormenge um ein für „Wirkung“ stehendes „w“ ergänzt und als $\mathbf{F}^w$, $\mathbf{F}^w \subseteq \mathbf{F}$, notiert wird:

Definition 3.4 (Potenzielle Wirkung)

Sei $\mathcal{K}_{\mathbf{F}}$ eine Konfigurationstabelle über eine Faktormenge $\mathbf{F}$. Ein Literal W eines Faktors $W \in \mathbf{F}$ ist genau dann eine potenzielle Wirkung und der Faktor W damit Element der Menge $\mathbf{F}^w \subseteq \mathbf{F}$, wenn es relativ zu $\mathcal{K}_{\mathbf{F}}$ einfache Minimale Theorien von W gibt und W nicht durch kausales Vorwissen als Wirkung ausgeschlossen ist.

Eine mögliche, Fehlschluss produzierende Variation beschränkt sich auf unberücksichtigte Ursachen dieser potenziellen Wirkungen. Zwar besitzen gemäss dem Kausalitätsprinzip auch Literale von geprüften Faktoren aus $\mathbf{F}$, die keine potenziellen Wirkungen sind, Ursachen, die darüber hinaus variieren können und unberücksichtigt bleiben. So gibt es beispielsweise im ersten Szenario ebenfalls variierende und unberücksichtigte Ursachen von D . Derartige unberücksichtigte

Ursachen von Literalen, die nicht potenzielle Wirkungen sind, führen deshalb nicht zu einem Fehlschluss, weil für entsprechende Literale wie D gar keine einfachen Minimalen Theorien abgeleitet werden. Die Variationen solcher unberücksichtigten Ursachen sind sogar ausdrücklich erwünscht, sofern die Ursachen für andere Literale von Faktoren der geprüften Menge nur indirekt über D kausal relevant sind, da diese Variationen die nötige Diversität in der Konfigurationstabelle erzeugen, sodass sich die Difference-Maker-Qualität von D für H offenbaren kann. Damit zeigt sich, dass die Bestimmung der Fehlschluss provozierenden Ursachen relativ ist und stark von der Menge $\mathbf{F}^w$ abhängt. Eine unberücksichtigte Ursache X eines Literals W des Faktors $W \in \mathbf{F}$ kann im Prinzip alleine deshalb bei einer Kausalanalyse problematisch sein und bei einer zweiten Untersuchung derselben Konfigurationstabelle nicht, weil sich die beiden Analysen einzig derart voneinander unterscheiden, dass in der ersten Analyse W als potenzielle Wirkung betrachtet ($W \in \mathbf{F}^w$) und in der zweiten aufgrund von kausalem Vorwissen als solche ausgeschlossen wird ($W \notin \mathbf{F}^w$). Geht man beispielsweise beim zweiten Szenario davon aus, dass kausales Vorwissen M als Wirkung ausschliesst, wird aus der Analyse von $\mathcal{K}_{3,9}$ nur die einfache Minimale Theorie $I + J \leftrightarrow L$ kausal interpretiert. Das entsprechende Kausalmodell ist eine Teilstruktur der tatsächlichen Kausalstruktur Ω_2 . Die Variation der unberücksichtigten Ursache Z kann vor diesem Hintergrund keinen kausalen Fehlschluss provozieren und stellt demnach kein Problem dar.

Unberücksichtigte Ursachen von potenziellen Wirkungen, die unbemerkt variieren können, haben das Potenzial, einen kausalen Fehlschluss zu erzeugen. Literale und Konjunktionen von Literalen mit obigen Fehlschluss provozierenden Eigenschaften werden als *potenzielle Störursachen* bezeichnet und allgemein folgendermassen definiert:

Definition 3.5 (Potenzielle Störursachen)

Sei $\mathbf{F}$ eine Faktormenge und W eine potenzielle Wirkung, das heisst, der Faktor W ist Element der Menge $\mathbf{F}^w \subseteq \mathbf{F}$. Sei weiter $\mathbf{G}$ die Menge von Faktoren, deren Literale kausal relevant für W und die nicht in $\mathbf{F}$ enthalten sind. Relativ zu $\mathbf{F}$ ist eine potenzielle Störursache von W ein aus Literalen von Faktoren aus $\mathbf{G}$ bestehender Konjunktionsterm $\sigma = S_1 \cdots S_n$, $n \geq 1$, für den gilt:

- (a) σ bildet zusammen mit oder ohne Literale von Faktoren aus $\mathbf{F}$ eine komplexe Ursache von W .
- (b) Für jedes Literal S_i , $1 \leq i \leq n$, des Konjunktionsterms σ gilt:
 - (b1) S_i ist relativ zu $\mathbf{F} \cup \{S_i\}$ direkt kausal relevant für W ,
 - (b2) die Literale der Faktoren aus $\mathbf{F}$ sind nicht kausal relevant für S_i .

Es gilt hervorzuheben, dass in den Szenarien 1 bis 3 die abgeleiteten Minimalen Theorien auch tatsächlich Minimale Theorien der entsprechenden Konfigurationstabellen sind, weshalb auch jedes adäquate Schlussverfahren diese Minimalen Theorien ausgeben muss. Daher können diese Fehlschlüsse nicht auf Fehler im analytischen Moment zurückgeführt werden, sondern sind das Produkt von mangelhaften Daten. Weil weiter die Güte angemessen ist, wonach ebenfalls Mess- oder Kalibrierungsfehler ausgeschlossen werden können, sind es die variierenden potenziellen Störursachen, die diese Datenmängel erzeugen.

3.4. Homogenität

Kann nicht ausgeschlossen werden, dass potenzielle Störursachen beliebig variieren, lässt sich auch das durch mangelhafte Daten entstehende Risiko eines kausalen Fehlschlusses nicht ausschliessen. Zu jeder Konfigurationstabelle könnten alternative Kausalmodelle konstruiert werden, die den aus der Kausalanalyse resultierenden Kausalmodellen widersprechen und letztere somit infrage stellen würden. Die Plausibilität, dass den Daten tatsächlich ein solches alternatives Kausalmodell zugrunde liegt, liesse sich dabei jeweils durch potenzielle Störursachen rechtfertigen, die, wie in den obigen drei Szenarien, zwischen verschiedenen Konfigurationen derart variieren, dass die entsprechenden Fehlschlüsse produziert werden. Damit solche Einwände, die eine Kausalanalyse gehaltlos machen, ausgeschlossen werden können, muss die Variation von potenziellen Störursachen durch gewisse Bedingungen eingeschränkt werden.

3.4.1. Homogenisierung der potenziellen Störursachen

Sei $\mathcal{K}_F$ die untersuchte Konfigurationstabelle über eine Faktormenge F und F^w die Menge der Faktoren aus F , deren Literale potenzielle Wirkungen sind. Eine naheliegende Forderung an die potenziellen Störursachen zur Verhinderung der obigen Fehlschlüsse ist, deren Variation derart einzuschränken, dass ihr kausaler Einfluss auf die potenziellen Wirkungen bei allen Konfigurationen derselbe ist (vgl. bspw. Graßhoff & May (2001), S. 203 ff.; Baumgartner (2006), S. 192 ff.). Konkret wird verlangt, dass die potenziellen Störursachen der Literale von Faktoren aus F^w *homogenisiert* sind und $\mathcal{K}_F$ damit *homogen* ist. Eine mögliche Variante dieser Homogenität ist in Definition 3.6 gegeben (vgl. u. a. Baumgartner (2009), S. 80). Diese Variante wird deshalb als „strikte“ Homogenität bezeichnet, weil in der Folge eine schwächere Form davon präsentiert wird (vgl. S. 128).

Definition 3.6 (Strikte Homogenität)

Sei $\mathcal{K}_F$ eine Konfigurationstabelle über die Faktormenge F und W ein beliebiger Faktor

aus $\mathbf{F}^w$. Sei weiter der aus Literalen $S_1, \dots, S_k$, $k \geq 1$, bestehende Konjunktionsterm σ eine beliebige potenzielle Störursache von W . σ ist in $\mathcal{K}_F$ genau dann strikt homogenisiert, wenn gilt: σ ist bei einer Konfiguration genau dann anwesend, wenn σ auch bei allen anderen Konfigurationen anwesend ist. $\mathcal{K}_F$ ist genau dann strikt homogen, wenn jede potenzielle Störursache jeder potenziellen Wirkung strikt homogenisiert ist.

Ist eine potenzielle Störursache strikt homogenisiert, variiert sie nicht gemäss Definition 3.3. Die Absicht der Homogenisierung aller potenziellen Störursachen der potenziellen Wirkungen ist, dass nicht unberücksichtigte Ursachen allfällige Variationen im Instantiierungsverhalten eines Faktors $W \in \mathbf{F}^w$ hervorrufen und, sofern Regelmässigkeiten erkennbar sind, diese auf tatsächliche Ursachen von W , deren Faktoren in $\mathbf{F}$ enthalten sind, zurückgeführt werden können.

Am ersten Szenario veranschaulicht, ergeben sich aus der strikten Homogenisierung der potenziellen Störursachen X und Y gemäss Definition 3.6 vier mögliche Konfigurationstabellen, die in Tabelle 3.15 zusammengestellt sind. Für alle vier Konfigurationstabellen wird gemäss dem ersten Schritt des analytischen Moments alleine der Faktor H nicht durch (WD1) und (WD2) ausgeschlossen. Für jede der vier Konfigurationstabellen lässt sich ein einziges RDN-Bikonditional von H ableiten. Bei allen vier RDN-Bikonditionalen handelt es sich um einfache Minimale Theorien von H , wobei (3.6) aus $\mathcal{K}_{3.15a}$ folgt, (3.7) aus $\mathcal{K}_{3.15b}$, (3.8) aus $\mathcal{K}_{3.15c}$ und (3.9) aus $\mathcal{K}_{3.15d}$:

$$F + G \leftrightarrow H \quad (3.6)$$

$$E + F + G \leftrightarrow H \quad (3.7)$$

$$D + F + G \leftrightarrow H \quad (3.8)$$

$$D + E + F + G \leftrightarrow H \quad (3.9)$$

Bei der Bildung der dazugehörigen Kausalmodelle zeigt sich, dass jedes Modell einer Teilstruktur der tatsächlichen Kausalstruktur Ω_1 aus Abbildung 3.6 entspricht (vgl. S. 113). Es wird kein kausaler Fehlschluss mehr erzeugt.

Für Szenario 2 ergeben sich unter der Voraussetzung, dass die potenzielle Störursache Z gemäss Definition 3.6 strikt homogenisiert ist, die Konfigurationen aus Tabelle 3.16.

Bei der strikt homogenen Konfigurationstabelle $\mathcal{K}_{3.16a}$ können mit L und M zwei Faktoren ermittelt werden, für die weder (WD1) noch (WD2) gilt. Für L und M existiert jeweils genau ein RDN-Bikonditional. Werden die beiden RDN-Bikonditionale konjunktiv verknüpft, ergibt sich die folgende (komplexe) Minimale Theorie von $\mathcal{K}_{3.16a}$:

$$(I + J \leftrightarrow L) \cdot (J + K \leftrightarrow M) \quad (3.10)$$

$\mathcal{K}_{3.15a}$	D	E	F	G	H
c_1	1	1	1	1	1
c_2	1	1	1	0	1
c_3	1	1	0	1	1
c_4	1	1	0	0	0
c_5	1	0	1	1	1
c_6	1	0	1	0	1
c_7	1	0	0	1	1
c_8	1	0	0	0	0
c_9	0	1	1	1	1
c_{10}	0	1	1	0	1
c_{11}	0	1	0	1	1
c_{12}	0	1	0	0	0
c_{13}	0	0	1	1	1
c_{14}	0	0	1	0	1
c_{15}	0	0	0	1	1
c_{16}	0	0	0	0	0

(a) X und Y sind konstant abwesend.

$\mathcal{K}_{3.15b}$	D	E	F	G	H
c_1	1	1	1	1	1
c_2	1	1	1	0	1
c_3	1	1	0	1	1
c_4	1	1	0	0	1
c_5	1	0	1	1	1
c_6	1	0	1	0	1
c_7	1	0	0	1	1
c_8	1	0	0	0	0
c_9	0	1	1	1	1
c_{10}	0	1	1	0	1
c_{11}	0	1	0	1	1
c_{12}	0	1	0	0	1
c_{13}	0	0	1	1	1
c_{14}	0	0	1	0	1
c_{15}	0	0	0	1	1
c_{16}	0	0	0	0	0

(b) X ist konstant abwesend, Y konstant anwesend.

$\mathcal{K}_{3.15c}$	D	E	F	G	H
c_1	1	1	1	1	1
c_2	1	1	1	0	1
c_3	1	1	0	1	1
c_4	1	1	0	0	1
c_5	1	0	1	1	1
c_6	1	0	1	0	1
c_7	1	0	0	1	1
c_8	1	0	0	0	1
c_9	0	1	1	1	1
c_{10}	0	1	1	0	1
c_{11}	0	1	0	1	1
c_{12}	0	1	0	0	0
c_{13}	0	0	1	1	1
c_{14}	0	0	1	0	1
c_{15}	0	0	0	1	1
c_{16}	0	0	0	0	0

(c) X ist konstant anwesend, Y konstant abwesend.

$\mathcal{K}_{3.15d}$	D	E	F	G	H
c_1	1	1	1	1	1
c_2	1	1	1	0	1
c_3	1	1	0	1	1
c_4	1	1	0	0	1
c_5	1	0	1	1	1
c_6	1	0	1	0	1
c_7	1	0	0	1	1
c_8	1	0	0	0	1
c_9	0	1	1	1	1
c_{10}	0	1	1	0	1
c_{11}	0	1	0	1	1
c_{12}	0	1	0	0	1
c_{13}	0	0	1	1	1
c_{14}	0	0	1	0	1
c_{15}	0	0	0	1	1
c_{16}	0	0	0	0	0

(d) X und Y sind konstant anwesend.

Tabelle 3.15.: Vor dem Hintergrund der Kausalstruktur aus Abbildung 3.6 strikt homogenisierte Konfigurationstabellen.

$\mathcal{K}_{3.16a}$	I	J	K	L	M
c_1	1	1	1	1	1
c_2	1	1	0	1	1
c_3	1	0	1	1	1
c_4	1	0	0	1	0
c_5	0	1	1	1	1
c_6	0	1	0	1	1
c_7	0	0	1	0	1
c_8	0	0	0	0	0

(a) Z ist konstant abwesend.

$\mathcal{K}_{3.16b}$	I	J	K	L	M
c_1	1	1	1	1	1
c_2	1	1	0	1	1
c_3	1	0	1	1	1
c_4	1	0	0	1	1
c_5	0	1	1	1	1
c_6	0	1	0	1	1
c_7	0	0	1	0	1
c_8	0	0	0	0	1

(b) Z ist konstant anwesend.

Tabelle 3.16.: Vor dem Hintergrund der Kausalstruktur aus Abbildung 3.9 strikt homogenisierte Konfigurationstabellen.

Bei der strikt homogenen Konfigurationstabelle $\mathcal{K}_{3.16b}$ lässt sich nur L nicht durch (WD1) und (WD2) ausschliessen und schliesslich die einfache Minimale Theorie $I + J \leftrightarrow L$ ableiten. Insgesamt entsprechen die sich daraus ergebenden Kausalmodelle relativ zu $\mathbf{N} = \{I, J, K, L, M\}$ der datengenerierenden Kausalstruktur aus Abbildung 3.9 oder einem Teil davon (vgl. S. 115). Ein kausaler Fehlschluss wird nicht mehr produziert.

Beim dritten Szenario ergeben sich bei einer strikten Homogenisierung der potenziellen Störursachen von C gemäss Definition 3.6 die Konfigurationen $\mathcal{K}_{3.17a}$ und $\mathcal{K}_{3.17b}$ aus Tabelle 3.17. Bei beiden strikt homogenen Konfigurationstabellen $\mathcal{K}_{3.17a}$ und $\mathcal{K}_{3.17b}$ kann die Analyse frühzeitig abgebrochen werden, weil für alle Faktoren entweder (WD1) oder (WD2) gilt. Es können keine RDN-Bikonditionale und damit auch keine Minimalen Theorien gefolgert werden. Aus den beiden Konfigurationstabellen ergeben sich keine Kausalmodelle, womit auch keine kausalen Fehlschlüsse provoziert werden können.

$\mathcal{K}_{3.17a}$	A	B	C
c_1	1	1	0
c_2	1	0	0
c_3	0	1	0
c_4	0	0	0

(a)

$\mathcal{K}_{3.17b}$	A	B	C
c_1	1	1	1
c_2	1	0	1
c_3	0	1	1
c_4	0	0	1

(b)

Tabelle 3.17.: Vor dem Hintergrund der kausalen Unabhängigkeit zwischen A , B und C strikt homogenisierte Konfigurationstabellen.

$\mathcal{K}_{3.18a}$					$\mathcal{K}_{3.18b}$	A	B	C		$\mathcal{K}_{3.18c}$	A	B	C
					c_1	1	1	0		c_1	1	0	0
					c_2	1	0	0		c_2	0	1	0
					c_3	0	1	0		c_3	0	0	0
					c_4	0	0	0		c_4	0	0	1
					c_5	0	0	1		c_5	0	1	1
	A	B	C										
c_1	1	1	1										
(a)					(b)					(c)			

Tabelle 3.18.: Weitere Beispiele strikt homogener Konfigurationstabellen für Szenario 3.

Die strikte Homogenisierung der potenziellen Störursachen von C aus dem dritten Szenario und die damit verbundene Konsequenz, dass basierend auf den Konfigurationen $\mathcal{K}_{3.17a}$ und $\mathcal{K}_{3.17b}$ keine potenziellen Wirkungen bestimmt werden können, macht auf den als *triviale Homogenität* bezeichneten Grenzfall aufmerksam, für den $\mathbf{F}^w = \emptyset$ gilt. Gemäss Definition 3.6 ist eine Konfigurationstabelle strikt homogen, wenn jede potenzielle Störursache der Literale von Faktoren aus $\mathbf{F}^w$ homogenisiert ist. Diese Bedingung ist trivialerweise für all jene Konfigurationstabellen erfüllt, für welche $\mathbf{F}^w = \emptyset$ gilt. Daraus folgt, dass unter anderem auch die zum dritten Szenario gehörigen Konfigurationstabellen aus Tabelle 3.18 strikt homogen sind. Alle Faktoren jeder Konfigurationstabelle aus Tabelle 3.18 erfüllen entweder Bedingung (WD1) oder (WD2) des ersten Hauptschrittes des analytischen Moments (vgl. S. 105). Für keine der Konfigurationstabellen werden einfache Minimale Theorien und damit auch keine Kausalmodelle abgeleitet.

Die Homogenitätsbedingung soll so wenig restriktiv formuliert werden wie möglich. Dadurch wird nicht nur der Raum der analysierbaren Konfigurationstabellen grösser. Auch im Hinblick auf eine praktische Umsetzung ist die Formulierung von Anforderungen sinnvoll, die das Instantiierungsverhalten von Faktoren so wenig wie möglich einschränkt. Denn je strenger die Bedingungen an das Instantiierungsverhalten der in einer Kausalanalyse involvierten Faktoren sind, desto unwahrscheinlicher wird es, solche Bedingungen in der Praxis vorzufinden beziehungsweise in einer Studie zu realisieren. Das Bestreben, die Anforderungen in dieser Hinsicht möglichst minimal zu halten, zeigt sich bei strikter Homogenität einerseits dadurch, dass alle Ursachen von Literalen, die keine potenziellen Wirkungen sind, beliebig variieren können. Andererseits ist dieses Bestreben auch bei den potenziellen Störursachen sichtbar, die tatsächlich homogenisiert werden müssen. Denn ist eine potenzielle Störursache strikt homogenisiert, bedeutet dies nicht, dass jedes darin enthaltene kausal relevante Literal bei allen Konfigurationen gleich gegeben beziehungsweise nicht gegeben sein muss. Innerhalb einer potenziellen Störursache sind durchaus Variationen zwischen Konfigurationen möglich. Nimmt man beispielsweise obiges Szenario 2 und ersetzt die potenzielle Störursache Z durch die unberücksichtigte komplexe Ursache Z_1Z_2 gemäss Abbildung 3.19, ist die potenzielle Störursache Z_1Z_2 bereits dann

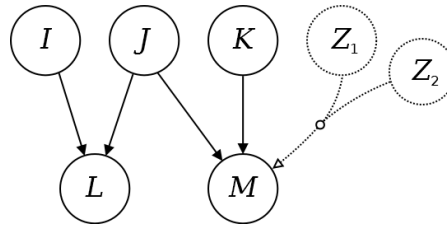


Abbildung 3.19.: Kausalstruktur aus Abbildung 3.9 des zweiten Szenarios, wobei die potenzielle Störursache Z durch eine komplexe Ursache Z_1Z_2 ersetzt wird.

strikt homogenisiert, wenn bei jeder Konfiguration mindestens einer der beiden Faktoren Z_1 oder Z_2 abwesend ist. Bei einer Konfiguration kann es sich dabei beispielsweise um beide Faktoren handeln, bei einer anderen nur um Z_1 oder nur um Z_2 .

Trotz dieser Freiheiten, die einer potenziellen Störursache bezüglich ihres Auftrittsverhaltens bei einer Homogenisierung gemäss Definition 3.6 bleiben, lässt sich Homogenität noch schwächer formulieren, ohne dass Fehlschlüsse möglich werden. Sichtbar wird dies beispielsweise bei Szenario 1: Relativ zu den Faktoren der Menge $\mathbf{P} = \{D, E, F, G, H\}$ sind die Konfigurationen aus Tabelle 3.20a identisch mit den Konfigurationen $\mathcal{K}_{3.15a}$ aus Tabelle 3.15 und jene aus Tabelle 3.20b identisch mit den Konfigurationen $\mathcal{K}_{3.15c}$. Obwohl also die potenzielle Störursache X jeweils nicht strikt homogenisiert ist, resultiert aus den Konfigurationen der Tabelle 3.20 relativ zu $\mathbf{P}$ kein kausaler Fehlschluss.

Der entscheidende Unterschied zwischen dem zu einem Fehlschluss führenden Auftrittsverhalten von X bei der Konfigurationstabelle 3.8 (vgl. S. 114) und den unproblematischen Variationen von X in Tabelle 3.20 lässt sich mithilfe des kausalen Effekts festmachen, der mit der Variation der potenziellen Störursache einhergeht. Konkret haben die Variationen von X aus Tabelle 3.20, im Gegensatz zur Variation von X aus Tabelle 3.8, keinen erkennbaren Effekt auf die Konfigurationen über $\mathbf{P} = \{D, E, F, G, H\}$, der allein auf X zurückgeht. Etwas präziser ausgedrückt führt bei keinen zwei Konfigurationen der Faktoren aus $\mathbf{P}$ der Tabellen 3.20a und 3.20b allein die Variation von X zu einer Variation von H . Derartige Variationen von potenziellen Störursachen werden in der Folge auch kurz als *nicht ausschlaggebend* bezeichnet.

In der Konfigurationstabelle 3.20a variiert zwar X zwischen der Konfiguration c_{16} und den restlichen Konfigurationen, dies hat jedoch keine Variation von H zur Folge. Die Variation von X bleibt also ohne Einfluss auf H und ist somit auch nicht ausschlaggebend. Bei der Konfigurationstabelle 3.20b variiert X ebenfalls nicht ausschlaggebend, da es nicht alleine die Variation von X ist, die zu einer Variation von H führt. Genauso ist es jeweils die Variation von D , die mitverantwortlich gemacht werden muss (vgl. c_4 und c_{12} sowie c_8 und c_{16}). Die Variationen von

	<i>D</i>	<i>E</i>	<i>F</i>	<i>G</i>	<i>H</i>	<i>X</i>	<i>Y</i>		<i>D</i>	<i>E</i>	<i>F</i>	<i>G</i>	<i>H</i>	<i>X</i>	<i>Y</i>
<i>c</i> ₁	1	1	1	1	1	0	0	<i>c</i> ₁	1	1	1	1	1	0	0
<i>c</i> ₂	1	1	1	0	1	0	0	<i>c</i> ₂	1	1	1	0	1	0	0
<i>c</i> ₃	1	1	0	1	1	0	0	<i>c</i> ₃	1	1	0	1	1	0	0
<i>c</i> ₄	1	1	0	0	0	0	0	<i>c</i> ₄	1	1	0	0	1	1	0
<i>c</i> ₅	1	0	1	1	1	0	0	<i>c</i> ₅	1	0	1	1	1	0	0
<i>c</i> ₆	1	0	1	0	1	0	0	<i>c</i> ₆	1	0	1	0	1	0	0
<i>c</i> ₇	1	0	0	1	1	0	0	<i>c</i> ₇	1	0	0	1	1	0	0
<i>c</i> ₈	1	0	0	0	0	0	0	<i>c</i> ₈	1	0	0	0	1	1	0
<i>c</i> ₉	0	1	1	1	1	0	0	<i>c</i> ₉	0	1	1	1	1	0	0
<i>c</i> ₁₀	0	1	1	0	1	0	0	<i>c</i> ₁₀	0	1	1	0	1	0	0
<i>c</i> ₁₁	0	1	0	1	1	0	0	<i>c</i> ₁₁	0	1	0	1	1	0	0
<i>c</i> ₁₂	0	1	0	0	0	0	0	<i>c</i> ₁₂	0	1	0	0	0	0	0
<i>c</i> ₁₃	0	0	1	1	1	0	0	<i>c</i> ₁₃	0	0	1	1	1	0	0
<i>c</i> ₁₄	0	0	1	0	1	0	0	<i>c</i> ₁₄	0	0	1	0	1	0	0
<i>c</i> ₁₅	0	0	0	1	1	0	0	<i>c</i> ₁₅	0	0	0	1	1	0	0
<i>c</i> ₁₆	0	0	0	0	0	1	0	<i>c</i> ₁₆	0	0	0	0	0	0	0

(a)

(b)

Tabelle 3.20.: Zu Szenario 1 gehörige Konfigurationstabellen, aus denen relativ zu $\mathbf{P} = \{D, E, F, G, H\}$ kein kausaler Fehlschluss resultiert, obwohl die potenzielle Störursache X jeweils nicht strikt homogenisiert ist.

X werden somit durch jene von D überdeckt und die sich daraus ergebenden Regelmässigkeiten auf das positive Literal des berücksichtigten Faktors D zurückgeführt, das kausal relevant für H ist.

Im Gegensatz dazu variiert X angesichts der Konfigurationen c_4 und c_8 bei der Konfigurationstabelle 3.8 ausschlaggebend, das heisst, zwischen c_4 und c_8 führt allein die Variation von X dazu, dass H variiert. D ist zwar als Teil der komplexen Ursache DX von H notwendig dafür, dass bei c_4 H gegeben ist. D ist jedoch in beiden Konfigurationen konstant anwesend, weshalb die Variation von H alleine auf X zurückgeführt werden kann. Auch die zwischen c_4 und c_8 bestehende Variation des für H kausal relevanten Literals E ändert nichts daran. Diese Variation bleibt ohne Einfluss auf die Variation von H , da dazu das (unberücksichtigte) Ko-Literal Y gegeben sein müsste.⁴ Letzteres ist aber nicht der Fall. Insgesamt kann damit bei der Konfigurationstabelle 3.8 ausschliesslich die potenzielle Störursache X für gewisse Variationen von H

⁴Wäre übrigens Y bei c_4 gegeben, könnte der kausale Fehlschluss trotzdem nicht verhindert werden, zumal die untersuchte Konfigurationstabelle unverändert bliebe. Mit den Konfigurationen c_4 und c_{12} würde nun die potenzielle Störursache Y ausschlaggebend variieren. Keine ausschlaggebenden Variationen der potenziellen Störursachen X und Y wären bspw. dann gegeben, wenn X bei c_4 gegeben und Y sowohl bei c_4 wie auch bei c_8 anwesend wäre.

	<i>I</i>	<i>J</i>	<i>K</i>	<i>L</i>	<i>M</i>	<i>Z</i>		<i>I</i>	<i>J</i>	<i>K</i>	<i>L</i>	<i>M</i>	<i>Z</i>
c_1	1	1	1	1	1	1	c_1	1	1	1	1	1	0
c_2	1	1	0	1	1	0	c_2	1	1	0	1	1	0
c_3	1	0	1	1	1	1	c_3	1	0	1	1	1	0
c_4	1	0	0	1	0	0	c_4	1	0	0	1	1	1
c_5	0	1	1	1	1	1	c_5	0	1	1	1	1	0
c_6	0	1	0	1	1	0	c_6	0	1	0	1	1	0
c_7	0	0	1	0	1	1	c_7	0	0	1	0	1	0
c_8	0	0	0	0	0	0	c_8	0	0	0	0	1	1

(a)

(b)

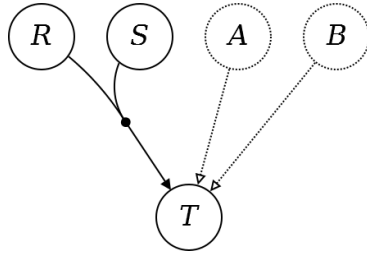
Tabelle 3.21.: Zu Szenario 2 gehörige Konfigurationen, aus denen relativ zu $\mathbf{N} = \{I, J, K, L, M\}$ kein kausaler Fehlschluss resultiert. Die potenzielle Störursache Z ist nicht strikt homogenisiert.

verantwortlich gemacht werden. Als Konsequenz ergibt sich ein verzerrtes Wirkungsbild mit abgeleiteten Regelmässigkeiten, die relativ zur geprüften Faktormenge einen kausalen Fehlschluss produzieren.

Für die potenzielle Störursache Z in Tabelle 3.11 aus dem zweiten Szenario kann analog argumentiert werden (vgl. S. 116). M variiert zwischen den Konfigurationen c_4 und c_8 alleine deshalb, weil Z variiert. Zum Vergleich lassen sich erneut Konfigurationen anführen, bei denen Z variiert und trotzdem keine kausalen Fehlschlüsse folgen, weil die Variation von Z nicht ausschlaggebend ist. Entsprechende Beispiele sind in Tabelle 3.21 dargestellt, wobei die Konfigurationstabelle 3.21a relativ zu $\mathbf{N} = \{I, J, K, L, M\}$ identisch mit den Konfigurationen $\mathcal{K}_{3.16a}$ aus Tabelle 3.16 ist und die Konfigurationen aus Tabelle 3.21b identisch mit $\mathcal{K}_{3.16b}$ sind.

Es bleibt eine zusätzliche Verfeinerung der ausschlaggebenden Variation einer potenziellen Störursache, die es zur Verhinderung von kausalen Fehlschlüssen einzubauen gilt. Dieser Verfeinerung zufolge muss die ausschlaggebende Variation auf die geprüfte Faktormenge und die Faktoren, deren Literale in der potenziellen Störursache enthalten sind, relativiert werden. Das heisst, Variationen von obiger potenzieller Störursache X beziehungsweise Z sind dann unproblematisch, wenn es keine zwei Konfigurationen gibt, bei denen X relativ zu $\mathbf{P} \cup \{X\}$ beziehungsweise Z relativ zu $\mathbf{N} \cup \{Z\}$ alleine zu einer Variation der entsprechenden potenziellen Wirkungen führt. Dieser relativierende Zusatz ist in den obigen Beispielen nicht konkret zur Anwendung gekommen. Das ändert sich, wenn die gleiche potenzielle Wirkung mehrere alternative potenzielle Störursachen besitzt. Würde man die alleinige Variation nicht zur geprüften Faktormenge und

Die entsprechende Konfigurationstabelle über $\mathbf{P} = \{D, E, F, G, H\}$, die identisch mit $\mathcal{K}_{3.15b}$ wäre, hätte keinen kausalen Fehlschluss zur Folge.



	<i>R</i>	<i>S</i>	<i>T</i>	<i>A</i>	<i>B</i>
<i>c</i> ₁	1	1	1	1	1
<i>c</i> ₂	1	0	1	1	1
<i>c</i> ₃	0	1	1	1	1
<i>c</i> ₄	0	0	0	0	0

Abbildung 3.22.: Konfigurationstabelle, die relativ zur geprüften Faktormenge $\mathbf{M} = \{R, S, T\}$ einen kausalen Fehlschluss provoziert.

zu den Faktoren, deren Literale potenzielle Störursachen bilden, in Beziehung setzen, würde eine potenzielle Störursache bereits dann nicht ausschlaggebend variieren, wenn eine andere potenzielle Störursache auf dieselbe Art variiert. Ein entsprechendes Beispiel ist in Abbildung 3.22 dargestellt. Gemäss der Konfigurationstabelle über $\mathbf{M} = \{R, S, T\}$ wird daraus die eindeutige einfache Minimale Theorie $R + S \leftrightarrow T$ abgeleitet, was einen kausalen Fehlschluss provoziert. Relativ zu $\{R, S, T, A, B\}$ führt weder *A* noch *B* alleine zu einer Variation von *T*, zumal neben *A* jeweils auch *B* variiert (und umgekehrt). Relativ zu $\{R, S, T, A\}$ beziehungsweise $\{R, S, T, B\}$ jedoch, ist sowohl alleine die Variation von *A* als auch alleine die Variation von *B* verantwortlich für die Variation von *T* (vgl. *c*₃ und *c*₄).

Sei $\mathbf{F}$ die geprüfte Faktormenge, die den Faktor $W \in \mathbf{F}^w$ enthält. Sei weiter σ eine aus Literalen der Faktoren $S_1, \dots, S_k$, $k \geq 1$, bestehende potenzielle Störursache von *W*. Allgemein gilt, dass, wenn es keine zwei Konfigurationen gibt, bei denen relativ zu $\mathbf{F} \cup \{S_1, \dots, S_k\}$ allein die Variation von σ zu einer Variation von *W* führt, σ entweder (1) für keine Variation von *W* in den untersuchten Konfigurationen verantwortlich ist oder (2) nur genau dann eine Variation von *W* erzeugt, wenn ebenfalls tatsächlich für *W* kausal relevante Literale von Faktoren aus $\mathbf{F}$ variieren. Gilt (1), können die potenziellen Störursachen nicht der Grund für das Vorhandensein von Konfigurationen sein, bei denen sich tatsächliche oder vermeintliche Difference-Maker-Qualitäten manifestieren. Letzteres ist jedoch notwendig dafür, dass überhaupt Bikonditionale abgeleitet werden, die einfache Minimale Theorien beschreiben. Gilt (2) und ergeben sich dadurch tatsächlich einfache Minimale Theorien, kann das entsprechende Auftrittsverhalten unter anderem durch die ebenfalls variierenden, kausal relevanten Faktoren aus der geprüften Faktormenge begründet werden.

Insgesamt lässt sich dieses trotz Variationen ebenfalls unproblematische Verhalten der potenziellen Störursachen durch folgende Definition der Homogenität einfangen:

Definition 3.7 (Homogenität)

Sei $\mathcal{K}_{\mathbf{F}}$ eine Konfigurationstabelle über die Faktormenge $\mathbf{F}$ und *W* ein beliebiger Faktor

aus $\mathbf{F}^W$. Sei weiter der aus Literalen $S_1, \dots, S_k$, $k \geq 1$, bestehende Konjunktionsterm σ eine beliebige potenzielle Störursache von W . σ ist in $\mathcal{K}_{\mathbf{F}}$ genau dann homogenisiert, wenn es keine zwei Konfigurationen gibt, bei denen relativ zu $\mathbf{F} \cup \{S_1, \dots, S_k\}$ allein die Variation von σ zu einer Variation von W führt. $\mathcal{K}_{\mathbf{F}}$ ist genau dann homogen, wenn jede potenzielle Störursache jeder potenziellen Wirkung homogenisiert ist.

Die Homogenisierung von potenziellen Störursachen gemäss Definition 3.7 schliesst die Homogenisierung gemäss Definition 3.6 mit ein, weshalb es sich um eine Abschwächung der strikten Homogenität handelt. Die Tragweite dieser Abschwächung wird sichtbar, wenn man sich nochmals das Beispiel aus Abbildung 3.22 vor Augen führt, bei dem T mehrere potenzielle Störursachen besitzt. Mit der strikten Homogenisierung als Bedingung an die potenziellen Störursachen würde man jeweils verlangen, dass, wenn eine potenzielle Störursache von T bei einer Konfiguration gegeben ist, immer dieselbe potenzielle Störursache von T auch bei allen anderen Konfigurationen gegeben ist. Mit der schwächeren Form aus Definition 3.7 kann bei jeder Konfiguration eine andere potenzielle Störursache von T anwesend sein, sofern beispielsweise bei allen Konfigurationen mindestens eine gegeben ist. Abbildung 3.23 etwa zeigt eine Kausalstruktur zusammen mit einem von vielen Instantiierungsverhalten der potenziellen Störursachen X , Y und Z , das relativ zu $\mathbf{A} = \{U, V, W\}$ zu keinem kausalen Fehlschluss führt, obwohl X , Y und Z variieren. Die Konfigurationen über $\mathbf{A} = \{U, V, W\}$, denen die Kausalstruktur links zugrunde liegt, sind homogen gemäss Definition 3.7.

Die Homogenität dient der Approximierung der Erweiterungsklausel in Entdeckungskontexten. So wäre beispielsweise die Konfigurationstabelle 2.14, die in Kapitel 2 zur Begründung der Einführung der Erweiterungsklausel angeführt wurde, nicht homogen (vgl. S. 52): Weil es eine einfache Minimale Theorie von C gibt, müssen dessen potenzielle Störursachen homogenisiert sein, was aufgrund der vorhandenen Variationen von C offensichtlich nicht der Fall ist. Ist die Homogenität gegeben, ist dies jedoch allgemein nicht gleichzusetzen mit der Einhaltung der Erweiterungsklausel, zumal nach wie vor Ambiguitäten möglich sind. Wie sich im Zusammenhang mit dem Kettenproblem gezeigt hat, können Minimale Theorien, die der Erweiterungsklausel ge-



Abbildung 3.23.: Die untersuchten Konfigurationen über $\mathbf{A} = \{U, V, W\}$ sind homogen gemäss Definition 3.7.

nügen, eindeutig kausal interpretiert werden. Im Gegensatz dazu stellen Minimale Theorien, die aus einer homogenen Konfigurationstabelle folgen, aufgrund von Ambiguitäten nicht zwangsläufig Kausalzusammenhänge dar. Beispielsweise ist die mit einer Kausalkette verträgliche Konfigurationstabelle aus Abbildung 3.4 homogen, trotzdem ist nicht eindeutig identifizierbar, ob dieser Konfigurationstabelle nun eine Kausalkette oder eine Common-Cause-Struktur zugrunde liegt (vgl. S. 102 ff.). Die Homogenität ist im Rahmen des kausalen Schliessens insofern als Ersatz für die Erweiterungsklausel zu verstehen, als beide kausale Fehlschlüsse verhindern. Bei der Homogenität werden dazu Massnahmen umgesetzt, die innerhalb der Analyse einer einzelnen Konfigurationstabelle (theoretisch) umsetzbar sind. Aufgrund des induktiven Charakters der Erweiterungsklausel geht mit der Verifizierung der Homogenität jedoch nicht ebenfalls zwangsläufig die Folgerung einer eindeutigen Kausalstruktur einher. Sofern Ambiguitäten aus der Analyse einer Konfigurationstabelle hervorgehen, kann mit der Verifizierung der Homogenität aber sichergestellt werden, dass eines der über die Bestimmung der Minimalen Theorien abgeleiteten Modelle der dahinterliegenden Kausalstruktur oder einem Teil davon entspricht beziehungsweise kein Fehlschluss generiert wird.

Sei $\mathcal{K}_F$ eine Konfigurationstabelle über eine Faktormenge F , die den Annahmen aus Abschnitt 3.1 genügt (vgl. S. 100 ff.). Sei weiter Ω jene Kausalstruktur, die $\mathcal{K}_F$ zugrunde liegt. Der Umstand, dass, wenn eine Konfigurationstabelle $\mathcal{K}_F$ über F homogen ist, kein kausaler Fehlschluss generiert wird, lässt sich folgendermassen nachweisen.

Erstens folgt aus der Homogenitätsannahme, dass (i) jedes Literal W eines Faktors aus F^w , das heisst jedes Literal, das eine potenzielle Wirkung ist, relativ zu F auch eine tatsächliche Wirkung ist. Würde sich eine einfache Minimale Theorie von einem Literal X eines Faktors aus F finden lassen, das gemäss Ω relativ zu F keine Wirkung ist, wäre die angenommene Homogenität von $\mathcal{K}_F$ verletzt. Indem X als potenzielle Wirkung betrachtet würde, müssten alle potenziellen Störursachen von X homogenisiert sein. Wäre X aber tatsächlich keine Wirkung relativ zu F , gäbe es keinen Faktor in F , dessen Literale kausal relevant für X sind. Das wiederum würde bedeuten, dass alle Ursachen von X potenzielle Störursachen wären, die homogenisiert werden müssten. Folglich müsste X entweder konstant an- oder konstant abwesend sein und gemäss der Bedingung (WD1) könnte man entsprechend keine einfache Minimale Theorie von X ableiten (vgl. S. 105).

Zweitens ist (ii) bei allen Konfigurationen aus $\mathcal{K}_F$, bei denen W gegeben ist, auch jeweils mindestens eine tatsächliche, aus Faktoren aus F bestehende Ursache von W gegeben. Dabei ist ϵ_i genau dann eine aus Faktoren aus F bestehende Ursache von W , wenn ϵ_i ein nur aus Literalen von Faktoren aus F bestehender Konjunktionsterm ist, der mit oder ohne latente Ko-Literalen eine komplexe Ursache von W bildet. Weil W eine potenzielle Wirkung ist, kann sie nicht in

allen Konfigurationen entweder konstant an- oder abwesend sein. Es gibt also mindestens eine Konfiguration c_i in $\mathcal{K}_F$, bei der W gegeben ist, und mindestens eine andere Konfiguration c_j , bei der W nicht gegeben ist. Wäre für den Auftritt von W in mindestens einer Konfiguration c_k eine potenzielle Störursache verantwortlich, ohne dass auch eine tatsächliche, aus Faktoren aus F bestehende Ursache gegeben wäre, gäbe es mit c_k und c_j mindestens zwei Konfigurationen, bei denen alleine die Variation einer potenziellen Störursache zu einer Variation von W führen würde. Damit wäre $\mathcal{K}_F$ nicht homogen.

Seien $\epsilon_1, \dots, \epsilon_n$, $n \geq 1$, tatsächliche, aus Faktoren aus F bestehende Ursachen von W , die in einer Konfiguration c_i zusammen mit W auftreten. Es gilt drittens, dass (iii) mindestens eine tatsächliche Ursache ϵ_i , $1 \leq i \leq n$, in $\mathcal{K}_F$ hinreichend für W ist. Wäre keine der Ursachen $\epsilon_1, \dots, \epsilon_n$ hinreichend für die potenzielle Wirkung W , könnte $\mathcal{K}_F$ nicht homogen sein. Denn es gäbe für jede Ursache ϵ_i , $1 \leq i \leq n$, mindestens eine Konfiguration c_j in $\mathcal{K}_F$, bei der ϵ_i gegeben und W nicht gegeben wäre.⁵ Eine solche Variation von W zwischen c_i und c_j könnte nur zustande kommen, wenn in c_j latente Ko-Literale von ϵ_i nicht gegeben sowie keine unberücksichtigten Alternativursachen von W gegeben wären und in c_i

- (1) alle latenten Ko-Literale von ϵ_i gegeben wären oder
- (2) mindestens eine unberücksichtigte Alternativursache von W gegeben wäre oder
- (3) weder (1) noch (2) gelten würde, das heisst auch in c_i latente Ko-Literale von ϵ_i nicht gegeben wären und die Anwesenheit von W in c_i auf andere Ursachen aus $\epsilon_1, \dots, \epsilon_n$ zurückzuführen wäre.

Bei den beiden Fällen (1) und (2) ginge die Variation von W alleine auf die Variation von potenziellen Störursachen zurück, womit $\mathcal{K}_F$ inhomogen wäre. Für mindestens eine der Ursachen $\epsilon_1, \dots, \epsilon_n$ müsste jedoch mindestens einer dieser zwei Fälle gelten, zumal W sonst in c_i nicht anwesend sein könnte. Folglich könnte $\mathcal{K}_F$ nicht homogen sein, wenn (iii) nicht gelten würde.

Aus den obigen zwei Punkten (ii) und (iii) folgt, dass in jeder Konfiguration aus $\mathcal{K}_F$, bei der W gegeben ist, auch mindestens eine aus Faktoren aus F bestehende tatsächliche Ursache von W gegeben ist, die sich in $\mathcal{K}_F$ als hinreichend für W herausstellt. Da es sich um eine tatsächliche Ursache von W handelt, wird sie gemäss dem Verursachungsbegriff auch minimal hinreichend für W sein. All diese tatsächlichen Ursachen bilden disjunktiv verknüpft eine notwendige Be-

⁵Ist in $\mathcal{K}_F$ keine der tatsächlichen, aus Faktoren aus F bestehenden Ursachen von W hinreichend für W , heisst das nicht, dass es in $\mathcal{K}_F$ keine hinreichenden Bedingungen von W geben kann. Durch Variationen potenzieller Störursachen können Konjunktionsterme hinreichend für W sein, die keine tatsächlichen Ursachen sind, sodass W trotzdem eine potenzielle Wirkung ist, für welche sich relativ zu $\mathcal{K}_F$ einfache Minimale Theorien finden lassen.

dingung von W . Befreit man die notwendige Bedingung von allfälligen redundanten Disjunkten, ergibt sich insgesamt mindestens eine aus minimal hinreichenden Bedingungen bestehende minimal notwendige Disjunktion von W , die ausschliesslich aus tatsächlichen Ursachen von W besteht. Analoges gilt für jede weitere potenzielle Wirkung, die gemäss (i) relativ zu $\mathbf{F}$ auch eine tatsächliche Wirkung ist, weshalb insgesamt eine Minimale Theorie aus der homogenen Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ folgt, die der tatsächlichen Kausalstruktur Ω nicht widerspricht.

3.4.2. Die Rolle von Homogenität für das analytische Moment

Sei $\mathcal{K}_{\mathbf{F}}$ eine Konfigurationstabelle über die Faktormenge $\mathbf{F}$, die den Annahmen aus Abschnitt 3.1 genügt (vgl. S. 100 ff.). Ist $\mathcal{K}_{\mathbf{F}}$ zusätzlich homogen, wird über einen korrekt arbeitenden Algorithmus, der aus $\mathcal{K}_{\mathbf{F}}$ alle Minimalen Theorien ableitet, garantiert kein kausaler Fehlschluss folgen.

Weil durch die Homogenität kausale Fehlschlüsse ausgeschlossen sind, können bei deren Voraussetzung weiter auch jene Fehlschlüsse nicht entstehen, die man möglicherweise darauf zurückführen möchte, dass empirisch mögliche Konfigurationen von Faktoren aus einer Faktormenge $\mathbf{F}$ in der untersuchten Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ über $\mathbf{F}$ fehlen. Ein Grund für solche Datenlücken kann etwa sein, dass Fälle gewisser Konfigurationen von Faktoren, die empirisch möglich sind, noch nie eingetroffen sind oder in einer jeweiligen Studie nicht gemessen beziehungsweise beobachtet wurden. Solche Fälle werden trivialerweise nicht erhoben und die entsprechenden Konfigurationen somit auch nicht Platz in der untersuchten Konfigurationstabelle finden. Ein typisches Beispiel ist Szenario 3 aus dem vorherigen Abschnitt. $\mathcal{K}_{3,12}$ enthält angesichts der tatsächlichen kausalen Unabhängigkeit zwischen A , B und C nicht alle empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{O} = \{A, B, C\}$ (vgl. S. 116). Relativ zu $\mathbf{O}$ gibt es nicht nur vier, sondern insgesamt acht dazugehörige Konfigurationen. Wie im vorherigen Abschnitt gezeigt, ist $\mathcal{K}_{3,12}$ aber vor allem auch nicht homogen. Werden die potenziellen Störursachen der potenziellen Wirkungen homogenisiert, löst sich das Risiko eines kausalen Fehlschlusses auf. Wie unter anderem die homogenen Konfigurationstabellen aus 3.17 zeigen, können diese dabei sogar durchaus in dem Sinne Lücken aufweisen, dass empirisch mögliche Konfigurationen nicht darin enthalten sind. Das Fehlen von Konfigurationen führt typischerweise dazu, dass keine Strukturen oder nur Teile der eigentlichen Kausalstruktur aus der untersuchten Konfigurationstabelle abgeleitet werden können. Zu einem kausalen Fehlschluss jedoch führt es nicht, sofern die Konfigurationstabelle homogen ist.

Eine weitere Konsequenz der angenommenen Homogenität betrifft gewisse Kennwerte, die vor allem bei der praktischen Umsetzung des konfiguralen kausalen Modellierens eine we-

sentliche Rolle einnehmen und etwas über die Aussagekraft von abgeleiteten Kausalmodellen aussagen. So gilt für jede der aus einer homogenen Konfigurationstabelle abgeleiteten einfachen Minimalen Theorien, dass das *Konsistenzmass* der Suffizienz der minimal hinreichenden Bedingungen und das Konsistenzmass der Notwendigkeit der minimal notwendigen Bedingungen jeweils 1 ist. Das Konsistenzmass ist eine Kennzahl, die bei QCA für den Umgang mit Fällen eingeführt wurde, die sich bezüglich eines Kausalmodells widersprechen (vgl. bspw. Ragin (2008), S. 45 ff.; Schneider & Wagemann (2012), S. 123 ff.). Es kann etwa vorkommen, dass in 39 von 40 untersuchten Fällen die Suffizienz eines Literals A für ein Literal W gestützt wird und nur ein dieser Suffizienz widersprechender Fall vorliegt, bei dem zwar A gegeben, nicht aber W gegeben ist. Wird eine strenge Interpretation der Suffizienz verwendet, wonach A genau dann als hinreichend für W angesehen wird, wenn immer, wenn A gegeben ist, auch W gegeben ist, müsste man die Suffizienz von A für W zurückweisen. In der Praxis würde man die Suffizienz von A für W aufgrund des einen dieser Suffizienz widersprechenden Falles jedoch kaum bestreiten wollen. Vielmehr würde man dafür argumentieren, dass dieser Fall ein ungewöhnlicher Ausreisser ist. Als Rechtfertigung dafür, die Suffizienz von A für W trotzdem als wahr betrachten zu können, wird das Konsistenzmass herangezogen. Das Konsistenzmass statet Aussagen wie „ A ist hinreichend für W “ mit einem zwischen 0 und 1 liegenden Kennwert aus, der sich aus den untersuchten Konfigurationen und deren Fallzahl berechnen lässt. Je näher der Wert bei 1 liegt, desto mehr wird die entsprechende Aussage von den Daten gestützt. Liegt der Wert bei 1, sind keine Fälle vorhanden, die der Aussage widersprechen. Berechnet wird das Konsistenzmass $K_{Z_1 \cdots Z_m \rightarrow W}$ für die Suffizienz einer Bedingung $Z_1 \cdots Z_m$, $m \geq 1$, von W bei der crisp-set Variante durch das folgende Verhältnis (vgl. Schneider & Wagemann (2012), S. 124 f.):

$$K_{Z_1 \cdots Z_m \rightarrow W} = \frac{i}{j}, \quad (3.11)$$

wobei i die Anzahl der „günstigen“ Fälle ist, bei denen sowohl $Z_1 \cdots Z_m$ als auch W gegeben sind und j die Anzahl Fälle, bei denen $Z_1 \cdots Z_m$ gegeben ist. Bei letzteren handelt es sich also um jene Fälle, bei denen $Z_1 \cdots Z_m$ zusammen mit oder ohne W anwesend ist. Um nun zu entscheiden, ob eine Aussage über die Suffizienz von $Z_1 \cdots Z_m$ für W als wahr betrachtet werden kann, wird in der Praxis vorab ein Schwellenwert bestimmt, der vom berechneten Wert überschritten werden muss. Der Zweck des Konsistenzmasses kann also verglichen werden mit dem des Signifikanzniveaus der statistischen Tests.

Ähnlich zum Konsistenzmass für die Suffizienz einer Bedingung $Z_1 \cdots Z_m$ von W kann mit der Notwendigkeit einer Bedingung $Z_1 \cdots Z_m$ von W verfahren werden. Dabei widersprechen genau jene Fälle der Notwendigkeit einer Bedingung $Z_1 \cdots Z_m$ für W , bei denen zwar W , nicht aber $Z_1 \cdots Z_m$ gegeben ist. Berechnet wird das Konsistenzmass $K_{Z_1 \cdots Z_m \leftarrow W}$ für die Notwendigkeit

einer Bedingung $Z_1 \cdots Z_m$ von W bei der crisp-set Variante durch das folgende Verhältnis (vgl. Schneider & Wagemann (2012), S. 140 f.):

$$K_{Z_1 \cdots Z_m \leftarrow W} = \frac{i}{k}, \quad (3.12)$$

wobei i erneut die Anzahl der „günstigen“ Fälle ist, bei denen sowohl $Z_1 \cdots Z_m$ als auch W gegeben sind und k die Anzahl Fälle, bei denen W mit oder ohne $Z_1 \cdots Z_m$ gegeben ist.

So wie sich das Konsistenzmass für die Suffizienz und Notwendigkeit einzelner Konjunktionsterme $Z_1 \cdots Z_m$ berechnen lässt, lässt es sich auch für Disjunktionen von Konjunktionen berechnen. Möchte man etwa das Konsistenzmass $K_{Z_1+Z_2 \rightarrow W}$ beziehungsweise $K_{Z_1+Z_2 \leftarrow W}$ für die Suffizienz beziehungsweise Notwendigkeit einer Disjunktion $Z_1 + Z_2$ für W berechnen, ergibt sich die Anzahl i der „günstigen“ Fälle aus der Anzahl jener Fälle, bei denen Z_1 oder Z_2 (oder beide) zusammen mit W anwesend sind. Der Wert j ergibt sich aus der Anzahl Fälle, bei denen Z_1 oder Z_2 (oder beide) gegeben sind. Der Wert k bleibt unverändert und entspricht der Anzahl Fälle, bei denen W anwesend ist. Es lässt sich unter anderem leicht feststellen, dass, wenn das jeweilige Konsistenzmass der Suffizienz zweier Konjunktionsterme für W 1 ist, auch das Konsistenzmass der Suffizienz der Disjunktion der beiden Terme 1 ist.

Wird Homogenität angenommen, wird das Konsistenzmass der Suffizienz der minimal hinreichenden Bedingungen und der disjunktiven Verknüpfungen solcher Bedingungen von abgeleiteten einfachen Minimalen Theorien immer 1 sein. Ist etwa $Z_1 \cdots Z_m$ eine in einer einfachen Minimalen Theorie von W enthaltene minimal hinreichende Bedingung, gibt es keinen Fall und damit auch keine Konfiguration, bei der $Z_1 \cdots Z_m$ gegeben und W nicht gegeben ist. Würde es eine solche Konfiguration geben, würde man $Z_1 \cdots Z_m$ nicht als hinreichend für W betrachten, zumal die Homogenität verletzt wäre. Man müsste beispielsweise davon ausgehen, dass im Gegensatz zu den die Suffizienz stützenden Konfigurationen in dieser Konfiguration potenzielle Störursachen von W derart variieren, dass die Anwesenheit von W gehemmt wird. Solche ausschlaggebenden Variationen von potenziellen Störursachen werden jedoch durch die Homogenität als nicht existent angenommen. Analog kann für das jeweils den Wert 1 annehmende Konsistenzmass der Notwendigkeit der minimal notwendigen Bedingungen von abgeleiteten einfachen Minimalen Theorien von W argumentiert werden. Weist eine solche minimal notwendige Bedingung ein Konsistenzmass der Notwendigkeit auf, das kleiner als 1 ist, gibt es eine Konfiguration mit Fallzahl ≥ 1 , bei der zwar W , nicht aber die minimal notwendige Bedingung gegeben ist. Unter diesen Gegebenheiten müsste man also beispielsweise davon ausgehen, dass diese der Notwendigkeit der Bedingung widersprechende Konfiguration das Produkt einer nicht homogenisierten

potenziellen Störursache von W ist. Solche Szenarien verletzen die angenommene Homogenität.⁶

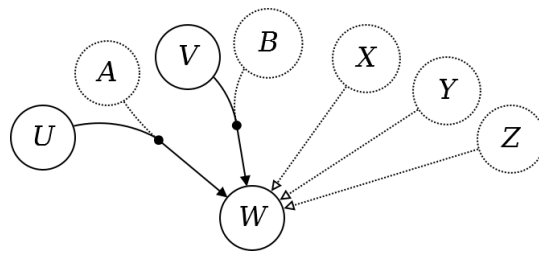
Es bleibt zu erwähnen, dass der Umstand, dass man in der Praxis auch bereit ist, bei einem Konsistenzmass < 1 kausal zu schliessen, zeigt, dass man auch bereit ist, bei verletzter Homogenität kausal zu schliessen. Die Homogenitätsanforderung ist nicht so zu verstehen, dass nur aus homogenen Konfigurationstabellen kausal geschlossen werden darf. Vielmehr liefert die Homogenität einfach eine Korrektheitsgarantie: Ist eine Konfigurationstabelle $\mathcal{K}_F$ homogen, ergibt sich garantiert kein kausaler Fehlschluss. Ist $\mathcal{K}_F$ nicht homogen, folgt dagegen nicht, dass ein Fehlschluss entsteht. Es ist nur nicht mehr garantiert, dass kein Fehlschluss produziert wird. In bestimmten Fällen mag man bereit sein, dieses Fehlschluss-Risiko einzugehen.

Mit der Annahme der Homogenität ergibt sich insgesamt, dass jede aus einer Konfigurationstabelle abgeleitete einfache Minimale Theorie der Definition 2.9 derart im strengen Sinne genügt, wie sie bei der Skizzierung des analytischen Moments und der damit verbundenen Wirkungsdeklaration angewendet wurde. Das heisst, unter diesen Voraussetzungen wird es keinen Fall und damit keine Konfiguration geben, die den abgeleiteten einfachen Minimalen Theorien widerspricht. Folglich wird die Berechnung dieser Kennwerte künftig vernachlässigt.

3.5. Maximale Diversität

Neben der Homogenität als zentraler Forderung beim konfigurationsalen kausalen Modellieren wird oft zusätzlich eine Art empirische Vollständigkeit der Konfigurationstabelle vorausgesetzt, damit sich ein Suchalgorithmus nicht nur darauf prüfen lässt, ob er keine kausalen Fehlschlüsse produziert, sondern auch, ob er alle tatsächlichen kausalen Abhängigkeiten zwischen den Literalen von Faktoren einer Menge F findet. Diese Art der empirischen Vollständigkeit einer Konfigurationstabelle $\mathcal{K}_F$ ist jedoch im Allgemeinen von jener zu unterscheiden, gemäss der alle empirisch möglichen Konfigurationen in $\mathcal{K}_F$ enthalten sind (vgl. Abschnitt 2.8.1, S. 50 f.). Der Grund dafür liegt erstens in der oftmals hohen kausalen Komplexität, die Kausalstrukturen aufweisen. Diese Kausalstrukturen werden zweitens innerhalb einer Analyse relativ zu Faktormengen untersucht, die meist nicht alle kausal verknüpften Faktoren aus der Struktur beinhalten. Wird beispielsweise für die drei Faktoren U , V und W eine Kausalanalyse durchgeführt, wobei

⁶Ein betreffend seine Interpretation weiterer wesentlicher QCA-Kennwert stellt das sogenannte *Abdeckungs-mass* dar, das etwas darüber aussagt, wie präsent hinreichende oder notwendige Bedingungen eines als Wirkung betrachteten Literals in den empirischen Daten sind. Eine detaillierte Beschreibung und Begründung des Konsistenz- und Abdeckungs-masses sowie des Zusammenhangs zwischen diesen beiden findet man u. a. bei Ragin (2008), Kapitel 3 oder Schneider & Wagemann (2012), Kapitel 5.



	A	B	C	U	V	W
c_1	1	1	1	1	1	1
c_2	0	0	0	1	1	0
c_3	1	1	1	1	0	1
c_4	0	1	0	1	0	0
c_5	1	1	1	0	1	1
c_6	1	0	0	0	1	0
c_7	1	1	1	0	0	1
c_8	1	1	0	0	0	0

Abbildung 3.24.: Datengenerierende Kausalstruktur Ω und Auswahl von dazugehörigen Konfigurationen.

U und V zwei Alternativursachen von W darstellen, muss davon ausgegangen werden, dass einerseits weitere Alternativursachen von W existieren und andererseits sowohl U als auch V Ko-Literale besitzen, die zusammen mit U oder V eine komplexe Ursache von W bilden. Abbildung 3.24 zeigt eine Auswahl dazugehöriger Kombinationen von an- und abwesenden Faktoren, die vor diesem Hintergrund realisierbar sind, wenn U und V jeweils um ein Ko-Literal ergänzt werden sowie eine weitere Alternativursache von W existiert.

Würde man verlangen, dass alle empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{Q} = \{U, V, W\}$ vorliegen, wären in der entsprechenden Konfigurationstabelle relativ zu $\mathbf{Q}$ alle $2^3 = 8$ logisch möglichen Konfigurationen der Faktoren U , V und W enthalten. Aus einer solchen der trivialen Homogenität genügenden Konfigurationstabelle lässt sich kein Kausalmodell folgern. Angesichts der typischerweise hohen kausalen Komplexität von Kausalstrukturen wäre ein derartiges Szenario die Regel, womit es beim kausalen Schliessen im Allgemeinen wenig Sinn ergibt, relativ zu einer Faktormenge zu verlangen, dass alle empirisch möglichen Konfigurationen der geprüften Faktoren in der untersuchten Konfigurationstabelle enthalten sind. Dabei muss jedoch zwischen dem begrifflichen Unterfangen der Definition der Kausalbeziehung aus Kapitel 2, die genau dies verlangt, und dem epistemologischen Unterfangen des kausalen Schliessens aus dem vorliegenden Kapitel unterschieden werden. Zur Bildung des Verursachungsbegriffs wurde eine Vollständigkeitsklausel in die Definition integriert, mit deren Hilfe sich ein Verursachungsbegriff formulieren lässt, der nicht auf das die Reduktivität der Definition gefährdende kausale Feld angewiesen ist (vgl. S. 53). Beim kausalen Schliessen wird diese Vollständigkeitsklausel im Allgemeinen nicht mehr vorausgesetzt und stattdessen die Homogenität angenommen, die sich ihrerseits basierend auf der Kausaltheorie definieren lässt.

Ein Literal X ist nur dann kausal relevant für ein anderes Literal W , wenn es ein Difference-Making-Paar von X bezüglich W gibt. Dieses Paar zeichnet sich dadurch aus, dass relativ zu

$\mathcal{K}_{3.25a}$	U	V	W
c_1	1	1	1
c_2	1	0	1
c_3	0	1	1
c_4	0	0	0
(a)			

$\mathcal{K}_{3.25b}$	A	B	C
c_1	1	1	1
c_2	1	0	1
c_3	0	1	1
c_4	0	0	0
c_5	1	1	0
c_6	1	0	0
c_7	0	1	0
c_8	0	0	1
(b)			

Tabelle 3.25.: Beispiele zweier maximal diverser Konfigurationstabellen über $\mathbf{Q} = \{U, V, W\}$ beziehungsweise $\mathbf{O} = \{A, B, C\}$.

einer Minimalen Theorie von W , die X enthält, alleine die Variation von X für die Variation von W verantwortlich gemacht werden kann. Damit also alle Ursachen von W aufgedeckt werden können, braucht es genügend Konfigurationen, die sich in paarweisen Vergleichen derart minimal voneinander unterscheiden, dass es von jedem der relevanten Literale Difference-Making-Paare bezüglich W gibt – kurz: Die Konfigurationstabelle muss eine genügend hohe Diversität aufweisen. Idealerweise ist diese Diversität derart hoch, dass die k Faktoren, die relativ zur untersuchten Faktormenge $\mathbf{F}$ keine potenziellen Wirkungen sind und deren Verhalten somit nicht vom Verhalten anderer Faktoren aus $\mathbf{F}$ abhängig ist, in allen 2^k logisch möglichen Kombinationen vorkommen. Trifft dies für eine Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ zu, wird $\mathcal{K}_{\mathbf{F}}$ als *maximal divers* bezeichnet. Anstelle der Forderung, dass alle empirisch möglichen Konfigurationen der Faktoren aus $\mathbf{F}$ in $\mathcal{K}_{\mathbf{F}}$ enthalten sind, wird also maximale Diversität von $\mathcal{K}_{\mathbf{F}}$ gefordert:

Definition 3.8 (Maximale Diversität)

Sei $\mathcal{K}_{\mathbf{F}}$ eine Konfigurationstabelle über eine Faktormenge $\mathbf{F}$. $\mathcal{K}_{\mathbf{F}}$ ist genau dann maximal divers, wenn relativ zu den k Faktoren aus $\mathbf{F} \setminus \mathbf{F}^w$, $k \geq 1$, nur für jede der 2^k logisch möglichen Kombinationen von an- und abwesenden Faktoren aus $\mathbf{F} \setminus \mathbf{F}^w$ eine Konfiguration in $\mathcal{K}_{\mathbf{F}}$ enthalten ist.

Relativ zu $\mathbf{Q}$ wäre beispielsweise die Konfigurationstabelle $\mathcal{K}_{3.25a}$ aus Tabelle 3.25 maximal divers. Relativ zu $\mathbf{O} = \{A, B, C\}$ aus dem dritten Szenario wäre angesichts der kausalen Unabhängigkeit zwischen A , B und C die Konfigurationstabelle $\mathcal{K}_{3.25b}$ maximal divers. Nicht maximal divers wären bezüglich des dritten Szenarios beispielsweise die drei Konfigurationstabellen aus Tabelle 3.18.

Wie die Homogenität ist auch die maximale Diversität von $\mathbf{F}^w$ abhängig. Das bedeutet konkret, dass es relativ zur gleichen Faktormenge mehrere maximal diverse Konfigurationstabellen geben kann. So wäre beispielsweise auch jene Konfigurationstabelle über $\mathbf{Q}$ maximal divers, die alle 2^3 logisch möglichen Kombinationen der drei Faktoren und damit alle empirisch möglichen Konfigurationen enthält. In diesem Fall wäre $\mathbf{F}^w = \emptyset$. Der entscheidende Unterschied zwischen maximaler Diversität und der Forderung, dass alle empirisch möglichen Konfigurationen vorhanden sind, zeigt sich, wenn man sie mit der Homogenität kombiniert. Angesichts der hohen kausalen Komplexität von Kausalstrukturen kann eine dazugehörige Konfigurationstabelle über eine Faktormenge $\mathbf{F}$ meistens nur genau dann alle empirisch möglichen Konfigurationen enthalten und gleichzeitig homogen sein, wenn es sich bei letzterer um die triviale Homogenität handelt, womit sich keine Kausalmodelle ableiten lassen. Im Gegensatz dazu ist beispielsweise $\mathcal{K}_{3,25a}$ vor dem Hintergrund jener Kausalzusammenhänge, gemäss denen U und V zwei Alternativursachen von W sind, maximal divers und homogen.⁷

Es gilt abschliessend nochmals zu bemerken, dass die maximale Diversität weder eine notwendige Ergänzung der Homogenität zur Verhinderung von kausalen Fehlschlüssen darstellt, noch – wie etwa obige Konfigurationstabelle $\mathcal{K}_{3,25b}$ vermuten liesse – eine mögliche Alternative zur Homogenität beschreibt. Einerseits geht aus dem vorherigen Abschnitt hervor, dass die Bedingung der Homogenität alleine ausreicht, um kausale Fehlschlüsse zu verhindern. Andererseits löst sich durch die maximale Diversität zwar der kausale Fehlschluss des dritten Szenarios auf, nicht aber jene der ersten beiden Szenarien. Sowohl $\mathcal{K}_{3,6}$ als auch $\mathcal{K}_{3,9}$ sind gemäss Definition 3.8 maximal divers und trotzdem wird ein kausaler Fehlschluss provoziert. Die maximale Diversität wird hier vor allem deshalb eingeführt, weil sie erlaubt, relativ zu einer Faktormenge $\mathbf{F}$ auf alle kausalen Zusammenhänge zu schliessen, die es relativ zu $\mathbf{F}$ gibt, ohne dass die entsprechende Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ komplett sein muss (vgl. Def. 2.14, S. 59). Des Weiteren ermöglicht die maximale Diversität, dass im Folgekapitel das mathematische Fundament anschaulich eingeführt werden kann.

⁷Es kann auch mehrere homogene und gleichzeitig maximal diverse Konfigurationstabellen über $\mathbf{F}$ mit identischer Menge $\mathbf{F}^w$ geben. Bezüglich $\mathbf{Q} = \{U, V, W\}$ wäre beispielsweise auch jene Konfigurationstabelle maximal divers und homogen, bei der W nur in c_1 und c_2 den Wert 1 annimmt und in c_3 und c_4 den Wert 0. In diesem Fall wären die Ko-Literale von V konstant abwesend. Insgesamt ist aber vor allem jener homogene Hintergrund interessant, gemäss dem alle latenten Ko-Literale jeweils gegeben und alle anderen potenziellen Störursachen jeweils nicht gegeben sind. Für diesen Hintergrund kann sich jedes tatsächlich kausal relevante Literal eines Faktors aus der geprüften Faktormenge in den Daten als Difference-Maker manifestieren. Es sind dementsprechend auch diese maximal diversen und homogen Konfigurationstabellen, wie $\mathcal{K}_{3,25a}$ eine ist, die in der vorliegenden Arbeit vor allem betrachtet werden.

3.6. In Kürze

Das vorliegende Kapitel beantwortet die Frage, welche Informationen konfigurationale Methoden zur Aufdeckung von kausalen Abhängigkeiten zwischen Literalen von Faktoren einer geprüften Menge $\mathbf{F}$ nutzen und wie diese aussehen. Dabei handelt es sich um Konfigurationen der Faktoren aus $\mathbf{F}$, die Kombinationen von an- und abwesenden Faktoren aus $\mathbf{F}$ wiedergeben, die empirisch möglich sind. In der Praxis werden diese durch die Betrachtung von Einzelfällen ermittelt, in denen Instanzen der verschiedenen Faktoren aus $\mathbf{F}$ koinzidieren. Inwiefern ein Ereignis (oder Zustand) dabei als Instanz eines Faktors bewertet werden kann, wird durch eine transparente Kalibrierung festgelegt. Die Auflistung aller Konfigurationen, von denen es (kalibrierte) Fälle gibt, ist eine Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ über $\mathbf{F}$. Wird auf die Werte 0 und 1 kalibriert, handelt es sich um eine crisp-set Konfigurationstabelle.

Konfigurationstabellen bilden den Ausgangspunkt des analytischen Moments und entsprechen jenen Informationen, die einem Algorithmus zur Aufdeckung von kausalen Abhängigkeiten als Input dienen. Damit eine solche Analyse erfolgreich sein kann, ist eine gewisse Datenqualität gefordert. Auch wenn unter anderem die Güte der Daten in angemessener Form vorliegt, sodass Mess- oder Kalibrierungsfehler ausgeschlossen werden können, können kausale Fehlschlüsse provoziert werden. Dabei liegt ein kausaler Fehlschluss genau dann vor, wenn jedes aus der untersuchten Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ abgeleitete Kausalmodell ω_i , $i \in \mathbb{N}$, der Kausalstruktur Ω , die $\mathcal{K}_{\mathbf{F}}$ tatsächlich zugrunde liegt, widerspricht.

Entstehen kann ein solcher kausaler Fehlschluss durch potenzielle Störursachen. Potenzielle Störursachen sind unberücksichtigte Ursachen von potenziellen Wirkungen, die unbemerkt zwischen Fällen verschiedener Konfigurationen variieren können. Je nachdem wie diese potenziellen Störursachen zwischen Konfigurationen variieren, produzieren sie vermeintliche Regelmäßigkeiten zwischen Literalen von Faktoren aus $\mathbf{F}$, die in die gefolgerten Kausalmodelle integriert werden und der tatsächlichen Kausalstruktur Ω widersprechen. Um solche kausalen Fehlschlüsse zu unterbinden, kann die Homogenität der Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ gefordert werden, womit sichergestellt wird, dass das Instantiierungsverhalten von potenziellen Störursachen nicht zu Fehlschluss provozierenden Variationen der potenziellen Wirkungen führt.

In der vorliegenden Arbeit steht das analytische Moment im Vordergrund. Das dabei zum Einsatz kommende Verfahren wird unter idealisierten Bedingungen entwickelt, damit Fehlschlüsse direkt auf fehlerhafte Verfahrensschritte zurückgeführt werden können. Konkret bedeutet dies, dass, sofern nicht explizit anders erwähnt, jede Konfigurationstabelle im weiteren Verlauf dieser Arbeit (a) eine crisp-set Konfigurationstabelle ist, (b) eine angemessene Güte aufweist, (c) für jede Konfiguration dieselbe Fallzahl hat sowie (c) homogen und (d) maximal divers ist.

4. Konfigurationsalgebra

Die Konfigurationstabellen bilden die Grundlage und den Ausgangspunkt des analytischen Moments zur Aufdeckung von unbekannten kausalen Abhängigkeiten mittels konfiguraler Methoden. Ein geeignetes Aufdeckungsverfahren muss in der Lage sein, basierend auf einer solchen Konfigurationstabelle $\mathcal{K}_F$ alle ableitbaren Minimalen Theorien auch tatsächlich abzuleiten, woraus sich schliesslich die mit $\mathcal{K}_F$ verträglichen Kausalmodelle ablesen lassen. Das mathematische Fundament zur Bewältigung dieser Aufgabe für crisp-set Konfigurationstabellen liefert dabei die Boolesche Algebra (vgl. u. a. Whitesitt (1964); Denis-Papin et al. (1974); Gumm & Poguntke (1981)). Im vorliegenden Kapitel wird dieses für das konfigurationale kausale Modellieren genutzte Fundament eingeführt und dessen Zusammenhang mit der Kausaltheorie sowie der empirischen Grundlage aus den Kapiteln 2 und 3 dargestellt.

4.1. Boolesche Algebren

Konfigurationale Methoden nutzen zur Aufdeckung von Kausalmodellen Boolesche Algebren, die auf die Arbeiten des Mathematikers und Logikers George Boole (1815–1864) zurückgehen. Boole veröffentlichte in seiner Schrift *The Mathematical Analysis of Logic* von 1847 erstmals einen Kalkül, der auf einem Verbund der Logik und Algebra fusst und die moderne mathematische Logik begründete (vgl. Boole (1847)). Eine von mehreren möglichen Definitionen einer Booleschen Algebra ist durch das folgende Axiomensystem gegeben¹:

Eine Menge B , ausgestattet mit zwei binären Operationen $\sqcap$ und $\sqcup$, einer einstelligen Operation $\neg$ und zwei ausgezeichneten (neutralen) Elementen 0 und 1 aus B , heisst Boolesche Algebra $B = (B, \sqcap, \sqcup, \neg, 0, 1)$, wenn folgende Gesetze gelten:

Assoziativgesetze:	$(a \sqcap b) \sqcap c = a \sqcap (b \sqcap c)$	$(a \sqcup b) \sqcup c = a \sqcup (b \sqcup c)$
Kommutativgesetze:	$a \sqcap b = b \sqcap a$	$a \sqcup b = b \sqcup a$
Absorptionsgesetze:	$a \sqcap (a \sqcup b) = a$	$a \sqcup (a \sqcap b) = a$
Distributivgesetze:	$(a \sqcup b) \sqcap c = (a \sqcap c) \sqcup (b \sqcap c)$	$(a \sqcap b) \sqcup c = (a \sqcup c) \sqcap (b \sqcup c)$

¹Die vorliegende Definition wurde zitiert aus Bronstein et al. (2001), S. 353.

Neutrale Elemente:	$a \sqcap 1 = a$	$a \sqcup 0 = a$
	$a \sqcap 0 = 0$	$a \sqcup 1 = 1$
Komplement:	$a \sqcap \neg a = 0$	$a \sqcup \neg a = 1$

Für die beiden Gesetze derselben Zeile gilt, dass sich das eine jeweils aus dem anderen Gesetz ergibt, indem man $\sqcap$ durch $\sqcup$, $\sqcup$ durch $\sqcap$, 0 durch 1 sowie 1 durch 0 ersetzt. Dabei handelt es sich um die sogenannte *Dualität*. Man sagt, dass die beiden Gesetze derselben Zeile dual zueinander sind.

Die Boolesche Algebra ist eine algebraische Struktur, die der Aussagenlogik und der Mengenalgebra übergeordnet ist. Letztere sind (mathematische) Interpretationen der Booleschen Algebra, die ihrerseits eine uninterpretierte mathematische Struktur darstellt. Bei einer Interpretation der Booleschen Algebra werden die obigen Operatoren und ausgezeichneten Elemente 0 und 1 auf bestimmte Weise gedeutet. So werden etwa in der Mengenalgebra die Operatoren $\sqcap$, $\sqcup$ sowie $\neg$ in dieser Reihenfolge als Durchschnitt, Vereinigung sowie Komplement und das Element 0 beziehungsweise 1 als leere beziehungsweise universelle Menge interpretiert. Entsprechende Schreibweisen verdeutlichen diese Interpretationen und beispielsweise anstelle von 0 wird die leere Menge typischerweise notiert als $\emptyset$. Die Boolesche Algebra ist auch für die technische Informatik von grosser Bedeutung. Konkret kommt in dieser Disziplin vor allem eine Interpretation der Booleschen Algebra zum Einsatz: die Schaltalgebra. Die Schaltalgebra ist eine Anwendung der Booleschen Algebra, mit deren Hilfe unter anderem Schaltkreise realisiert und minimiert werden können (vgl. bspw. Fabricius (1992); Nelson et al. (1995); Kumar (2014)). Die bei der Schaltalgebra verwendeten Methoden werden auch hier eine wichtige Rolle spielen. Damit man sich jedoch solchen Instrumenten bedienen kann, muss in einem ersten Schritt nachgewiesen werden, dass eine Anwendung dieser Methoden überhaupt gerechtfertigt ist. Dazu wird im Folgenden zuerst die mathematische Struktur eingeführt, die zur Aufdeckung von Kausalzusammenhängen verwendet wird und anschliessend demonstriert, dass diese Struktur den obigen Gesetzen genügt und damit eine Boolesche Algebra ist.

4.1.1. Die Konfigurationsalgebra als mathematisches Modell einer zweielementigen Booleschen Algebra

Die Beschreibung komplexer Kausalstrukturen aus Kapitel 2 hat unter anderem gezeigt, dass anwesende oder abwesende Faktoren kausal relevant für Wirkungen sein können und verschiedene Literale zusammen komplexe Ursachen bilden oder Teil verschiedener Alternativursachen sein können. Diese drei Verknüpfungsformen zusammen mit der kausalen Relevanz an sich, die sich

A	B	$A \cdot B$	A	B	$A + B$	A	$\bar{A}$
1	1	1	1	1	1	1	0
1	0	0	1	0	1	0	1
0	1	0	0	1	1		
0	0	0	0	0	0		
(a)			(b)			(c)	

Tabelle 4.1.: Die drei fundamentalen Verknüpfungen (a) der Konjunktion, (b) der Disjunktion und (c) des Komplements.

weiter in eine direkte und indirekte kausale Relevanz unterteilen lässt, reichen aus, um jede beliebige komplexe Kausalstruktur darzustellen. Analoges zeigt sich auf der empirischen Ebene bei den für die hier behandelten Methoden fundamentalen Konfigurationen und den daraus ableitbaren Minimalen Theorien: Ein Faktor kann instantiiert oder nicht instantiiert und verschiedene Faktoren können gemeinsam oder auch nur der eine oder andere instantiiert sein. Je nach Instantiierungsverhalten kann damit eine Instanz einer Wirkung einhergehen oder auch nicht. Daraus motiviert werden die drei Verknüpfungen (a) $+$, (b) $\cdot$ und (c) $\bar{}$ zwischen den zwei Elementen 0 und 1 aus Tabelle 4.1 definiert, die der Reihe nach als (a) Konjunktion, (b) Disjunktion und (c) Komplement bezeichnet werden. Zur besseren Darstellung werden dabei Grossbuchstaben als Platzhalter für 0 und 1 verwendet.

Obwohl die drei Verknüpfungen aus Tabelle 4.1 vorerst aus rein mathematischer Sicht betrachtet werden und demnach noch uninterpretiert bleiben, dürfte es mit Blick auf die vorhergehenden Kapitel offensichtlich sein, wie diese künftig interpretiert werden sollen. Wird 1 gedeutet als „tritt auf“, „ist gegeben“ oder „ist instantiiert“ und 0 interpretiert als „bleibt aus“, „ist nicht gegeben“ oder „ist nicht instantiiert“, kann die konjunktive Verknüpfung zwischen den zwei 0 oder 1 annehmenden Platzhalter A und B interpretiert werden als gemeinsames Auftreten zweier Faktoren A und B . Eine solche Koinzidenz ist genau dann gegeben, wenn sowohl A als auch B gegeben sind. Ist mindestens einer dieser beiden Faktoren abwesend, ist auch keine Koinzidenz von A und B gegeben. Diese Verknüpfung ist vor allem hinsichtlich der Darstellung von komplexen Ursachen wesentlich. Die disjunktive Verknüpfung von A und B kann interpretiert werden als alternatives Auftreten zweier Faktoren A und B . Die Disjunktion zwischen A und B ist genau dann gegeben, wenn mindestens einer der beiden Faktoren instantiiert ist. Diese Verknüpfung zielt vor allem darauf ab, die verschiedenen Alternativursachen einer Wirkung zum Ausdruck zu bringen. Das Komplement von A kann schliesslich interpretiert werden als das entgegengesetzte Auftretensverhalten des Faktors A . Das Komplement des Faktors A ist genau dann gegeben, wenn A nicht gegeben ist, und nicht gegeben, wenn A gegeben ist. Die Einführung des Komplements ist vor allem daraus motiviert, zwischen der kausalen Relevanz eines positiven und der

kausalen Relevanz eines negativen Literals unterscheiden zu können und diese Unterscheidung entsprechend zum Ausdruck zu bringen.

Die drei oben definierten Verknüpfungen werden abschliessend zusammen mit den zwei Elementen 0 und 1 zur sogenannten *Konfigurationsalgebra* zusammengefasst:

Definition 4.1 (*Konfigurationsalgebra*)

Die Struktur $(0, 1, +, \cdot, \neg)$, bestehend aus der Menge $\mathcal{B} = \{0, 1\}$, der konjunktiven Verknüpfung, der disjunktiven Verknüpfung und dem Komplement heisst Konfigurationsalgebra.²

Die Konfigurationsalgebra ist eine Boolesche Algebra, für die die eingangs des Kapitels beschriebenen Booleschen Gesetze gelten. Dies lässt sich bereits daran erkennen, dass sich die Konfigurationsalgebra lediglich im vorgesehenen Einsatzgebiet sowie der Wahl der typischerweise verwendeten Symbole von der zweiwertigen Aussagenlogik unterscheidet, die ihrerseits eine Boolesche Algebra bildet. Ein ausführlicher Nachweis, dass die Konfigurationsalgebra tatsächlich auch eine Boolesche Algebra ist, kann entsprechend analog geführt werden (vgl. bspw. Whitesitt (1964), S. 52 ff.). Konkret kann anhand von Wahrheitstabellen gezeigt werden, dass die Gesetze der Booleschen Algebra für die oben definierte Konfigurationsalgebra gelten. Dazu werden die Variablen a , b und c ersetzt durch 0 oder 1, $\sqcap$ ersetzt durch die oben definierte Verknüpfung $\cdot$, $\sqcup$ ersetzt durch $+$ und $\neg$ ersetzt durch $\neg$. Anschliessend werden die Gesetze für jede mögliche Belegung ausgewertet und so gezeigt, dass die Ausdrücke auf der linken Seite vom Gleichheitszeichen eines Gesetzes jeweils dieselben Werte annehmen wie jene auf der rechten Seite. Ist dem so, sagt man, dass die Ausdrücke äquivalent sind. Welche Werte die Verknüpfungen bei einer bestimmten Belegung annehmen lässt sich aus der Tabelle 4.1 ablesen. Bei der Reihenfolge der Auswertung hält man sich an die Klammersetzung und arbeitet diese von innen nach aussen ab. Eine Auswertung über alle möglichen Belegungen wird in Tabelle 4.2 beispielhaft für das Distributivgesetz durchgeführt, wobei die resultierenden Werte für den gesamten Ausdruck in Fettschrift dargestellt sind.³ Eine solche Tabelle wird als Wertetabelle bezeichnet. Wie bei den Definitionen der obigen drei Verknüpfungen bereits verwendet, werden dabei für die Platzhalter von 0 und 1 kursive Grossbuchstaben gewählt. Die Platzhalter selbst werden vorläufig als *Variablen* bezeichnet und nicht etwa, wie sich hinsichtlich des vorgesehenen Anwendungsgebietes anbieten würde, als Faktoren. Diese bewusste Unterscheidung zwischen Variablen und Faktoren soll verdeutlichen, dass die im vorliegenden Abschnitt beschriebenen Begriffe aus rein mathematischer Sicht betrachtet werden und noch keine Interpretation in die

²Zur Notation: Für die Menge von Werten werden in der Folge kalligraphische Grossbuchstaben in Fettschrift verwendet, damit keine Verwechslungen mit Mengen von Faktoren entstehen können. Der kalligraphische Grossbuchstabe $\mathcal{K}$ wird nicht benutzt und bleibt als Bezeichner für Konfigurationstabellen reserviert.

³Die entsprechenden Wertetabellen für den Nachweis der Gültigkeit der restlichen Booleschen Gesetze von Seite 140 findet man in Abschnitt A.2 des Anhangs (vgl. S. 247).

A	B	C	$(A + B)$	$\cdot$	C	$(A \cdot C)$	$+$	$(B \cdot C)$
1	1	1	1	1		1	1	1
1	1	0	1	0		0	0	0
1	0	1	1	1		1	1	0
1	0	0	1	0		0	0	0
0	1	1	1	1		0	1	1
0	1	0	1	0		0	0	0
0	0	1	0	0		0	0	0
0	0	0	0	0		0	0	0

Tabelle 4.2.: Nachweis der Gültigkeit des Distributivgesetzes für die Konfigurationsalgebra mittels Auswertung über alle Belegungen.

Anwendung des konfiguralen kausalen Modellierens vorgenommen wird. So sind A , B , C etc. vorerst bloss mathematische Objekte, genauer gesagt (Boolesche) Variablen, die Werte annehmen. Erst wenn die Variablen im Rahmen des konfiguralen kausalen Modellierens eingesetzt werden, werden sie als Faktoren bezeichnet, die ihrerseits natürliche Eigenschaften repräsentieren.⁴ Die Interpretation in die Anwendung wird in Abschnitt 4.2 vorgenommen.

4.1.2. Instanzfunktion

Die Konfigurationsalgebra wird genutzt, um das Instantiierungsverhalten von Faktoren mathematisch zu beschreiben und zu analysieren. Eine fundamentale Rolle spielt dabei die als *Instanzfunktion* bezeichnete Abbildung, die folgendermassen definiert ist:

Definition 4.2 (Instanzfunktion)

Eine Instanzfunktion von n Variablen ist eine Funktion $f : \mathcal{D} \rightarrow \mathcal{B}$ mit $\mathcal{D} \subseteq \mathcal{B}^n$, wobei $\mathcal{B} = \{0, 1\}$ und $\mathcal{B}^n$ das n -fache kartesische Produkt von $\mathcal{B}$ mit sich selbst ist.

Bezüglich der Notation werden für Instanzfunktionen kursive Kleinbuchstaben $f, g, \dots$ verwendet, damit keine Verwechslungen zwischen Variablen und Instanzfunktionen entstehen können. Für Variablen werden teils kursive Grossbuchstaben $A, B, C, \dots$ oder mit Indizes versehene kursive Grossbuchstaben, wie etwa $X_1, X_2, \dots, X_i, \dots$, verwendet. Die jeweilige Wahl hängt dabei vor allem von der Lesbarkeit ab. Im Gegensatz zu Mengen von Variablen, die mittels fetten Grossbuchstaben $\mathbf{A}, \mathbf{B}, \dots$ notiert werden, werden Definitions- und Wertebereiche einer beliebigen Instanzfunktion f , die Mengen von Werten darstellen, durch kalligraphische Grossbuchstaben in

⁴Vgl. S. 15 zum Begriff des Faktors.

A	B	$f_1(A, B)$	A	B	$f_2(A, B)$	A	B	$f_3(A, B)$	A	B	$f_4(A, B)$
1	1	1	1	1	1	1	1	1	1	1	1
1	0	0	1	0	1	1	0	0	1	0	0
0	1	0	0	1	1	0	1	1	0	1	0
0	0	0	0	0	0	0	0	1	0	0	1
(a)			(b)			(c)			(d)		

Tabelle 4.3.: Vier Beispiele von Instanzfunktionen, die mittels Wertetabellen dargestellt werden.

Fettschrift $\mathcal{D}, \mathcal{B}, \dots$ dargestellt. Auch diese Schreibweise dient der Vermeidung von möglichen Verwechslungen.

Eine Instanzfunktion ist eine *Boolesche Funktion* oder eine *partiell definierte Boolesche Funktion*, wobei die Wahl der Bezeichnung durch das vorgesehene Anwendungsgebiet des konfigurationsalen kausalen Modellierens motiviert ist und so auf eine gewisse Interpretationsweise hindeuten soll. Eine Boolesche Funktion bildet jedes aus Nullen und Einsen bestehende n -Tupel $(X_1, \dots, X_n) \in \mathcal{B}^n$, was einer geordneten Folge von Nullen und Einsen entspricht, auf einen Wert aus $\mathcal{B} = \{0, 1\}$ ab (vgl. u. a. Crama & Hammer (2011), S. 4 f.; Whitesitt & Stumpf (1973), S. 19 f.). Bei einer partiell definierten Booleschen Funktion wird eine solche Zuordnung nur für eine Teilmenge aller aus Nullen und Einsen bestehenden n -Tupeln aus $\mathcal{B}^n$ definiert (vgl. u. a. Crama & Hammer (2011), S. 511.). Instanzfunktionen lassen sich eindeutig mittels Wertetabellen darstellen, die die n -Tupel zusammen mit den entsprechenden Funktionswerten auflisten. Tabelle 4.3 zeigt vier Beispiele von Instanzfunktionen, die durch Wertetabellen dargestellt sind.

Die Wertetabellen der beiden Instanzfunktionen f_1 und f_2 sind identisch mit jenen der Verknüpfungen $+$ und $\cdot$. Sie werden entsprechend auch als Konjunktion und Disjunktion bezeichnet. Es soll dabei jedoch bemerkt werden, dass f_1 und f_2 sich insofern von den beiden Verknüpfungen $+$ und $\cdot$ aus Tabelle 4.1 unterscheiden, als f_1 und f_2 innerhalb der Konfigurationsalgebra definiert sind und die Konfigurationsalgebra selbst unter anderem durch $+$ und $\cdot$ definiert wird. Das heisst, im Unterschied zur konjunktiven und disjunktiven Verknüpfung aus Tabelle 4.1 werden bei den beiden als Konjunktion und Disjunktion bezeichneten Instanzfunktionen die Booleschen Gesetze bereits als gegeben vorausgesetzt. Die Instanzfunktion f_3 wird auch als Implikation oder Subjunktion und f_4 auch als Äquivalenz, Bijunktion oder XNOR-Funktion bezeichnet (vgl. u. a. Hoffmann (2016), S. 99; Böhme (1981), S. 149; Morgenstern (1997), S. 11).

Alternativ zur Wertetabelle können Instanzfunktionen durch sogenannte *Boolesche Ausdrücke* dargestellt werden (vgl. u. a. Crama & Hammer (2011), S. 10 f.; Gumm & Poguntke (1981), S. 24 f.). Diese Darstellungsform bietet sich vor allem bei höherstelligen Instanzfunktionen an,

wobei die Stelligkeit einer Instanzfunktion der Anzahl Argumente der Funktion entspricht. Während sich etwa die zweistelligen Instanzfunktionen aus Tabelle 4.3 mit ihren $2^2 = 4$ Belegungen übersichtlich durch eine Wertetabelle darstellen lassen, ist diese Darstellung bei Funktionen mit hoher Stelligkeit kaum mehr handhabbar. Für die Konfigurationsalgebra als zweielementige Boolesche Algebra lassen sich Booleschen Ausdrücke induktiv mittels der Verknüpfungen $+$, $\cdot$ und $\bar{}$ wie folgt definieren:

Definition 4.3 (Boolescher Ausdruck)

Sei $\mathcal{B} = \{0, 1\}$ und $X_1, \dots, X_n$ Variablen, die die Werte aus $\mathcal{B}$ annehmen. Ein Boolescher Ausdruck wird folgendermassen definiert:

- (a) Die Konstanten 0 und 1 sowie jede Variable X_i , $1 \leq i \leq n$, sind Boolesche Ausdrücke.
- (b) Sind μ und ν Boolesche Ausdrücke, so auch $\bar{\mu}$, $(\mu + \nu)$ und $(\mu \cdot \nu)$.
- (c) Boolesche Ausdrücke sind genau jene, die sich mit (a) und (b) in endlich vielen Schritten konstruieren lassen.

Eine nicht-negierte Variable X oder eine negierte Variable $\bar{X}$ in einem Booleschen Ausdruck werden als *Literale* bezeichnet, wobei X das positive und $\bar{X}$ das negative Literal darstellt. Weiter ist betreffend die Bildung von Booleschen Ausdrücken wichtig zu bemerken, dass nicht nur einzelne Literale disjunktiv oder konjunktiv miteinander verknüpft werden können beziehungsweise das Komplement nur auf ein Literal $\bar{X}$ angewendet werden kann. Die Verknüpfungen können auch auf mehrere Literale beinhaltende Ausdrücke angewendet werden. Auch $\overline{(X + Y)}$ ist beispielsweise ein Boolescher Ausdruck, der, wie von der Definition des Komplements aus Tabelle 4.1c vorgegeben, genau dann den Wert 1 annimmt, wenn $(X + Y)$ den Wert 0 annimmt.

Die obigen Instanzfunktionen aus Tabelle 4.3 lassen sich beispielsweise der Reihe nach durch folgende Boolesche Ausdrücke darstellen:

- (a) $A \cdot B$,
- (b) $A + B$,
- (c) $\bar{A} + B$,
- (d) $A \cdot B + \bar{A} \cdot \bar{B}$.

Dass der jeweilige Boolesche Ausdruck tatsächlich eine alternative Darstellung der entsprechenden, durch eine Wertetabelle dargestellten Instanzfunktion ist, lässt sich überprüfen, indem der

Boolesche Ausdruck über alle Belegungen ausgewertet wird. Die Auswertung erfolgt auf analoge Weise wie in Tabelle 4.2 in Zusammenhang mit dem Nachweis des Distributivgesetzes demonstriert (vgl. S. 144). Dabei gilt es zwei Konventionen zu berücksichtigen, die der besseren Lesbarkeit halber eingesetzt werden und bei den obigen Booleschen Ausdrücken teilweise bereits umgesetzt sind: Einerseits werden konjunktive Verknüpfungen wie $A \cdot B$ oftmals ohne explizite Angabe des Verknüpfungssymbols notiert und kurz geschrieben als AB . Andererseits wird oft auf Klammersetzungen verzichtet und die Auswertungsreihenfolge durch eine Bindungsstärke der einzelnen Verknüpfungen festgelegt, wobei immer zuerst die stärkere Bindung ausgewertet wird. Bezüglich der Bindungsstärke gilt wie in der Aussagenlogik, dass $\bar{}$ stärker bindet als $\cdot$ und $\cdot$ stärker bindet als $+$. Klammern werden damit oftmals nur noch dann eingesetzt, wenn man die Reihenfolge der Auswertung bei einem Booleschen Ausdruck anders priorisieren möchte als dies von der Konvention vorgegeben wird. Insgesamt ist somit etwa der Boolesche Ausdruck $AB + \bar{A}\bar{B}$ die kurze, typischerweise verwendete Form des obigen Ausdrucks $A \cdot B + \bar{A} \cdot \bar{B}$, der seinerseits eine kürzere Form des Ausdrucks $(A \cdot B) + ((\bar{A}) \cdot (\bar{B}))$ darstellt (vgl. u. a. Crama & Hammer (2011), S. 11; Gotthardt (2005), S. 37).

Jede beliebige Instanzfunktion kann bereits alleine mittels Verwendung der drei Verknüpfungen $+$, $\cdot$ und $\bar{}$ als Boolescher Ausdruck dargestellt werden.⁵ Man sagt, dass die drei Verknüpfungen zusammen ein vollständiges Operatorensystem oder eine *Verknüpfungsbasis* bilden (vgl. u. a. Böhme (1981), S. 149 f.). Sie werden deshalb auch als Basisverknüpfungen bezeichnet. Trotzdem werden teilweise zusätzlich sogenannte abgeleitete Verknüpfungen eingeführt, um Boolesche Ausdrücke unter anderem „eleganter und übersichtlicher“ formulieren zu können. So werden etwa zur Darstellung von f_3 beziehungsweise f_4 aus Tabelle 4.3 unter anderem auch die Symbole $\rightarrow$ beziehungsweise $\leftrightarrow$ eingesetzt, sodass sich die Booleschen Ausdrücke $\bar{A} + B$ beziehungsweise $AB + \bar{A}\bar{B}$ kurz schreiben lassen als $A \rightarrow B$ beziehungsweise $A \leftrightarrow B$ (vgl. u. a. Hoffmann (2016), S. 98 f.). Allgemein gilt für zwei beliebige Boolesche Ausdrücke μ und ν :

$$\begin{aligned}\mu \rightarrow \nu &:= \bar{\mu} + \nu, \\ \mu \leftrightarrow \nu &:= (\mu \cdot \nu) + (\bar{\mu} \cdot \bar{\nu}).\end{aligned}$$

Bezüglich der Bindungsstärke wird festgelegt, dass $\rightarrow$ und $\leftrightarrow$ schwächer binden als $\bar{}$, $\cdot$ und $+$, sodass beispielsweise der Ausdruck $A + B \leftrightarrow C$ äquivalent ist mit $(A + B) \leftrightarrow C$.

Neben $\rightarrow$ und $\leftrightarrow$ gibt es weitere oftmals verwendete abgeleitete Verknüpfungen (vgl. u. a. Hoffmann (2016), S. 98). Die explizite Einführung einer bestimmten abgeleiteten Verknüpfung hängt typischerweise vom Zweck und dem Anwendungsgebiet ab. So ist es denn auch kein Zufall, dass in der vorliegenden Arbeit zusätzlich die obigen beiden Verknüpfungen $\rightarrow$ und

⁵Tatsächlich würde bereits alleine $+$ zusammen mit $\bar{}$ oder $\cdot$ zusammen mit $\bar{}$ ausreichen, um alle möglichen Instanzfunktionen darzustellen (vgl. u. a. Böhme (1981), S. 150).

X	$g(X)$
1	1
0	0

Tabelle 4.4.: Einstellige Instanzfunktion g .

$\leftrightarrow$ definiert werden: Die Implikation $\mu \rightarrow \nu$ bildet das Prinzip der Suffizienz von μ für ν beziehungsweise das Prinzip der Notwendigkeit von ν für μ nach, und die Äquivalenz $\mu \leftrightarrow \nu$ bildet das Prinzip des Bikonditionals nach. Angesichts des kausaltheoretischen Rahmens der vorhergehenden Kapitel nehmen diese Verknüpfungen deshalb eine wesentliche Rolle ein.

Eine Instanzfunktion kann durch unendlich viele Boolesche Ausdrücke dargestellt werden. Beispielsweise die einstellige Instanzfunktion $g : \mathcal{B} \rightarrow \mathcal{B}$, $\mathcal{B} = \{0, 1\}$, die durch die Wertetabelle 4.4 definiert ist, lässt sich unter anderem durch X , $X + X$ oder $X + X + X$ ausdrücken. Aus diesem Grund muss grundsätzlich zwischen den beiden Begriffen der Instanzfunktion und des Booleschen Ausdrucks, der eine Instanzfunktion darstellt, unterschieden werden, weil zwar jeder Boolesche Ausdruck einer Instanzfunktion entspricht, nicht aber umgekehrt. Eine Instanzfunktion bildet Nullen und Einsen auf Null und Eins ab und ein Boolescher Ausdruck ist eine Form, diese Abbildung wiederzugeben. Es wird daher auch gesagt, dass ein Boolescher Ausdruck eine Instanzfunktion *repräsentiert* beziehungsweise ein Repräsentant dieser Instanzfunktion ist (vgl. u. a. Hoffmann (2016), S. 97 ff.). Weil ein Boolescher Ausdruck unter Berücksichtigung des Definitionsbereiches $\mathcal{D} \subseteq \mathcal{B}^n$ die entsprechende Instanzfunktion $f : \mathcal{D} \rightarrow \mathcal{B}$ aber eindeutig repräsentiert, werden Instanzfunktionen und ihre Repräsentationen durch Boolesche Ausdrücke künftig der Einfachheit halber oftmals direkt gleichgesetzt. Das heisst, dass anstelle einer Instanzfunktion f und ihrer Repräsentation durch einen Booleschen Ausdruck μ , wie beispielsweise $\mu = AB + \bar{A}\bar{B}$, oft kurz $f(A, B) = AB + \bar{A}\bar{B}$ geschrieben wird. Oder f wird einfach mit dem Komplementsymbol als $\bar{f}$ geschrieben, wenn damit die Instanzfunktion g dargestellt werden soll, die durch den Booleschen Ausdruck $\bar{\mu} = \overline{AB + \bar{A}\bar{B}}$ repräsentiert wird. Das heisst, $\bar{f}$ ist das Komplement der Instanzfunktion f , sodass $\bar{f}$ genau dann den Wert 1 beziehungsweise 0 annimmt, wenn f den Wert 0 beziehungsweise 1 annimmt. Dabei gilt wohlgemerkt, dass $\bar{f}$ denselben Definitionsbereich hat wie f .

Im Folgenden werden einige spezielle Formen von Booleschen Ausdrücken definiert, die allgemein in der Booleschen Algebra von grossem Nutzen sind und auch im Zusammenhang mit der Konfigurationsalgebra als mathematisches Fundament des konfiguralen kausalen Modellierens immer wieder Verwendung finden werden.

Disjunktive Normalformen und ihre Bestandteile

Eine besondere Form der Darstellung von Instanzfunktionen mittels Boolescher Ausdrücke ist die *disjunktive Normalform* (DNF), die analog zur in einfachen Minimalen Theorien vorkommenden Disjunktion aus einer Disjunktion von Konjunktionstermen besteht (vgl. bspw. Nelson et al. (1995), S. 94). Dabei ist ein Konjunktionsterm eine konjunktive Verknüpfung von Literalen, wobei das gleiche Literal nur als positives Literal oder nur als negatives Literal in der Verknüpfung vorhanden ist.

Definition 4.4 (*Disjunktive Normalform (DNF)*)

Sei $\mathcal{B} = \{0, 1\}$, $\mathcal{D} \subseteq \mathcal{B}^n$ und $f : \mathcal{D} \rightarrow \mathcal{B}$ eine Instanzfunktion von n Variablen $X_1, \dots, X_n$. Die Instanzfunktion f ist genau dann in disjunktiver Normalform (DNF), wenn es sich um eine Disjunktion von Konjunktionstermen handelt, wobei ein Konjunktionsterm eine Konjunktion ist, die die gleiche Variable X_i , $1 \leq i \leq n$, höchstens einmal entweder als positives oder negatives Literal enthält.

Enthält ein Konjunktionsterm jede Variable X_i der n Variablen entweder als positives Literal X_i oder als negatives Literal $\bar{X}_i$, wird der Konjunktionsterm als *Minterm* bezeichnet (vgl. u. a. Böhme (1981), S. 139 ff.). Ein Minterm wird weiter Minterm einer Instanzfunktion f genannt, wenn er für ein Element $(X_1, \dots, X_n)$ des Definitionsbereiches von f mit $f(X_1, \dots, X_n) = 1$ den Wert 1 annimmt. So sind beispielsweise AB und $\bar{A}\bar{B}$ Minterme der Instanzfunktion $f_4(A, B)$ aus Tabelle 4.3d.

Definition 4.5 (*Minterm*)

Sei $\mathcal{B} = \{0, 1\}$, $\mathcal{D} \subseteq \mathcal{B}^n$ und $f : \mathcal{D} \rightarrow \mathcal{B}$ eine Instanzfunktion von n Variablen $X_1, \dots, X_n$. Ein Konjunktionsterm $\mathbf{m}$, der jede Variable X_i , $1 \leq i \leq n$, entweder als positives Literal X_i oder als negatives Literal $\bar{X}_i$ enthält, heisst Minterm. Nimmt $\mathbf{m}$ für ein Element $(X_1, \dots, X_n) \in \mathcal{D}$ mit $f(X_1, \dots, X_n) = 1$ den Wert 1 an, nennt man $\mathbf{m}$ Minterm von f .

Ist eine Instanzfunktion f in disjunktiver Normalform dargestellt, deren Konjunktionsterme genau alle paarweise verschiedenen Minterme von f sind, spricht man von einer *kanonischen disjunktiven Normalform* (vgl. bspw. Nelson et al. (1995), S. 95). Mit $AB + \bar{A}\bar{B}$ hat man beispielsweise die kanonische disjunktive Normalform der Instanzfunktion $f_4(A, B)$ aus Tabelle 4.3d.

Definition 4.6 (*Kanonische Disjunktive Normalform (KDNF)*)

Sei $\mathcal{B} = \{0, 1\}$, $\mathcal{D} \subseteq \mathcal{B}^n$ und $f : \mathcal{D} \rightarrow \mathcal{B}$ eine Instanzfunktion von n Variablen. Die Instanzfunktion f ist genau dann in kanonischer disjunktiver Normalform (KDNF) dargestellt, wenn

es sich um eine DNF handelt, deren Konjunktionsterme genau alle paarweise verschiedenen Minterme von f sind.⁶

Jede Instanzfunktion $f : \mathcal{B}^n \rightarrow \mathcal{B}$ besitzt genau eine KDNF. Dieser Zusammenhang zwischen $f : \mathcal{B}^n \rightarrow \mathcal{B}$ und ihrer eindeutigen Darstellbarkeit durch eine KDNF ist bekannt als sogenannter *Hauptsatz der Booleschen Algebra* (vgl. u. a. Böhme (1981), S. 140). Die KDNF von f kann dabei mehr oder weniger direkt aus der Wertetabelle von f abgelesen werden (vgl. bspw. Staab (2007), S. 71 f.). Dazu betrachtet man die aus 0 und 1 bestehenden Elemente des Definitionsbereiches von f , für die f den Funktionswert 1 annimmt. Jedes dieser n -Tupel $(X_1, \dots, X_n)$ wird in einen aus n Literalen bestehenden Minterm übersetzt, indem man X_i beziehungsweise $\bar{X}_i$ in den Minterm aufnimmt, wenn an i -ter Stelle des n -Tupels $(X_1, \dots, X_n)$ eine 1 beziehungsweise eine 0 steht. Verknüpft man all diese (paarweise verschiedenen) Minterme von f schliesslich disjunktiv, ergibt sich daraus die KDNF von f . Tabelle 4.5 zeigt konkret die Bildung der Minterme von f_4 , die genau für jene Elemente $(A, B) \in \mathcal{B}^2$ aus dem Definitionsbereich von f_4 den Funktionswert 1 aufweist, für die $A = B$ gilt. Aus dem Tupel $(1, 1)$ wird der Minterm AB gebildet und aus dem Tupel $(0, 0)$ der Minterm $\bar{A}\bar{B}$. Verknüpft man beide Minterme konjunktiv, ergibt sich die gewünschte KDNF $AB + \bar{A}\bar{B}$ von f_4 .

A	B	$f_4(A, B)$	Minterm von f_4
1	1	1	AB
1	0	0	
0	1	0	
0	0	1	$\bar{A}\bar{B}$

Tabelle 4.5.: Bildung der Instanzfunktion $f_4(A, B) = AB + \bar{A}\bar{B}$ in KDNF.

Während jede Instanzfunktion $f : \mathcal{B}^n \rightarrow \mathcal{B}$ genau eine kanonische DNF besitzt, kann sie verschiedene äquivalente DNF aufweisen. Zwei Boolesche Ausdrücke μ und ν sind genau dann äquivalent, wenn sie dieselbe Instanzfunktion $f : \mathcal{B}^n \rightarrow \mathcal{B}$ repräsentieren beziehungsweise wenn für alle Belegungen der Variablen $X_1, \dots, X_n$ der Instanzfunktion gilt, dass μ und ν identische Werte annehmen (vgl. bspw. Hoffmann (2016), S. 101; Whitesitt (1964), S. 39). So lässt sich etwa die Instanzfunktion $f_2(A, B) = A + B$ aus Tabelle 4.3 über die drei Minterme AB , $A\bar{B}$ und $\bar{A}B$ auch darstellen als $f_2(A, B) = AB + A\bar{B} + \bar{A}B$. Bei beiden Booleschen Ausdrücken $A + B$ und $AB + A\bar{B} + \bar{A}B$ handelt es sich um disjunktive Normalformen, die dieselbe Funktion f_2 repräsentieren, wobei letztere zusätzlich kanonisch ist. Die beiden Booleschen Ausdrücke sind äquivalent und nehmen für jede Belegung von A und B identische Werte an.

⁶In der Literatur zur Booleschen Algebra ist mit der disjunktiven Normalform oft die hier als kanonisch bezeichnete disjunktive Normalform gemeint (vgl. z. B. Whitesitt (1964), S. 36 f.; Gumm & Pogutke (1981), S. 31 f.).

Eine weitere speziell ausgezeichnete disjunktive Normalform ist die sogenannte *disjunktive Minimalform* (DMF). Eine Instanzfunktion ist in DMF dargestellt, wenn es keine äquivalente DNF gibt, die weniger Disjunkte aufweist und deren Anzahl Literale pro Konjunktionsterm, falls die Anzahl disjunktiv verknüpfter Konjunktionsterme identisch ist, kleiner ist (vgl. u. a. Garnier & Taylor (2002), S. 476). Für die Instanzfunktion f_2 mit der Wertetabelle 4.3b ist $A + B$ die DMF. Die DMF kann beispielsweise gefunden werden, indem man die KDNF mittels der Booleschen Gesetzen so umformt, dass sie minimal wird:

$$\begin{aligned} f_2(A, B) &= AB + A\bar{B} + \bar{A}B \\ &= AB + A\bar{B} + AB + \bar{A}B \\ &= A(B + \bar{B}) + B(A + \bar{A}) \\ &= A \cdot 1 + B \cdot 1 \\ &= A + B \end{aligned}$$

Es soll an dieser Stelle erwähnt sein, dass im Unterschied zur speziell ausgezeichneten KDNF eine Instanzfunktion mehrere disjunktive Minimalformen besitzen kann. Im nächsten Kapitel werden entsprechende Beispiele präsentiert.

Mit den speziell ausgezeichneten Normalformen gibt es neben den Mintermen ebenfalls weitere speziell ausgezeichnete (Boolesche) Terme. Einer dieser Terme ist der *Implikant* (vgl. u. a. Kunz & Stoffel (1997), S. 9; Becker & Molitor (2008), S. 135). Ein Konjunktionsterm einer als Boolescher Ausdruck dargestellten Instanzfunktion f , für den gilt, dass er für mindesten ein Element aus dem Definitionsbereich von f den Wert 1 annimmt und wenn er den Wert 1 annimmt, auch der Funktionswert 1 ist, wird als *Implikant* von f bezeichnet. Von einem solchen Konjunktionsterm wird gesagt, dass er f impliziert.

Definition 4.7 (Implikant)

Sei $\mathcal{B} = \{0, 1\}$, $\mathcal{D} \subseteq \mathcal{B}^n$ und $f : \mathcal{D} \rightarrow \mathcal{B}$ eine Instanzfunktion von n Variablen. Ein aus Literalen von Variablen aus f bestehender Konjunktionsterm, für den gilt, dass

- (a) er für mindestens ein $(X_1, \dots, X_n) \in \mathcal{D}$ den Wert 1 annimmt sowie
- (b) wenn er den Wert 1 annimmt, auch der Funktionswert 1 ist,

heißt Implikant von f .

Die drei Minterme AB , $A\bar{B}$ und $\bar{A}B$ obiger Instanzfunktion f_2 sind somit Implikanten dieser Funktion. Hat man einen Implikanten einer Instanzfunktion, der nicht weiter verkürzt werden kann, wird er Primterm oder *Primimplikant* genannt (vgl. bspw. Eschermann (1993), S. 82;

Kunz & Stoffel (1997), S. 9). Mit $A + B$ liegt beispielsweise eine aus den zwei Primimplikanten A und B bestehende DMF von f_2 vor. Ein Primimplikant kann wie folgt definiert werden:

Definition 4.8 (Primimplikant)

Sei $\mathcal{B} = \{0, 1\}$, $\mathcal{D} \subseteq \mathcal{B}^n$ und $f : \mathcal{D} \rightarrow \mathcal{B}$ eine Instanzfunktion. Ein Implikant von f , der kein Implikant von f mehr ist sobald mindestens ein (beliebiges) Literal aus dem Implikanten entfernt wird, heisst Primimplikant von f .

Primimplikanten einer Instanzfunktion f sind kürzeste Konjunktionsterme, die f implizieren. Ähnlich wie ein Implikant von f die Instanzfunktion f impliziert, implizieren Minterme von f Primimplikanten von f . So gilt etwa für den Minterm AB der Instanzfunktion f_2 und die beiden Primimplikanten A und B von f_2 , dass wenn AB den Wert 1 annimmt, A und B den Wert 1 annehmen. Von den Primimplikanten ausgehend wird gesagt, dass der Primimplikant A beziehungsweise der Primimplikant B den Minterm AB *abdeckt*. Ein Primimplikant deckt einen Minterm genau dann ab, wenn der Primimplikant Teil des Minterms ist. Eine Abdeckung eines Konjunktionsterms durch einen verkürzten oder zumindest gleich langen Konjunktionsterm lässt sich also daran erkennen, dass letzterer in ersterem enthalten ist oder beide identisch sind. Wird ein Minterm einer Instanzfunktion nur von einem Primimplikanten abgedeckt, wird dieser Primimplikant auch Kernprimterm oder *Kernprimimplikant* genannt (vgl. bspw. Eschermann (1993), S. 82):

Definition 4.9 (Kernprimimplikant)

Sei $\mathcal{B} = \{0, 1\}$, $\mathcal{D} \subseteq \mathcal{B}^n$ und $f : \mathcal{D} \rightarrow \mathcal{B}$ eine Instanzfunktion. Sei weiter $\mathbf{m}$ ein Minterm dieser Instanzfunktion f . Ein Primimplikant von f heisst Kernprimimplikant von f , wenn $\mathbf{m}$ bis auf diesen Primimplikanten von keinem anderen Primimplikanten von f abgedeckt wird.

So wie ein Implikant einen Minterm abdecken kann, deckt eine Disjunktion von Implikanten eine Instanzfunktion genau dann ab, wenn es für jeden Minterm der Instanzfunktion einen Implikanten gibt, der den Minterm abdeckt. Dies gilt deshalb, weil die Disjunktion der Minterme selbst eine Abdeckung der Instanzfunktion ist. Dementsprechend ist etwa $AB + A\bar{B} + \bar{A}B$ oder $A + B$ eine Abdeckung von f_2 . Entspricht eine Abdeckung wie $A + B$ von f_2 der Darstellung einer Instanzfunktion f in DMF, das heisst, ist eine Abdeckung von f keine Abdeckung mehr, sobald ein Implikant daraus entfernt wird, sagt man auch, dass sie *irredundant* oder redundanzfrei ist (vgl. u. a. Kunz & Stoffel (1997), S. 9)⁷:

⁷Im Zusammenhang mit der Abdeckung wird künftig typischerweise der Begriff „irredundant“ verwendet, um all-fällige Verwechslungen mit dem noch folgenden Begriff der redundanzfreien disjunktiven Normalform vorzubeugen (vgl. Def. 4.15, S. 170.). Denn es wird sich herausstellen, dass jede redundanzfreie disjunktive Normalform auch eine irredundante Abdeckung ist, aber nicht jede irredundante Abdeckung eine redundanzfreie disjunktive Normalform.

Definition 4.10 (Irredundante Abdeckung)

Sei $\mathcal{B} = \{0, 1\}$, $\mathcal{D} \subseteq \mathcal{B}^n$ und $f : \mathcal{D} \rightarrow \mathcal{B}$ eine Instanzfunktion von n Variablen. Eine disjunktive Verknüpfung von Implikanten von f , sodass wenn f den Wert 1 annimmt, auch die Disjunktion von Implikanten den Wert 1 annimmt, heisst Abdeckung von f . Eine Abdeckung, die keine Abdeckung mehr ist, sobald mindestens ein Implikant daraus entfernt wird, heisst irredundante Abdeckung von f .

Totale und partielle Instanzfunktionen

Neben den verschiedenen Darstellungsformen von Instanzfunktionen werden künftig auch die Unterscheidung zwischen einer vollständig definierten oder *totalen* Instanzfunktion und einer unvollständig definierten, partiell definierten oder kurz *partiellen Instanzfunktion* sowie damit verknüpfte Bereiche des Definitions- und des Wertebereichs eine wesentliche Rolle spielen. Der Unterschied zwischen einer totalen und einer partiellen Instanzfunktion lässt sich durch folgende Definition ausdrücken (vgl. u. a. Becker & Molitor (2008), S. 26):

Definition 4.11 (Totale und partielle Instanzfunktion)

Sei $\mathcal{B} = \{0, 1\}$ und $\mathcal{D} \subset \mathcal{B}^n$. Eine Instanzfunktion $f : \mathcal{B}^n \rightarrow \mathcal{B}$ heisst totale Instanzfunktion (von n Variablen) und eine Instanzfunktion $f : \mathcal{D} \rightarrow \mathcal{B}$ heisst partielle Instanzfunktion (von n Variablen über $\mathcal{D}$).

Die betrachteten Beispiele von Instanzfunktionen aus vorigem Abschnitt sind jeweils totale Instanzfunktionen, da sie für alle möglichen Belegungen aus $\mathcal{B}^n$ ausgewertet werden. So sind etwa die Instanzfunktionen aus Tabelle 4.3 total, da sie für alle vier möglichen Belegungen aus $\mathcal{B}^2$ einen Funktionswert ausgeben. Im Gegensatz dazu ist eine partielle Instanzfunktion nur für einen Teil aus $\mathcal{B}^n$ definiert. Dargestellt durch eine Wertetabelle wird dies zum Ausdruck gebracht, indem ein Teil der 2^n Zeilen gar nicht erst in der Tabelle erscheint. Die Instanzfunktion g mit der Wertetabelle 4.6 ist ein Beispiel einer solchen partiellen Instanzfunktion, für die die Zuweisung des Wertes 0 zu A und B sowie ein dazugehöriger Funktionswert nicht definiert sind.

A	B	$g(A, B)$
1	1	1
1	0	0
0	1	0

Tabelle 4.6.: Beispiel einer partiellen Instanzfunktion der Variablen A und B .

Im Zusammenhang mit der Unterteilung von totalen und partiellen Instanzfunktionen werden bestimmte Bereiche der Definitions- und Wertebereiche speziell gekennzeichnet:

Definition 4.12 (Instantiierungs-, Noninstantiierungs- und ND-Bereich)

Sei $\mathcal{B} = \{0, 1\}$ und $\mathcal{D} \subseteq \mathcal{B}^n$. Sei weiter $f : \mathcal{D} \rightarrow \mathcal{B}$ eine Instanzfunktion, die total oder partiell ist.

- Der Bereich $\mathcal{I}(f) = \{X \in \mathcal{D} \mid f(X) = 1\}$ heisst Instantiierungsbereich von f .
- Der Bereich $\mathcal{NI}(f) = \{X \in \mathcal{D} \mid f(X) = 0\}$ ist der Noninstantiierungsbereich von f .
- $\mathcal{I}(f)$ und $\mathcal{NI}(f)$ bilden zusammen den Definitionsbereich $\mathcal{D}(f)$ von f , das heisst, $\mathcal{D}(f) = \mathcal{I}(f) \cup \mathcal{NI}(f) = \mathcal{D}$.
- Schliesslich ist $\mathcal{ND}(f) = \mathcal{B}^n \setminus \mathcal{D} = \{X \in \mathcal{B}^n \mid X \notin \mathcal{D}\}$ der nicht definierte Bereich von f .

In der Schaltalgebra entsprechen die drei Mengen $\mathcal{I}(f)$, $\mathcal{NI}(f)$ und $\mathcal{ND}(f)$ der Reihe nach der ON-Menge, der OFF-Menge und dem Don't-Care-Bereich (vgl. u. a. Becker & Molitor (2008), S. 26). Der Unterschied zu den hier definierten Mengen liegt in den Bezeichnungen und den Interpretationen, die dem hier vorliegenden Kontext angepasst werden. Wie sich später zeigen wird, ist diese Anpassung auch deshalb sinnvoll, weil gerade die Bezeichnung „Don't Care“ im vorliegenden Kontext etwas irreführend ist.

Veranschaulicht anhand der partiellen Instanzfunktion g aus Tabelle 4.6 entspricht die Zuweisung von 1 sowohl zu A als auch zu B dem Instantiierungsbereich $\mathcal{I}(g)$ von g . Die Zuweisung von 1 beziehungsweise 0 zu A und 0 beziehungsweise 1 für B beschreibt den Noninstantiierungsbereich $\mathcal{NI}(g)$ von g . Schliesslich ist die Zuweisung von 0 sowohl zu A als auch zu B im nicht definierten Bereich $\mathcal{ND}(g)$ von g . Werden diese Belegungen mittels Tupeln dargestellt, wobei die erste Stelle jeweils der Wertezuweisung zu A und die zweite Stelle der Wertezuweisung zu B entspricht, können die drei Mengen $\mathcal{I}(g)$, $\mathcal{NI}(g)$ und $\mathcal{ND}(g)$ geschrieben werden als

- $\mathcal{I}(g) = \{(1, 1)\}$,
- $\mathcal{NI}(g) = \{(1, 0), (0, 1)\}$ und
- $\mathcal{ND}(g) = \{(0, 0)\}$.

Eine totale Instanzfunktion ist bereits alleine durch ihren Instantiierungsbereich oder durch ihren Noninstantiierungsbereich vollständig bestimmt, da der nicht definierte Bereich jeweils identisch

mit der leeren Menge ist. Bei einer partiellen Instanzfunktion f sind dazu zwei der drei Mengen $\mathcal{I}(f)$, $\mathcal{NI}(f)$ oder $\mathcal{ND}(f)$ notwendig (vgl. u. a. Becker & Molitor (2008), S. 26). So ist f etwa folgendermassen vollständig bestimmt:

$$f(X_1, \dots, X_n) = \begin{cases} 1, & (X_1, \dots, X_n) \in \mathcal{I}(f) \\ 0, & (X_1, \dots, X_n) \in \mathcal{NI}(f) \end{cases} \quad (4.1)$$

Im konkreten Fall der obigen partiellen Instanzfunktion g mit der Wertetabelle 4.6 wäre eine vollständige Bestimmung gegeben durch:

$$g(A, B) = \begin{cases} 1, & (A, B) \in \mathcal{I}(g) = \{(1, 1)\} \\ 0, & (A, B) \in \mathcal{NI}(g) = \{(1, 0), (0, 1)\} \end{cases} \quad (4.2)$$

Analog zu einer totalen Instanzfunktion lässt sich eine partielle Instanzfunktion mit Booleschen Ausdrücken beschreiben, die jedoch im Gegensatz zur totalen Instanzfunktion nicht über alle Belegungen, sondern nur über einen Teil davon ausgewertet werden. Dementsprechend gilt es bei einer mit Booleschen Ausdrücken dargestellten partiellen Instanzfunktion besonders auf den Definitionsbereich zu achten. So lässt sich etwa die durch die Wertetabelle 4.6 repräsentierte partielle Instanzfunktion g unter Berücksichtigung des Definitionsbereiches $\mathcal{D}(g) = \mathcal{I}(g) \cup \mathcal{NI}(g) = \{(1, 1), (1, 0), (0, 1)\} = \mathcal{B}^2 \setminus \{(0, 0)\}$ folgendermassen repräsentieren:

$$g(A, B) = AB, \quad (A, B) \in \mathcal{B}^2 \setminus \{(0, 0)\} \quad (4.3)$$

Im Umgang mit partiellen Instanzfunktionen werden oft sogenannte *Erweiterungen* dieser Funktionen verwendet (vgl. u. a. Boros et al. (2000); Crama et al. (1988); Crama & Hammer (2011)). Eine Erweiterung einer partiellen Instanzfunktion f ist eine totale Instanzfunktion f' , die (1) für jene Belegungen den Funktionswert 1 annimmt, bei denen f ebenfalls den Funktionswert 1 annimmt, (2) für jene Belegungen den Funktionswert 0 annimmt, bei denen f ebenfalls den Funktionswert 0 annimmt und (3) für jene Belegungen entweder 0 oder 1 annimmt, die ausserhalb des Definitionsbereiches von f liegen:

Definition 4.13 (Erweiterung)

Sei $\mathcal{B} = \{0, 1\}$, $\mathcal{D} \subset \mathcal{B}^n$ und $f : \mathcal{D} \rightarrow \mathcal{B}$ eine partielle Instanzfunktion über $\mathcal{D}$. Eine Erweiterung von f ist eine totale Instanzfunktion f' , für die gilt: $\mathcal{I}(f) \subseteq \mathcal{I}(f')$ und $\mathcal{NI}(f) \subseteq \mathcal{NI}(f')$.

Von den möglichen Erweiterungen einer partiellen Instanzfunktion f gibt es zwei, die speziell als f^+ und f^- ausgezeichnet werden können und folgendermassen definiert sind (vgl. u. a. Crama et al. (1988), S. 302 f.):

$$f^+(X_1, \dots, X_n) = \begin{cases} 1, & (X_1, \dots, X_n) \notin \mathcal{NI}(f) \\ 0, & (X_1, \dots, X_n) \in \mathcal{NI}(f) \end{cases} \quad (4.4)$$

A	B	$g^+(A, B)$	A	B	$g^-(A, B)$
1	1	1	1	1	1
1	0	0	1	0	0
0	1	0	0	1	0
0	0	1	0	0	0

(a)
(b)

Tabelle 4.7.: Die beiden speziell ausgezeichneten Erweiterungen g^+ und g^- der partiellen Instanzfunktion g mit Wertetabelle 4.6.

und

$$f^-(X_1, \dots, X_n) = \begin{cases} 1, & (X_1, \dots, X_n) \in \mathcal{I}(f) \\ 0, & (X_1, \dots, X_n) \notin \mathcal{I}(f) \end{cases} \quad (4.5)$$

Bei der Erweiterung f^+ von f wird die partielle Instanzfunktion f derart erweitert, dass für die Belegungen des Bereiches $\mathcal{ND}(f)$ jeweils der Funktionswert 1 ausgegeben wird. Für die Erweiterung f^- von f ist der Funktionswert 0 bei den Belegungen des Bereiches $\mathcal{ND}(f)$. Konkret anhand der partiellen Instanzfunktion g aus Tabelle 4.6 veranschaulicht, ist g^+ durch die Wertetabelle 4.7a und g^- durch die Wertetabelle 4.7b dargestellt.

4.2. Die Konfigurationsalgebra als algebraisches Fundament konfiguraler Methoden

Es stellt sich nun die Frage, inwiefern die Konfigurationsalgebra eingesetzt werden kann, um aus Konfigurationstabellen Minimale Theorien abzuleiten. Die Konfigurationsalgebra wurde durch die Betrachtung der verschiedenen Verknüpfungsformen von an- und abwesenden Faktoren im kausaltheoretischen Kontext motiviert und kann basierend auf dieser Interpretation als mathematische Struktur aufgefasst werden, die das Instantiierungsverhalten von Faktoren erfasst und beschreibt. Mit der Konfigurationsalgebra wird unter anderem ermöglicht, Konfigurationstabellen als Instanzfunktionen zu interpretieren, diesen Instanzfunktionen Boolesche Ausdrücke zuzuordnen und letztere mit den Gesetzen der Booleschen Algebra umzuformen. In diesem Abschnitt soll dieser Zusammenhang zwischen dem konfiguralen kausalen Modellieren und der Konfigurationsalgebra vermittelt werden. Damit wird unter anderem die Grundlage dafür geschaffen, um mithilfe der Konfigurationsalgebra systematisch Kausalmodelle zu finden.

4.2.1. Instanzfunktionen als Darstellungen des Instantiierungsverhaltens von Faktoren

Der Zusammenhang zwischen einer Konfigurationstabelle und einer Instanzfunktion soll konkret anhand eines Beispiels veranschaulicht werden. Betrachtet wird dazu die Kausalstruktur aus Abbildung 4.8 mit der dazugehörigen homogenen und maximal diversen Konfigurationstabelle $\mathcal{K}_{4,8}$ über die Faktoren der Menge $\mathbf{G} = \{A, B, C\}$.

Von den insgesamt $2^3 = 8$ logisch möglichen Kombinationen von an- und abwesenden Faktoren aus $\mathbf{G}$ gibt es gemäss $\mathcal{K}_{4,8}$ vier Konfigurationen, die in Fällen instantiiert sind. Von den restlichen vier Konfigurationen gibt es keine Fälle. Die vier in $\mathcal{K}_{4,8}$ enthaltenen Konfigurationen lassen sich verstehen als Instantiierungsbereich $\mathcal{I}(f)$ einer Instanzfunktion f dreier Variablen A , B und C . Das heisst, f nimmt genau dann den Wert 1 an (und 0 sonst), wenn A , B und C eine der vier Belegungen annimmt, die den Konfigurationen aus $\mathcal{K}_{4,8}$ entsprechen. Diese Auffassung, die in der Folge oftmals kurz auch notiert wird als $\mathcal{I}(f) = \mathcal{K}_{4,8}$, ermöglicht die Transformation von $\mathcal{K}_{4,8}$ in die Instanzfunktion f , die in Tabelle 4.9b als Wertetabelle dargestellt ist.

Die hier mittels eines konkreten Beispiels dargestellte Transformation einer Konfigurationstabelle in eine Instanzfunktion kann für jede beliebige Konfigurationstabelle $\mathcal{K}_{\mathbf{G}}$ durchgeführt werden. So gesehen ist $\mathcal{K}_{\mathbf{G}}$ über $\mathbf{G} = \{X_1, \dots, X_n\}$ mit k Konfigurationen nichts anderes als eine tabellarische Darstellung des Instantiierungsbereiches $\mathcal{I}(f)$ einer Instanzfunktion $f : \mathcal{B}^n \rightarrow \mathcal{B}$ von n Variablen. Präsentiert man f detailliert durch eine Wertetabelle, nimmt sie genau für jene k Belegungen aus $\mathcal{B}^n$ den Funktionswert 1 an, die identisch mit den Konfigurationen aus $\mathcal{K}_{\mathbf{G}}$ sind. Für alle anderen Belegungen nimmt sie den Funktionswert 0 an. Wiederum kurz notiert gilt $\mathcal{I}(f) = \mathcal{K}_{\mathbf{G}}$.

Aus rein mathematischer Sicht ist eine Instanzfunktion von n Variablen vorerst nichts anderes als eine Funktion, die n -Tupel von Nullen und Einsen auf die Werte 0 und 1 abbildet. In der konkreten Anwendung des konfiguralen kausalen Modellierens jedoch, bei der eine Instanz-



Abbildung 4.8.: Kausalstruktur mit dazugehöriger homogener und maximal diverser Konfigurationstabelle $\mathcal{K}_{4,8}$.

				A	B	C	$f(A, B, C)$
				1	1	1	1
				1	1	0	0
				1	0	1	1
				1	0	0	0
				0	1	1	1
				0	1	0	0
				0	0	1	0
				0	0	0	1
$\mathcal{K}_{4,8}$	A	B	C				
c_1	1	1	1				
c_2	1	0	1				
c_3	0	1	1				
c_4	0	0	0				
(a)				(b)			

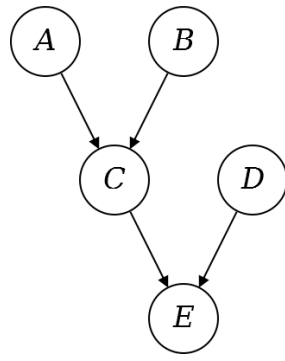
Tabelle 4.9.: Transformation der homogenen und maximal diversen Konfigurationstabelle $\mathcal{K}_{4,8}$ aus Abbildung 4.8 in eine Instanzfunktion f .

funktion über obige Transformation einer Konfigurationstabelle definiert wird, bringt sie das Auftrettsverhalten von Faktoren gemäss der Konfigurationstabelle zum Ausdruck. Dabei werden die Variablen der Instanzfunktion als Faktoren interpretiert, die bei einer Wertezuordnung von 1 als instantiiert oder gegeben und bei einer Wertezuordnung von 0 als nicht instantiiert oder nicht gegeben betrachtet werden. Aus diesem Grund werden Variablen künftig als Faktoren bezeichnet, wenn eine Instanzfunktion primär im Anwendungskontext des konfigurationsalen kausalen Modellierens betrachtet wird. Wird eine solche Instanzfunktion weiter durch ihre Wertetabelle repräsentiert, kann diese Wertetabelle im Anwendungskontext des konfigurationsalen kausalen Modellierens gedeutet werden als *Verhaltensmuster* von n verschiedenen Faktoren. Konkret wird als Verhaltensmuster also eine Wertetabelle bezeichnet, die eine Instanzfunktion f repräsentiert, die ihrerseits durch die Transformation einer Konfigurationstabelle $\mathcal{K}_F$ zustande kommt oder kurz: für die $\mathcal{I}(f) = \mathcal{K}_F$ gilt.

Eine Instanzfunktion f kann bekanntlich durch verschiedene Darstellungstypen repräsentiert werden. Neben Wertetabellen ist vor allem die Darstellung als Boolescher Ausdruck von grossem Nutzen. Mit obiger Transformation der Konfigurationen in eine Instanzfunktion, kann nun auch hier alternativ ein Boolescher Ausdruck verwendet werden. Der Instanzfunktion f mit der Wertetabelle 4.9 etwa kann gemäss dem Hauptsatz der Booleschen Algebra ein Boolescher Ausdruck in kanonischer disjunktiver Normalform zugeordnet werden, indem die Minterme von f disjunktiv verknüpft werden (vgl. S. 150):

$$f(A, B, C) = ABC + A\bar{B}C + \bar{A}BC + \bar{A}\bar{B}\bar{C}, \quad (A, B, C) \in \mathcal{B}^3 \quad (4.6)$$

Angeichts der obigen Transformation von $\mathcal{K}_{4,8}$ in die Instanzfunktion f und der damit einhergehenden Auffassung, dass f das Instantiierungsverhalten von Faktoren wiedergibt, kann (4.6),



$\mathcal{K}_{4,10}$	A	B	C	D	E
c_1	1	1	1	1	1
c_2	1	1	1	0	1
c_3	1	0	1	1	1
c_4	1	0	1	0	1
c_5	0	1	1	1	1
c_6	0	1	1	0	1
c_7	0	0	0	1	1
c_8	0	0	0	0	0

Abbildung 4.10.: Kausalkette mit der dazugehörigen homogenen und maximal diversen Konfigurationstabelle $\mathcal{K}_{4,10}$ über $\mathbf{H} = \{A, B, C, D, E\}$.

das heisst f repräsentiert mittels eines Booleschen Ausdrucks, als eine zwischen A , B und C geltende *Verhaltensregel* gedeutet werden, die das Verhaltensmuster der Faktoren beschreibt. Als Verhaltensregel wird ein Boolescher Ausdruck bezeichnet, der eine Instanzfunktion f repräsentiert, für die $\mathcal{I}(f) = \mathcal{K}_F$ gilt. Obige Verhaltensregel (4.6) bringt zum Ausdruck, dass A , B und C gemeinsam auftreten oder A und C gemeinsam auftreten, während B abwesend ist oder B und C gemeinsam auftreten, während A abwesend ist oder keiner der drei Faktoren anwesend ist. Die Verhaltensregel ist genau dann erfüllt – das heisst der Funktionswert ist genau dann 1 – wenn die drei Faktoren eine der vier Belegungen annehmen, die in $\mathcal{K}_{4,8}$ als Konfigurationen enthalten sind.

Insgesamt kann demnach eine Instanzfunktion f im Kontext des konfiguralen kausalen Modellierens verstanden werden als (mathematische) Darstellung des Instantiierungsverhaltens von Faktoren. Dieses Verhalten kann jeweils durch Boolesche Ausdrücke oder mittels Wertetabellen repräsentiert werden. In der konkreten Anwendung des konfiguralen kausalen Modellierens entspricht dabei ein Boolescher Ausdruck einer Verhaltensregel und eine Wertetabelle einem Verhaltensmuster der Faktoren. Der Zusammenhang zwischen einer Instanzfunktion (a), dargestellt durch eine Wertetabelle, und derselben Instanzfunktion (b), dargestellt durch einen Booleschen Ausdruck, kann also so gedeutet werden, dass (a) gewissermassen im Detail angibt, wie die Instanzen verschiedener Faktoren gegeben beziehungsweise nicht gegeben sein können, damit die durch (b) dargestellte Verhaltensregel erfüllt oder verletzt ist.

Von dieser Interpretation ausgehend sollen zwei weitere Beispiele betrachtet werden. Zum einen die homogene und maximal diverse Konfigurationstabelle $\mathcal{K}_{4,10}$ über $\mathbf{H} = \{A, B, C, D, E\}$ aus Abbildung 4.10, der eine Kausalkette zugrunde liegt.

	$\mathcal{K}_{4,11}$	X	Y
$\bigcirc X$	c_1	1	1
$\bigcirc Y$	c_2	1	0
	c_3	0	1
	c_4	0	0

Abbildung 4.11.: Homogene und maximal diverse Konfigurationstabelle $\mathcal{K}_{4,11}$ über $\mathbf{J} = \{X, Y\}$, wobei X und Y in keinem kausalen Zusammenhang stehen.

Die folgende Instanzfunktion g in KDNF aus (4.7) gibt die Verhaltensregel wieder, der die Faktoren aus $\mathbf{H} = \{A, B, C, D, E, F\}$ genügen:

$$g(A, B, C, D, E) = ABCDE + ABC\bar{D}E + A\bar{B}CDE + A\bar{B}C\bar{D}E + \bar{A}BCDE + \bar{A}BC\bar{D}E + \bar{A}\bar{B}CDE + \bar{A}\bar{B}C\bar{D}E, \quad (A, B, C, D, E) \in \mathcal{B}^5 \quad (4.7)$$

Als weiteres Beispiel soll zum anderen die homogene und maximal diverse Konfigurationstabelle $\mathcal{K}_{4,11}$ über $\mathbf{J} = \{X, Y\}$ aus Abbildung 4.11 betrachtet werden, wobei X und Y in keinem kausalen Zusammenhang stehen. Letzteres zeigt sich im entsprechenden Graphen dadurch, dass er keine Kanten enthält.

Die folgende Instanzfunktion h in KDNF beschreibt das Instantiierungsverhalten der beiden Faktoren X und Y :

$$h(X, Y) = XY + X\bar{Y} + \bar{X}Y + \bar{X}\bar{Y}, \quad (X, Y) \in \mathcal{B}^2 \quad (4.8)$$

4.2.2. Die Booleschen Gesetze als Schlussregeln des konfiguralen kausalen Modellierens

Mit der Auffassung des konfiguralen kausalen Modellierens als Anwendung der Konfigurationsalgebra, können nun mittels der Booleschen Gesetze Boolesche Ausdrücke wie $XY + X\bar{Y} + \bar{X}Y + \bar{X}\bar{Y}$ der Instanzfunktion h aus (4.8) analysiert und auf äquivalente Ausdrücke umgeformt werden. Auf diese Weise lässt sich ein Boolescher Ausdruck derart umformen, dass die resultierende Darstellung jene Informationen in den Vordergrund rückt, die relativ zur Fragestellung von besonderem Interesse sind. Die Instanzfunktion h aus (4.8) lässt sich etwa folgendermassen umformen:

$$h(X, Y) = XY + X\bar{Y} + \bar{X}Y + \bar{X}\bar{Y} \quad (4.9)$$

$$= X(Y + \bar{Y}) + \bar{X}(Y + \bar{Y}) \quad (4.10)$$

$$= X + \bar{X} \quad (4.11)$$

$$= 1 \quad (4.12)$$

Hinsichtlich der Anwendung der Konfigurationsalgebra im Kontext des konfiguralen kausalen Modellierens kann der aus der Konstanten 1 bestehende Boolesche Ausdruck (4.12) so interpretiert werden, dass die Verhaltensregel, die das Instantiierungsverhalten von X und Y beschreibt, immer und unabhängig vom Verhalten von X und Y erfüllt ist. Indem es keine Rolle spielt wie sich die beiden Faktoren verhalten und sie demnach frei und beliebig variieren können, besteht auch keine Abhängigkeit zwischen ihnen. Insgesamt wurde damit die Instanzfunktion h ausgehend von der KDNF in eine Form gebracht, die im vorliegenden Kontext die wesentlichen Informationen direkt zum Ausdruck bringt.

Aus rein mathematischer Sicht lässt sich mittels der Booleschen Gesetze ein Boolescher Ausdruck in einen anderen äquivalenten Booleschen Ausdruck umformen. Im Anwendungskontext des konfiguralen kausalen Modellierens, bei dem Instanzfunktionen das Instantiierungsverhalten von Faktoren wiedergeben, können diese Gesetze nun gedeutet werden als Regeln des konfiguralen kausalen Schlussfolgerns.

Der Anstoss für die Rechtfertigung einer solchen Deutung ergibt sich bereits aus den drei grundlegenden Verknüpfungsformen der Konjunktion, Disjunktion und dem Komplement, die einerseits den Booleschen Gesetzen genügen und die andererseits durch die Verknüpfungsformen innerhalb von komplexen Kausalstrukturen sowie dem damit einhergehenden Instantiierungsverhalten von Faktoren motiviert wurden. Im Folgenden soll die Legitimation dieser Deutung der Booleschen Gesetze als Schlussregeln des konfiguralen kausalen Modellierens nochmals anhand der Betrachtung der einzelnen Gesetze auf Seite 140 der vorliegenden Arbeit plausibel gemacht werden.

Für die Assoziativ- und Kommutativgesetze wird erneut die Instanzfunktion $f(A, B, C) = ABC + A\bar{B}\bar{C} + \bar{A}BC + \bar{A}\bar{B}C$ aus (4.6) der Konfigurationen $\mathcal{K}_{4,8}$ aus Tabelle 4.9 betrachtet. Gibt diese Instanzfunktion f in KDNF eine Verhaltensregel wieder, bringt etwa die Instanzfunktion $f'(A, B, C) = CBA + C\bar{B}\bar{A} + \bar{C}A\bar{B} + \bar{C}\bar{A}B$ genau dieselbe Verhaltensregel zum Ausdruck. Die Instanzfunktionen f und f' sind äquivalent, sodass sie die gleiche Aussage über das Instantiierungsverhalten machen. Es kann allgemein gesagt werden, dass die Reihenfolge der Literale innerhalb einer Konjunktion oder von Konjunktionstermen innerhalb einer Disjunktion keine Rolle spielt und lediglich entscheidend ist, welche Literale in welchen Verknüpfungen enthalten sind und in welchen nicht. Das ist ein wesentliches Merkmal der konfiguralen Methoden, deren Fokus zur Analyse von kausalen Abhängigkeiten auf dem sich in geeigneter raumzeitlicher Nähe zeigenden An- und Abwesenheitsverhalten von Faktoren liegt. Zusammenfassend ergibt sich daraus die Gültigkeit der Assoziativ- und Kommutativgesetze.

Für die Betrachtung der Absorptionsgesetze wird davon ausgegangen, dass auf ähnliche Weise wie im obigen Fall von f eine Konfigurationstabelle über $\mathbf{H}$, $A, B \in \mathbf{H}$, in eine Instanzfunktion g in DNF transformiert worden ist, die unter anderem den Ausdruck $A + AB$ beinhaltet. Demnach ist die Verhaltensregel, die das Instantiierungsverhalten der Faktoren aus $\mathbf{H}$ beschreibt, unter anderem dann erfüllt, wenn (1) A gegeben ist oder wenn (2) A zusammen mit B gegeben ist. Immer wenn (2) erfüllt ist, ist auch (1) erfüllt und (2) ist als Spezialfall von (1) bereits in (1) enthalten. Sofern also (1) als Bedingung in der Verhaltensregel enthalten ist, beschreibt (2) eine überflüssige Information, auf die verzichtet und die aus dem Ausdruck entfernt werden kann. Dies entspricht genau dem Absorptionsgesetz $A + AB = A$. Auf analoge Weise kann für den Fall $A(A + B) = A$ argumentiert werden. Dieser Ausdruck besagt, dass die Verhaltensregel dann erfüllt ist, wenn (1) A gegeben ist und wenn (2) A oder B gegeben sind. Aufgrund der konjunktiven Bedingung, nach der sowohl (1) als auch (2) erfüllt sein müssen, muss gemäss (1) A gegeben sein. Ist A gegeben, folgt daraus die Gültigkeit von (2), womit die Verhaltensregel auf die Bedingung (1) reduziert werden kann.

Die Gültigkeit der Distributivgesetze innerhalb der Analyse von Konfigurationstabellen kann mit dem Verweis auf die Interpretation der konjunktiven und disjunktiven Verknüpfung direkt abgeleitet werden: Ist eine Verhaltensregel dann erfüllt, wenn ein Faktor A oder ein Faktor B jeweils mit einem Faktor C gegeben ist, ist dies gleichbedeutend mit den beiden Teilaussagen, wonach A zusammen mit C gegeben ist oder B zusammen mit C gegeben ist – kurz: $(A + B)C = AC + BC$. Gleiches gilt für $(A + B)C = AC + BC$.

Für die neutralen Elemente werden zwei Literale A und B konjunktiv verknüpft und der Faktor B als konstant anwesend beziehungsweise abwesend betrachtet, das heisst, das Literal B kann mit dem Wert 1 beziehungsweise 0 gleichgesetzt werden. Ist eine Verhaltensregel dann erfüllt, wenn A zusammen mit B gegeben ist, wobei B konstant instantiiert ist, ist die Erfüllung dieser Bedingung nur von A abhängig und AB ist genau dann wahr, wenn A wahr ist. Ist umgekehrt der Faktor B konstant abwesend, kann die Bedingung AB nie erfüllt sein. Insgesamt gilt $A \cdot 1 = A$ und $A \cdot 0 = 0$. Ähnliche Überlegungen kann man sich bei der disjunktiven Verknüpfung von A und B machen, die als Bedingung Teil einer Verhaltensregel ist. Ist der Faktor B konstant abwesend, kann die Bedingung nur dann erfüllt sein, wenn der Faktor A gegeben ist. Das heisst, die Erfüllung der Bedingung ist alleine von A abhängig, indem die Disjunktion genau dann wahr ist, wenn A wahr ist. Ist der Faktor B jedoch umgekehrt immer gegeben, so ist immer mindestens einer der beiden Faktoren anwesend und damit auch die Bedingung $A + B$ immer erfüllt. Insgesamt gelten also $A + 0 = A$ und $A + 1 = 1$.

Abschliessend bleibt noch das Komplement. Bei crisp-set Konfigurationstabellen werden Ereignisse (oder Zustände) einem Faktor zugeteilt oder nicht zugeteilt. Es ist somit nicht möglich eine

Verhaltensregel zu erfüllen, gemäss der ein Faktor gleichzeitig sowohl gegeben als auch nicht gegeben sein soll. Jedoch ist einer der beiden Zustände zwingend und es ist immer der Fall, dass ein und derselbe Faktor gegeben oder nicht gegeben ist. Insgesamt gilt also $A \cdot \bar{A} = 0$ und $A + \bar{A} = 1$.

4.2.3. Minimale Theorien als spezielle Darstellungsform einer Instanzfunktion

Ähnlich wie die obige Instanzfunktion h lassen sich auch f und g mittels der Booleschen Gesetze umformen. Die Instanzfunktion f aus (4.6) lässt sich etwa folgendermassen umformen, wobei bei Schritt (1) das sogenannte *De Morgansche Gesetz* $\overline{AB} = \overline{A + B}$ angewendet wird, das sich aus den Booleschen Gesetzen ergibt und bei Schritt (2) der oben definierte Äquivalenzoperator verwendet wird⁸:

$$f(A, B, C) = ABC + A\bar{B}C + \bar{A}BC + \bar{A}\bar{B}\bar{C} \quad (4.13)$$

$$= ABC + A\bar{B}C + ABC + \bar{A}BC + \bar{A}\bar{B}\bar{C} \quad (4.14)$$

$$= AC(B + \bar{B}) + BC(A + \bar{A}) + \bar{A}\bar{B}\bar{C} \quad (4.15)$$

$$= AC + BC + \bar{A}\bar{B}\bar{C} \quad (4.16)$$

$$\stackrel{(1)}{=} (A + B)C + \overline{(A + B)}\bar{C} \quad (4.17)$$

$$\stackrel{(2)}{=} A + B \leftrightarrow C \quad (4.18)$$

Die Instanzfunktion g lässt sich unter anderem folgendermassen umformen⁹:

$$g(A, B, C, D, E) = ABCDE + ABC\bar{D}\bar{E} + \bar{A}\bar{B}CDE + \bar{A}\bar{B}\bar{C}\bar{D}\bar{E} \quad (4.19)$$

$$+ \bar{A}BCDE + \bar{A}BC\bar{D}\bar{E} + \bar{A}\bar{B}\bar{C}DE + \bar{A}\bar{B}\bar{C}\bar{D}\bar{E}$$

$$= ACE + ACDE + BCE + BCDE + \bar{A}\bar{B}\bar{C}DE + \bar{A}\bar{B}\bar{C}\bar{D}\bar{E} \quad (4.20)$$

$$= (AC + BC + \bar{A}\bar{B}\bar{C}) \cdot (CE + DE + \bar{C}\bar{D}\bar{E}) \quad (4.21)$$

$$= ((A + B)C + \overline{(A + B)}\bar{C}) \cdot ((C + D)E + \overline{(C + D)}\bar{E}) \quad (4.22)$$

$$= (A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E) \quad (4.23)$$

⁸Aufgrund der geltenden Dualität gibt es zwei De Morgansche Gesetze. Neben dem obigen gilt das zu diesem Gesetz duale Gesetz $\overline{(\bar{A}\bar{B})} = \bar{A} + \bar{B}$. Analog zum Nachweis des Distributivgesetzes aus Tabelle 4.2 und allen weiteren Booleschen Gesetzen lässt sich die Gültigkeit der De Morganschen Gesetze u. a. durch eine Auswertung über alle Belegungen nachweisen.

⁹Um die einzelnen Rechenschritte bei der Umformung von g besser nachvollziehen zu können, empfiehlt es sich in entgegengesetzter Richtung zu rechnen.

Eine weitere Umformung von g ist die folgende:

$$g(A, B, C, D, E) = ABCDE + ABC\bar{D}E + A\bar{B}CDE + A\bar{B}C\bar{D}E + \bar{A}BCDE + \bar{A}BC\bar{D}E + \bar{A}\bar{B}CDE + \bar{A}\bar{B}C\bar{D}E \quad (4.24)$$

$$= ACE + ABCE + ACDE + ABCE + BCE + BCDE + \bar{A}\bar{B}CDE + \bar{A}\bar{B}C\bar{D}E \quad (4.25)$$

$$= (AC + BC + \bar{A}\bar{B}C) \cdot (AE + BE + DE + \bar{A}\bar{B}\bar{D}\bar{E}) \quad (4.26)$$

$$= ((A + B)C + \overline{(A + B)}\bar{C}) \cdot ((A + B + D)E + \overline{(A + B + D)}\bar{E}) \quad (4.27)$$

$$= (A + B \leftrightarrow C) \cdot (A + B + D \leftrightarrow E) \quad (4.28)$$

Hinsichtlich der Anwendung der Konfigurationsalgebra im Kontext des konfiguralen kausalen Modellierens besagt der Ausdruck (4.18), dass C genau dann gegeben ist, wenn A oder B gegeben sind. Mit (4.23) wird ausgedrückt, dass C genau dann gegeben ist, wenn A oder B gegeben sind, sowie E genau dann gegeben ist, wenn C oder D gegeben sind. Der Ausdruck (4.28) schliesslich besagt, dass C genau dann gegeben ist, wenn A oder B gegeben sind, sowie E genau dann gegeben ist, wenn A , B oder D gegeben sind.

Es dürfte dem Leser aufgefallen sein, dass alle drei Booleschen Ausdrücke (4.18), (4.23) und (4.28) genau die Form von Minimalen Theorien aufweisen, die aus den jeweiligen Konfigurationstabellen gefolgert werden können (vgl. u. a. S. 70 f. für (4.23) und (4.28)). Somit wurden die Instanzfunktionen in KDNF, die als noch wenig aussagekräftige Verhaltensregeln zwischen den Faktoren gedeutet werden können, auf die im vorliegenden Kontext wesentlichen Teile reduziert, sodass sich daraus unter Berücksichtigung des kausaltheoretischen Rahmens Kausalmodelle ablesen lassen. Eine Minimale Theorie kann demnach gedeutet werden als eine Instanzfunktion dargestellt als Boolescher Ausdruck, der bezüglich der Form gewissen Bedingungen genügt.

Sowie sich eine Instanzfunktion einer Konfigurationstabelle beispielsweise in DNF darstellen lässt, lässt sie sich je nachdem auch in eine Form bringen, aus der direkt Minimale Theorien abgelesen werden können. Von einer derart präsentierten Instanzfunktion wird künftig gesagt, dass sie in *MT-Form* oder kurz in *MTF* dargestellt ist. So sind etwa die Instanzfunktionen $f(A, B, C) = A + B \leftrightarrow C$ oder $g(A, B, C, D, E) = (A + B \leftrightarrow C) \cdot (C + D \leftrightarrow E)$ der obigen beiden Beispiele in MTF.

Im Folgenden wird eine allgemeine mathematische Definition der MT-Form einer Instanzfunktion f vorgestellt. Dazu werden zunächst speziell ausgezeichnete Instanzfunktionen definiert, über welche die MT-Form gebildet wird.

Innere Instanzfunktion f_{X_i} von f mit abhängiger Variable X_i

Was eine Minimale Theorie ist, ist innerhalb des kausaltheoretischen Rahmens mit der Definition 2.8 eindeutig beschrieben. Konkret handelt es sich jeweils um einen Ausdruck ψ der Form

$$\psi = (\phi_{X_1} \leftrightarrow X_1) \cdot \dots \cdot (\phi_{X_k} \leftrightarrow X_k), \quad (4.29)$$

wobei jeder Term $(\phi_{X_i} \leftrightarrow X_i)$, $1 \leq i \leq k$, ein RDN-Bikonditional ist. Bei jedem in (4.29) enthaltenen Ausdruck ϕ_{X_i} handelt es sich um eine disjunktive Normalform, die aus minimal hinreichenden Bedingungen besteht, die disjunktiv verknüpft minimal notwendig für X_i sind. Das Literal X_i selbst ist nicht in ϕ_{X_i} enthalten. Ferner gilt für Literale X_i und X_j , von denen RDN-Bikonditionale $(\phi_{X_i} \leftrightarrow X_i)$ beziehungsweise $(\phi_{X_j} \leftrightarrow X_j)$ in ψ enthalten sind, dass die entsprechenden Faktoren X_i und X_j paarweise verschieden sind. Schliesslich ist der gesamte Ausdruck ψ strukturell minimal, sodass ein Ausdruck ψ' , der bis auf mindestens ein fehlendes Konjunkt $(\phi_{X_i} \leftrightarrow X_i)$ identisch mit ψ ist, nicht äquivalent zu ψ ist.

Eine solche Minimale Theorie ψ kann man nun allgemein mathematisch mithilfe der Begriffe aus Abschnitt 4.1.2 als MT-Form einer Instanzfunktion f ausdrücken. Um dies zu bewerkstelligen, wird eine Instanzfunktion f von n Variablen, die durch einen Booleschen Ausdruck der Form ψ aus (4.29) repräsentiert werden kann, in Teilfunktionen zerlegt, die ihrerseits Instanzfunktionen sind. Dabei wird jeder Boolesche Ausdruck ϕ_{X_i} aus ψ als Repräsentant einer solchen, als innere Instanzfunktion f_{X_i} von f bezeichneten Teilfunktion betrachtet, der mit X_i durch die Äquivalenzverknüpfung verbunden ist. Zumal X_i selbst nicht in ϕ_{X_i} enthalten ist, handelt es sich bei f_{X_i} um eine Instanzfunktion von $n - 1$ Variablen. Gekennzeichnet wird eine innere Instanzfunktion jeweils durch einen Index, der jener Variable entspricht, deren Literal mit dem Booleschen Ausdruck ϕ_{X_i} , der die innere Instanzfunktion repräsentiert, über die Äquivalenz verknüpft ist.

Basierend auf diesen Überlegungen wird etwa die Instanzfunktion $f(A, B, C) = A + B \leftrightarrow C$ verstanden als Komposition $f(f_C(A, B), C)$, wobei der die innere Instanzfunktion $f_C(A, B)$ repräsentierende Boolesche Ausdruck $A + B$ mit C über die Äquivalenz miteinander verknüpft sind:

$$f(A, B, C) = \underbrace{A + B}_{f_C} \leftrightarrow C \quad (4.30)$$

Die Instanzfunktion g weist zwei innere Instanzfunktionen g_C und g_E auf. Dabei lässt sich erstere durch den Booleschen Ausdruck $A + B$ und letztere durch die Booleschen Ausdrücke $C + D$ oder durch $A + B + D$ darstellen:

$$g(A, B, C, D, E) = \underbrace{(A + B \leftrightarrow C)}_{g_C} \cdot \underbrace{(C + D \leftrightarrow E)}_{g_E} \quad (4.31)$$

A	B	$C = f_C(A, B)$
1	1	1
1	0	1
0	1	1
0	0	0

Tabelle 4.12.: Wertetabelle der inneren Instanzfunktion f_C der Instanzfunktion f mit Wertetabelle 4.9.

oder

$$g(A, B, C, D, E) = \underbrace{(A + B \leftrightarrow C)}_{g_C} \cdot \underbrace{(A + B + D \leftrightarrow E)}_{g_E} \quad (4.32)$$

Lässt sich eine Instanzfunktion f von n Variablen und dem Instantiierungsbereich $\mathcal{I}(f)$ durch einen Ausdruck der Form ψ aus 4.29 darstellen, gilt für jeden Teilausdruck $(\phi_{X_i} \leftrightarrow X_i)$, $1 \leq i \leq k \leq n$, dass er für alle Belegungen $(X_1, \dots, X_n) \in \mathcal{I}(f)$ den Wert 1 annimmt. Aufgrund der Äquivalenzverknüpfung gilt damit für die entsprechende innere Instanzfunktion f_{X_i} , dass sie für eine beliebige Belegung $(X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n)$ genau dann den Funktionswert 1 beziehungsweise 0 annimmt, wenn $(X_1, \dots, X_{i-1}, 1, X_{i+1}, \dots, X_n)$ beziehungsweise $(X_1, \dots, X_{i-1}, 0, X_{i+1}, \dots, X_n)$ in $\mathcal{I}(f)$ enthalten ist. Ist weder $(X_1, \dots, X_{i-1}, 1, X_{i+1}, \dots, X_n)$ noch $(X_1, \dots, X_{i-1}, 0, X_{i+1}, \dots, X_n)$ in $\mathcal{I}(f)$ enthalten, ist f_{X_i} nicht definiert für die entsprechende Belegung $(X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n)$, zumal sonst die Äquivalenz verletzt wäre. Solche Belegungen sind Teil von $\mathcal{ND}(f_{X_i})$.

Eine innere Instanzfunktion f_{X_i} kann direkt über die Wertetabelle von f dargestellt werden, indem für alle Belegungen des Instantiierungsbereiches von f die X_i -Spalte mit dem Funktionswert von f_{X_i} gleichgesetzt wird. Die innere Instanzfunktion f_C aus (4.30) lässt sich demnach durch die Wertetabelle 4.12 darstellen.

Die inneren Instanzfunktionen g_C und g_E lassen sich durch die beiden Wertetabellen aus Tabelle 4.13 darstellen.

Ist f eine Instanzfunktion von n Variablen $X_1, \dots, X_n$, gibt eine innere Instanzfunktion f_{X_i} die Beziehung zwischen einer abhängigen Variable X_i und den unabhängigen Variablen $X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n$ wieder. Das heisst, f_{X_i} beschreibt die Belegungen von X_i anhand der Variablen $X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n$. Dabei bleibt zu erwähnen, dass sich nicht jede Variable als abhängige Variable einer solchen Beziehung ansehen lässt. Wie sich etwa bei der Instanzfunktion f aus Tabelle 4.9 zeigt, lässt sich keine innere Instanzfunktion f_A von f mit abhängiger Variable A bilden, zumal bei dieser Funktion unter anderem für die Belegung 1 von B und C sowohl der Funktionswert 1 als auch 0 zugewiesen wird. Da diese Funktion das Zweiwertigkeitsprinzip verletzt, handelt es sich um keine Instanzfunktion. Wie für jede Instanzfunktion gilt auch für eine beliebige innere Instanzfunktion f_{X_i} $\mathcal{I}(f_{X_i}) \cap \mathcal{ND}(f_{X_i}) = \emptyset$.

A	B	D	E	$C = g_C(A, B, D, E)$	A	B	C	D	$E = g_E(A, B, C, D)$
1	1	1	1	1	1	1	1	1	1
1	1	0	1	1	1	1	1	0	1
1	0	1	1	1	1	0	1	1	1
1	0	0	1	1	1	0	1	0	1
0	1	1	1	1	0	1	1	1	1
0	1	0	1	1	0	1	1	0	1
0	0	1	1	0	0	0	0	1	1
0	0	0	0	0	0	0	0	0	0

(a)
(b)

Tabelle 4.13.: Wertetabelle der inneren Instanzfunktionen g_C und g_E der Instanzfunktion g .

Allgemein lässt sich eine innere Instanzfunktion f_{X_i} einer Instanzfunktion f mit abhängiger Variable X_i folgendermassen definieren:

Definition 4.14 (*Innere Instanzfunktion f_{X_i} von f mit abhängiger Variable X_i*)

Sei $\mathcal{B} = \{0, 1\}$ und f eine Instanzfunktion $f : \mathcal{B}^n \rightarrow \mathcal{B}$ von n Variablen. Die Instanzfunktion $f_{X_i} : \mathcal{D} \rightarrow \mathcal{B}$ über $\mathcal{D} \subseteq \mathcal{B}^{n-1}$, für die gilt

$$f_{X_i}(X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n) = \begin{cases} 1, & (X_1, \dots, X_{i-1}, 1, X_{i+1}, \dots, X_n) \in \mathcal{I}(f) \\ 0, & (X_1, \dots, X_{i-1}, 0, X_{i+1}, \dots, X_n) \in \mathcal{I}(f) \end{cases} \quad (4.33)$$

heisst innere Instanzfunktion f_{X_i} von f mit abhängiger Variable X_i .

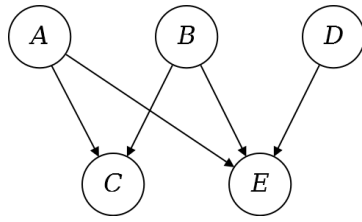
Aus dieser Definition ist unter anderem leicht zu erkennen, dass wenn ϕ_{X_i} ein Boolescher Ausdruck ist, der f_{X_i} repräsentiert, der Boolesche Ausdruck $\phi_{X_i} \leftrightarrow X_i$ für die Belegungen im Instanzierungsbereich $\mathcal{I}(f)$ von f den Wert 1 annimmt. Weiter folgt aus der Definition unter anderem, dass genau dann $(X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n) \in \mathcal{D} \subseteq \mathcal{B}^{n-1}$, wenn nicht sowohl $(X_1, \dots, X_{i-1}, 1, X_{i+1}, \dots, X_n) \in \mathcal{NI}(f)$ als auch $(X_1, \dots, X_{i-1}, 0, X_{i+1}, \dots, X_n) \in \mathcal{NI}(f)$ gilt. Schliesslich gilt es zu erwähnen, dass eine innere Instanzfunktion f_{X_i} von f nur für eine Instanzfunktion f gebildet werden kann, die total ist. Denn ist f partiell, lässt sie sich nicht mehr eindeutig durch die Komposition innerer Instanzfunktionen beschreiben. Wäre $f : \mathcal{D} \rightarrow \mathcal{B}$ etwa eine partielle Instanzfunktion über $\mathcal{D} \subseteq \mathcal{B}^3$, für die beispielsweise $(0, 0, 0) \in \mathcal{ND}(f)$ gilt, würde für jede innere Instanzfunktion f_{X_i} von f gelten, dass $(0, 0) \in \mathcal{ND}(f_{X_i})$. Konkret von f_{X_1} ausgehend, wäre daraus jedoch nicht bestimmbar, ob (i) $(0, 0, 0) \in \mathcal{ND}(f)$ gilt oder (ii) $(0, 0, 0) \in \mathcal{NI}(f)$ und $(1, 0, 0) \in \mathcal{NI}(f)$ gilt. Analoges würde für allfällige innere Instanzfunktionen f_{X_2} oder f_{X_3} gelten, womit eine Komposition von inneren Instanzfunktionen die entsprechende (partielle) Instanzfunktion f nicht eindeutig beschreiben könnte. Wäre im Gegensatz dazu f total, würde aus $(0, 0) \in \mathcal{ND}(f_{X_1})$ eindeutig (ii) folgen.

Unter Berücksichtigung der Interpretation der Konfigurationsalgebra im Rahmen des konfigurationsalen kausalen Modellierens, beschreibt eine innere Instanzfunktion f_{X_i} gemäss Definition 4.14 das Instantiierungsverhalten eines Faktors X_i in Abhängigkeit vom Instantiierungsverhalten der von X_i paarweise verschiedenen Faktoren $X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n$. So unterscheidet sich beispielsweise die obige innere Instanzfunktion $f_C(A, B) = A + B$ insofern von einer üblichen Instanzfunktion $f(A, B) = A + B$ ohne Index, als erstere nicht nur etwas über die Verknüpfung von A und B aussagt, sondern darüber hinaus auch etwas über C vermittelt. Indem bei f_C die Funktionswerte mit dem Instantiierungsverhalten von C korrespondieren, kann f_C nicht nur als gültige Verhaltensregel zwischen A und B verstanden werden, sondern zusätzlich als Beziehung zwischen dem im Index stehenden abhängigen Faktor C und den unabhängigen Faktoren A und B . Konkret bringt die innere Instanzfunktion $f_C(A, B) = A + B$ zum Ausdruck, dass (1) A oder B instantiiert sein können, und (2) dass genau dann, wenn A oder B instantiiert sind, C instantiiert ist. So kommt es denn auch, dass eine innere Instanzfunktion f_{X_i} jeweils nur über $\mathcal{I}(f)$ definiert ist und damit etwa g_E mit der Wertetabelle 4.13b eine partielle Instanzfunktion darstellt. Die innere Instanzfunktion g_E ist beispielsweise nicht definiert für die Belegung $(1, 0, 0, 0)$, das heisst für die Belegung, bei der A den Wert 1 und B, C sowie D den Wert 0 annehmen. Wäre g_E für diese Belegung definiert, würde dies implizieren, dass es Fälle gibt, bei denen der Faktor A , nicht aber B, C und D instantiiert sind. Belegungen, bei denen sich die vier Faktoren derart verhalten und die allesamt in $\mathcal{NI}(g)$ enthalten sind, sind angesichts der tatsächlichen Kausalstruktur nie instantiiert, da bei instantiiertem Faktor A unter anderem auch jeweils C instantiiert ist.

Redundanzfreie disjunktive Normalform (RDNF)

Ausgehend von einer Instanzfunktion f , die das Auftrittsverhalten von Faktoren einer Konfigurationstabelle $\mathcal{K}_G$ über $\mathbf{G} = \{X_1, \dots, X_n\}$ ausdrückt, gibt eine innere Instanzfunktion f_{X_i} in DNF eine aus hinreichenden Bedingungen bestehende notwendige Bedingung von X_i wieder. Enthält die DNF von f_{X_i} zusätzlich keine redundanten Teile, wie beispielsweise der Boolesche Ausdruck $A + B$ von f_C , gibt die innere Instanzfunktion eine aus minimal hinreichenden Bedingungen bestehende minimal notwendige Bedingung von X_i wieder – kurz: die DNF von f_{X_i} stellt eine RDN-Bedingung von X_i dar.

Eine solche RDN-Bedingung von X_i lässt sich nun leicht mittels Verwendung der Begriffe aus Abschnitt 4.1.2 über die innere Instanzfunktion f_{X_i} definieren und so in das hier präsentierte mathematische Fundament einbetten. Im Zusammenhang mit Instanzfunktionen wird dabei wie bisher jeweils von einer *Form* und nicht von einer Bedingung gesprochen, sodass f_{X_i} in RDN-Form sich interpretieren lässt als RDN-Bedingung von X_i .



$A :=$ „Entzünden eines Feuers“,
 $B :=$ „Anstellen einer Glühlampe“,
 $C :=$ „Helligkeit“,
 $D :=$ „Aneinanderreiben der Handflächen“,
 $E :=$ „Wärme“.

Abbildung 4.14.: Common-Cause-Struktur, die Alternativursachen von Helligkeit und Wärme wiedergibt.

Im ersten Moment könnte man vielleicht vermuten, dass eine RDN-Form dasselbe ist wie die in Abschnitt 4.1.2 eingeführte disjunktive Minimalform (DMF) und sich deshalb dieser gängige Begriff der DMF als Definition einer RDN-Form anbietet. Wie im folgenden gezeigt wird, ist die DMF dazu jedoch nicht geeignet. Einerseits wurde der Begriff der RDN-Bedingung bereits im kausaltheoretischen Rahmen etabliert und andererseits ist der Begriff der Minimalform derart durch bestimmte, im Kontext des konfiguralen kausalen Modellierens unpassende Eigenschaften vorbelastet, dass sich leicht Missverständnisse ergeben können. Konkret ist dem gängigen Begriff der Minimalform zufolge eine disjunktive Normalform genau dann minimal, wenn jede äquivalente Darstellung mindestens die gleiche Anzahl Disjunkte aufweist und, falls diese Anzahl identisch ist, die Anzahl Literale pro Disjunkt mindestens gleich gross ist. Im Fall der inneren Instanzfunktion g_E , die durch $\phi_{E,1} = A + B + D$ oder durch $\phi_{E,2} = C + D$ dargestellt werden kann, würde man unter Verwendung dieser allgemeinen Definition nur $\phi_{E,2}$ als Minimalform betrachten. Bei vielen Fragestellungen anderer Anwendungsbereiche der Booleschen Algebra mag ein derartiger Begriff der disjunktiven Minimalform zwar angemessen sein, aber im hier vorliegenden Kontext greifen solche rein quantitativen Eigenschaften zu kurz.¹⁰ Kausalstrukturen lassen sich allgemein nicht durch diese disjunktive Minimalform einfangen. Abbildung 4.14 etwa zeigt ein einfaches Beispiel einer Kausalstruktur zusammen mit den entsprechenden Faktordefinitionen, bei der obiger Boolescher Ausdruck $\phi_{E,1}$ die Ursachen von E korrekt abbildet. Das Beispiel gibt die kausalen Zusammenhänge wieder, wonach das Entzünden eines Feuers (A) und das Anstellen einer Glühlampe (B) zwei Alternativursachen einerseits für Helligkeit (C) und andererseits für Wärme (E) sind. Das Aneinanderreiben der Handflächen (D) ist eine weitere Alternativursache für Wärme (E).

¹⁰Ein Anwendungsgebiet, für das der gängige Begriff der disjunktiven Minimalform angemessen ist, ist etwa die Digitaltechnik. Mit der Schaltalgebra werden in dieser Disziplin Boolesche Funktionen als Schaltungen gedeutet und mithilfe der Booleschen Gesetze in Minimalform gebracht, um so kostengünstige und platzsparende Hardware herstellen zu können (vgl. u. a. Nelson et al. (1995)).

Eine dazugehörige homogene und maximal diverse Konfigurationstabelle über $\mathbf{F} = \{A, B, C, D, E\}$ ist identisch mit $\mathcal{K}_{4,10}$, wovon bekanntlich die Booleschen Ausdrücke (4.31) und (4.32) Minimale Theorien darstellen. In diesem Fall fängt der Boolesche Ausdruck $\phi_{E,1} = A + B + D$, der keine disjunktive Minimalform darstellt, die Ursachen von E ein und nicht die ebenfalls aus $\mathcal{K}_{4,10}$ gefolgerte disjunktive Minimalform $\phi_{E,2} = C + D$. Dieses Beispiel verdeutlicht, dass bei Booleschen Ausdrücken, die Kausalstrukturen einfangen, nicht entscheidend ist, ob sie eine möglichst kleine Anzahl Literale aufweisen, sondern ob die Ausdrücke redundanzfrei sind. Dabei ist eine solche redundanzfreie disjunktive Normalform offensichtlich nicht gleichzusetzen mit der gängigen disjunktiven Minimalform.¹¹

Indem die Funktionswerte von f_{X_i} mit X_i korrespondieren, entspricht im Kontext des konfigurationsalen kausalen Modellierens ein Implikant von f_{X_i} einer hinreichenden Bedingung von X_i . Ist ein solcher Implikant zusätzlich ein Primimplikant von f_{X_i} , stellt er eine minimal hinreichende Bedingung von X_i dar (vgl. Def. 4.7 und Def. 4.8). Weiter entspricht eine aus disjunktiv verknüpften Implikanten bestehende Abdeckung von f_{X_i} genau einer Disjunktion von hinreichenden Bedingungen von X_i , die als Ganzes notwendig für X_i ist. Ist die Abdeckung von f_{X_i} ferner irredundant, liegt eine minimal notwendige Bedingung von X_i vor (vgl. Def. 4.10). Insgesamt lässt sich demnach eine RDN-Bedingung von X_i beschreiben als eine aus Primimplikanten bestehende irredundante Abdeckung von f_{X_i} , wobei f_{X_i} total (vgl. bspw. f_C aus Tabelle 4.12) oder partiell (vgl. bspw. g_C bzw. g_E aus Tabelle 4.13) sein kann. Allgemein ergibt sich für die redundanzfreie disjunktive Normalform folgende Definition¹²:

Definition 4.15 (Redundanzfreie disjunktive Normalform (RDNF))

Sei $\mathcal{B} = \{0, 1\}$, $\mathcal{D} \subseteq \mathcal{B}^n$ und $f : \mathcal{D} \rightarrow \mathcal{B}$ eine Instanzfunktion. Die Instanzfunktion f ist genau dann in redundanzfreier disjunktiver Normalform (RDNF) dargestellt, wenn es sich bei dieser Darstellung um eine disjunktive Normalform handelt, für die gilt:

- (a) Jeder der disjunktiv verknüpften Konjunktionsterme ist ein Primimplikant von f .

¹¹Eine klare begriffliche Abgrenzung zur Vermeidung von Missverständnissen ist auch deshalb sinnvoll, weil ebenfalls der Begriff der Minimalform isoliert teils sehr unterschiedlich genutzt wird. So wird beispielsweise in der Schaltalgebra oft zusätzlich eine Kostenfunktion berücksichtigt und darauf basierend anhand der niedrigsten Kosten die Minimalform ermittelt (vgl. u. a. Keller & Paul (1997), S. 107; Staab (2007), S. 78 f.). Die Kosten werden üblicherweise mittels der Anzahl Verknüpfungszeichen berechnet, wodurch die Minimalform genau dann gegeben ist, wenn es sich um einen Booleschen Ausdruck mit minimaler Anzahl solcher Zeichen handelt. In diesem Fall ist die Minimalform nicht zwangsläufig in DNF dargestellt, da etwa $X_1X_2 + X_1X_3$ mit Kosten von 3 grösser ist als die äquivalente Darstellung $X_1(X_2 + X_3)$ mit Kosten 2.

¹²Die Definition der RDNF wird nicht auf innere Instanzfunktionen beschränkt, weil auch Instanzfunktionen in RDNF präsentiert werden können, die keine inneren Instanzfunktionen darstellen. Selbstverständlich werden aber im weiteren Verlauf der Arbeit vor allem die RDN-Formen von inneren Instanzfunktionen von besonderem Interesse sein.

- (b) Die Primimplikanten bilden disjunktiv verknüpft eine irredundante Abdeckung von f .

Zur Erinnerung soll nochmals darauf hingewiesen werden, dass es sich bei obiger Definition der RDN-Form um die rein mathematische Entsprechung jenes Ausdrucks handelt, der innerhalb des kausaltheoretischen Rahmens als RDN-Bedingung bezeichnet wird. Erstere wird also noch nicht im Kontext des konfiguralen kausalen Modellierens interpretiert, weshalb etwa das die Identität $f(X) = X$ repräsentierende Literal X alleine bereits eine RDNF darstellen kann. Eine dazu korrespondierende vermeintliche RDN-Bedingung X des Literals X ist im Kontext des konfiguralen kausalen Modellierens deshalb ausgeschlossen, weil der kausaltheoretische Rahmen sie nicht zulässt. Von demselben rein mathematischen Standpunkt aus wird auch die noch folgende Definition der MT-Form betrachtet.

MT-Form (MTF)

Mit den Definitionen 4.14 und 4.15 sind zwei wesentliche Begriffe für die Definition der MT-Form einer Instanzfunktion eingeführt. Konkret lassen sich damit RDN-Bikonditionale ausdrücken, die allenfalls einfache Minimale Theorien sind: Ist f eine Instanzfunktion und f_{X_i} eine innere Instanzfunktion von f mit abhängiger Variable X_i , entspricht der Boolesche Ausdruck $\phi_{X_i} \leftrightarrow X_i$, wobei ϕ_{X_i} eine RDNF von f_{X_i} ist, im Kontext des konfiguralen kausalen Modellierens einem RDN-Bikonditional von X_i . Um nun Minimale Theorien abbilden zu können, werden solche Booleschen Ausdrücke konjunktiv verknüpft und es wird verlangt, dass diese Konjunktion strukturell minimal ist. Ist ψ eine solche Konjunktion, ist letzteres bekanntlich genau dann der Fall, wenn kein Konjunkt wie $(\phi_{X_i} \leftrightarrow X_i)$ aus ψ entfernt werden kann, ohne dass der resultierende Boolesche Ausdruck ψ' seine Äquivalenz zu ψ verliert. Zudem wird verlangt, dass für dieselben Variable X_i nicht mehrere unterschiedliche Konjunkte $(\phi_{X_i,1} \leftrightarrow X_i), \dots, (\phi_{X_i,m} \leftrightarrow X_i)$, $m > 1$, in ψ enthalten sind. Zusammen mit der Berücksichtigung dieser Redundanzfreiheit lässt sich die MT-Form folgendermassen definieren:

Definition 4.16 (MT-Form)

Sei $\mathcal{B} = \{0, 1\}$ und $f : \mathcal{B}^n \rightarrow \mathcal{B}$ eine Instanzfunktion von n Variablen. Die Instanzfunktion f ist genau dann in MTF dargestellt, wenn es sich bei der Darstellung um einen Booleschen Ausdruck ψ der folgenden Form handelt:

$$\psi = (\phi_{X_1} \leftrightarrow X_1) \cdot \dots \cdot (\phi_{X_l} \leftrightarrow X_l), 1 \leq i \leq l \leq n, \quad (4.34)$$

wobei für ein beliebiges in ψ enthaltenes Konjunkt $(\phi_{X_j} \leftrightarrow X_j)$, $i \leq j \leq l$, gilt:

- (a) ϕ_{X_j} ist eine RDNF der inneren Instanzfunktion f_{X_j} von f mit abhängiger Variable X_j ,

- (b) wird $(\phi_{X_j} \leftrightarrow X_j)$ aus ψ entfernt, ist der resultierende Term ψ' nicht mehr äquivalent zu ψ ,
- (c) $X_j \neq X_k$ für jede beliebige Variable X_k , $i \leq k \leq l$ und $j \neq k$, für die ebenfalls ein Boolescher Ausdruck $(\phi_{X_k} \leftrightarrow X_k)$ in ψ enthalten ist.

Weil die MT-Form einer Instanzfunktion f über innere Instanzfunktionen von f definiert wird, wobei innere Instanzfunktionen von f nur für eine totale Instanzfunktion f gebildet werden, wird auch von f , die in MT-Form ist, explizit verlangt, dass sie total ist (vgl. S. 167). Ist f eine Instanzfunktion von n Variablen, die im Kontext des konfiguralen kausalen Modellierens das Auftretensverhalten von Faktoren einer Konfigurationstabelle $\mathcal{K}_G$ über $G = \{X_1, \dots, X_n\}$ ausdrückt, entspricht f in MT-Form einer Minimalen Theorie von $\mathcal{K}_G$.

4.3. In Kürze

Im vorliegenden Kapitel wurde das mathematische Fundament des von crisp-set Konfigurationstabellen ausgehenden konfiguralen kausalen Modellierens beschrieben. Das konfigurale kausale Modellieren ist eine Anwendung der Konfigurationsalgebra, die ihrerseits eine zweielementige Boolesche Algebra ist. Die Konfigurationsalgebra ist eine dem Zweiwertigkeitsprinzip genügende mathematische Struktur, die aus den Elementen 0 und 1, der Konjunktion, der Disjunktion sowie dem Komplement besteht und den Booleschen Gesetzen genügt.

Interpretiert man 1 als „ist instantiiert“ oder „ist gegeben“ und 0 als „ist nicht instantiiert“ oder „ist nicht gegeben“, lässt sich mit der Konfigurationsalgebra das Instantiierungsverhalten von Faktoren beschreiben und analysieren. Aus mathematischer Sicht handelt es sich bei diesem Instantiierungsverhalten um Instanzfunktionen, die Elemente aus $\mathcal{D} \subseteq \mathcal{B}^n$ auf $\mathcal{B}$ abbilden, wobei $\mathcal{B} = \{0, 1\}$ und $\mathcal{B}^n$ das n -fache kartesische Produkt von $\mathcal{B}$ mit sich selbst ist. Instanzfunktionen sind also Boolesche Funktionen oder partiell definierte Boolesche Funktionen.

Instanzfunktionen lassen sich durch verschiedene Darstellungstypen wie Wertetabellen oder Boolesche Ausdrücke repräsentieren. Im Kontext des konfiguralen kausalen Modellierens entspricht eine tabellarische Darstellung einer Instanzfunktion dem Verhaltensmuster, dem Faktoren folgen und ein Boolescher Ausdruck der Verhaltensregel, der die Faktoren dabei genügen. So gesehen können umgekehrt sowohl Konfigurationstabellen als auch daraus ableitbare Minimale Theorien gedeutet werden als unterschiedliche Darstellungen derselben Instanzfunktion. Konkret entspricht eine Konfigurationstabelle $\mathcal{K}_G$ über $G = \{X_1, \dots, X_n\}$ dem Instantiierungsbe-

reich $\mathcal{I}(f)$ einer Instanzfunktion f von n Variablen und die aus $\mathcal{K}_G$ ableitbaren Minimalen Theorien speziellen Darstellungen von Booleschen Ausdrücken, die f repräsentieren. Diese spezielle Darstellungsform von f mit Booleschen Ausdrücken wird als MT-Form bezeichnet und lässt sich mithilfe gängiger Begriffe der Booleschen Algebra definieren.

Ausgehend von dieser Deutung wird ein direkter mathematischer Zugang für das konfigurationale kausale Modellieren ermöglicht, indem Konfigurationstabellen unter anderem in Instanzfunktionen transformiert und mit den Booleschen Gesetzen in MT-Form gebracht werden können. Daraus lassen sich schliesslich direkt Minimale Theorien ablesen. Wie sich im Folgekapitel zeigen wird, ergibt sich damit auch die Möglichkeit von einer Vielzahl an Methoden zu profitieren, die im Zusammenhang mit anderen Anwendungen einer zweielementigen Booleschen Algebra entwickelt wurden und nun unter anderem für eine systematische Suche nach Minimalen Theorien eingesetzt werden können. Schliesslich wird sich ebenfalls zeigen, dass mit der Einführung der Konfigurationsalgebra ein Fundament vorliegt, mit dessen Verwendung sich die verschiedenen bereits existierenden Aufdeckungsverfahren direkt vergleichen lassen.

5. Das Aufdeckungsverfahren csCNA+

Die in den vorhergehenden Kapiteln beschriebene Kausaltheorie, die empirische Grundlage sowie die Konfigurationsalgebra als mathematisches Fundament des konfiguralen kausalen Modellierens liefern das Gerüst für die Entwicklung von Verfahren, mit denen innerhalb des analytischen Moments systematisch Kausalmodelle aus crisp-set Konfigurationstabellen erschlossen werden können.¹ Im vorliegenden Kapitel wird ein neues, als csCNA+ bezeichnetes Verfahren für die Kausalanalyse solcher Konfigurationstabellen eingeführt.

Für die Entwicklung von csCNA+ werden zuerst die bereits existierenden Verfahren aus QCA und CNA für den crisp-set Fall betrachtet, die in der Folge entsprechend auch als csQCA und csCNA bezeichnet werden. Sowohl csQCA als auch csCNA wurden für das konfigural-kausale Modellieren entwickelt, weisen unter anderem aber unterschiedliche Suchziele und unterschiedliche dazu eingesetzte Suchalgorithmen auf. Die grundlegenden Prinzipien, die zu diesen Unterschieden führen, werden unter Verwendung der Konfigurationsalgebra ausgearbeitet und die algorithmischen Implementierungen der Prinzipien miteinander verglichen. Teile der sich daraus ergebenden Resultate werden zentrale Bausteine von csCNA+ darstellen. Wie der Name dabei vermuten lässt, orientiert sich csCNA+ insgesamt am Vorgehen von csCNA (vgl. u. a. Baumgartner (2009)), das dasselbe Ziel verfolgt, indem versucht wird Kausalstrukturen beliebiger Komplexität aufzudecken. Angesichts der im kausaltheoretischen Rahmen eingeführten strukturellen Redundanzen und der damit einhergehenden Weiterentwicklung der Kausaltheorie (vgl. Abschnitt 2.7, S. 44 ff.), werden dabei Ergänzungen vorgenommen. Inspiriert durch csQCA werden des Weiteren bestehende Datenverarbeitungsschritte recheneffizienter gestaltet oder ersetzt. Schliesslich wird ein systematisches Vorgehen zur Folgerung von indirekten Kausalbeziehungen als formalisierter Schritt in den Algorithmus integriert.

Um csCNA+ auf seine Korrektheit prüfen und allfällige Fehlschlüsse direkt auf Fehlverhalten der verwendeten Algorithmen zurückführen zu können, wird das Verfahren unter einer idealen Datenlage entwickelt. Konkret werden homogene und maximal diverse Konfigurationstabellen betrachtet, deren Güte angemessen ist und alle Konfigurationen dieselbe Fallzahl aufweisen

¹Zum analytischen Moment vgl. Abschnitt 3.2, S. 102 ff.

(vgl. Kapitel 3). Hinsichtlich der Komplexität der untersuchten Kausalstrukturen werden keine Annahmen getroffen, womit es mit csCNA+ unter anderem auch möglich sein wird, Common-Cause-Strukturen, Kausalketten oder Feedbackstrukturen aufzudecken.

5.1. Das analytische Moment

Das Ziel des hier präsentierten Verfahrens csCNA+ ist, alleine ausgehend von einer homogenen und maximal diversen Konfigurationstabelle $\mathcal{K}_F$ über $F = \{X_1, \dots, X_n\}$ systematisch jedes Kausalmodell ω_i aufzudecken, das mit $\mathcal{K}_F$ verträglich ist. Die aus dem analytischen Moment hervorgehenden Kausalmodelle sind gemäss Definition 3.1 und dem kausaltheoretischen Rahmen kausal interpretierte Minimale Theorien oder Systeme von Minimalen Theorien von $\mathcal{K}_F$. Letzteres ist genau dann der Fall, wenn eine Minimale Theorie von $\mathcal{K}_F$ mindestens eine aus Literalen von Faktoren aus F bestehende kausale Verkettung $(X_i, \dots, X_j, \dots, X_k)$ einfängt. In diesem Fall muss für die entsprechende kausale Interpretation die Minimale Theorie aufgrund der Intransitivität zusammen mit zusätzlichen Minimalen Theorien ein System von Minimalen Theorien bilden, wobei für jede indirekte kausale Relevanz eines Literals X_i für ein anderes Literal X_k eine einfache Minimale Theorien von X_k im System enthalten sein muss, von der X_i Teil ist. Weiter sind in diesem Zusammenhang auch allfällige Switching-Literale als solche auszuweisen, sofern sich herausstellt, dass die Minimalen Theorien zwar auf eine indirekte kausale Relevanz von X_i für X_j hindeuten, X_i aber kein Difference-Maker von X_j ist.² Die Menge solcher Switching-Literale von Faktoren aus F wird künftig mit einem „s“ als F^s gekennzeichnet, wobei $F^s \subset F$ gilt. Insgesamt wird also ein mit einer Konfigurationstabelle $\mathcal{K}_F$ verträgliches Kausalmodell ω_i folgendermassen charakterisiert:

$$\omega_i := [\Psi_i; \beta_1, \dots, \beta_l; F^s], \quad (5.1)$$

wobei Ψ_i eine kausal interpretierte Minimale Theorie ist, $\beta_1, \dots, \beta_l$, $l \geq 0$, einfache Minimale Theorien sind, die zur Bestätigung indirekter Zusammenhänge aus Ψ_i herangezogen werden, und F^s die Menge von Switching-Literalen der Faktoren aus F ist. Es kann sowohl $l = 0$ als auch $F^s = \emptyset$ gelten. Zudem lässt sich ein Kausalmodell grundsätzlich alleine mittels Ψ_i und $\beta_1, \dots, \beta_l$ vollständig beschreiben. Die Hinzunahme von F^s soll jedoch die Lesbarkeit etwas fördern. Neben der obigen Darstellung in Form eines Systems von Ausdrücken, wird ein Kausalmodell üblicherweise mittels eines kausal wohlgeformten Hypergraphen dargestellt (vgl. Abschnitt 2.1.4 und Abschnitt 2.10.2), der sich direkt aus $[\Psi_i; \beta_1, \dots, \beta_l; F^s]$ ablesen lässt.

²Zum Begriff des Switching-Literals vgl. S. 75 f.

Mithilfe der Konfigurationsalgebra lässt sich ein wesentlicher Teil dieses Vorhabens nun derart mathematisch ausdrücken, dass man $\mathcal{K}_{\mathbf{F}}$ als Instantiierungsbereich $\mathcal{I}(f)$ einer Instanzfunktion f auffasst, deren MT-Formen gesucht werden. Aufgrund der angenommenen maximalen Diversität gemäss Definition 3.8 der homogenen Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ über $\mathbf{F} = \{X_1, \dots, X_n\}$, sind dabei relativ zu $\mathbf{F}$ in $\mathcal{K}_{\mathbf{F}}$ alle Kombinationen von an- und abwesenden Faktoren enthalten, die bei gegebener Homogenität gemäss Definition 3.7 realisierbar sind. Konkret bedeutet das, dass Konfigurationen, die nicht in $\mathcal{K}_{\mathbf{F}}$ enthalten sind, entweder empirisch nicht möglich sind oder nur unter inhomogenen Bedingungen zustande kommen können. Die entsprechende Instanzfunktion f , die das Instantiierungsverhalten der Faktoren aus der homogenen und maximal diversen Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ mathematisch ausdrückt, ist somit total gemäss Definition 4.11 und nimmt genau für jene Belegungen den Wert 1 an (und sonst 0), die einer Konfiguration aus $\mathcal{K}_{\mathbf{F}}$ entsprechen. Mit $\mathcal{I}(f) = \mathcal{K}_{\mathbf{F}}$ gilt für die totale Instanzfunktion f :

$$f(X_1, \dots, X_n) = \begin{cases} 1, & (X_1, \dots, X_n) \in \mathcal{I}(f) \\ 0, & (X_1, \dots, X_n) \notin \mathcal{I}(f) \end{cases} \quad (5.2)$$

Wie im vorhergehenden Kapitel beschrieben, lassen sich damit unter den gegebenen Voraussetzungen im Prinzip Minimale Theorien einer Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ finden, indem man durch Umformungen der KDNF von f mittels der Booleschen Gesetze nach den MT-Formen von f sucht. Die KDNF der totalen Instanzfunktion f ist dabei gemäss dem Hauptsatz der Booleschen Algebra eindeutig und kann über die Bestimmung der Minterme von f direkt aus $\mathcal{I}(f)$ beziehungsweise $\mathcal{K}_{\mathbf{F}}$ abgelesen werden.

Zur nochmaligen Veranschaulichung der fundamentalen Beziehung zwischen Konfigurationstabelle, Konfigurationsalgebra, Instanzfunktion und Minimalen Theorien werden die beiden Beispiele aus den Abbildungen 5.1 und 5.2 betrachtet, die in der Folge dieses Kapitels immer wieder aufgegriffen werden. Dabei gilt es folgende Randbemerkung vorausszuschicken: Indem die Konfigurationsalgebra im vorliegenden Kapitel durchwegs im konkreten Anwendungskontext des konfiguralen kausalen Modellierens betrachtet wird, wird auch im Zusammenhang mit Instanzfunktionen nun meist von Faktoren gesprochen, die natürliche Eigenschaften repräsentieren, und nicht mehr von (uninterpretierten) Variablen, wie dies im vorhergehenden Kapitel noch gehandhabt wurde.

Die beiden Abbildungen 5.1 und 5.2 zeigen die datengenerierenden Kausalstrukturen zusammen mit den dazugehörigen homogenen und maximal diversen Konfigurationstabellen $\mathcal{K}_{5.1}$ und $\mathcal{K}_{5.2}$, die untersucht werden.

Gerade beim Beispiel aus Abbildung 5.1 dürfte auffallen, dass es sich um eine relativ vertrackte Struktur handelt. Der Grund für die Wahl derartiger Beispiele liegt in der Nachvollziehbarkeit

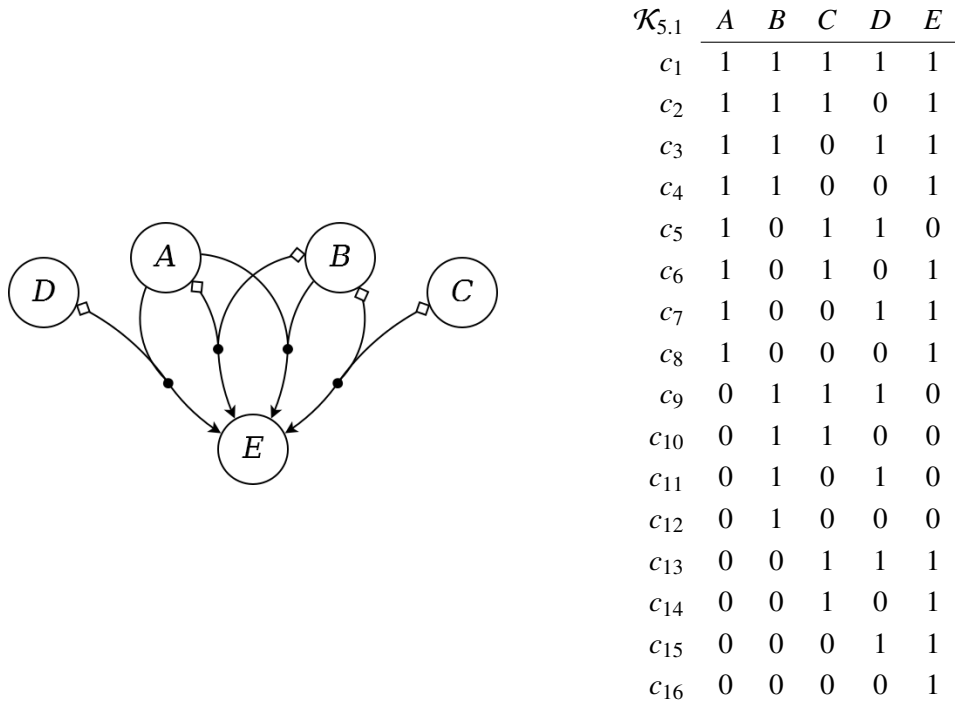


Abbildung 5.1.: Datengenerierende Kausalstruktur zusammen mit der dazugehörigen Konfigurationstabelle $\mathcal{K}_{5,1}$ über $\mathbf{G} = \{A, B, C, D, E\}$. $\mathcal{K}_{5,1}$ ist homogen und maximal divers.

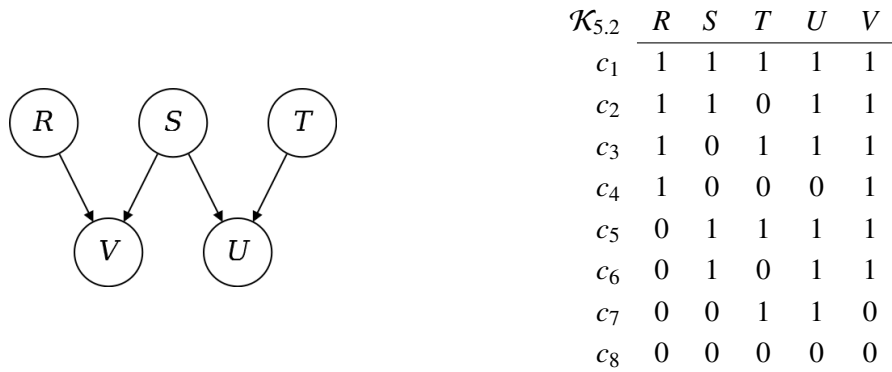
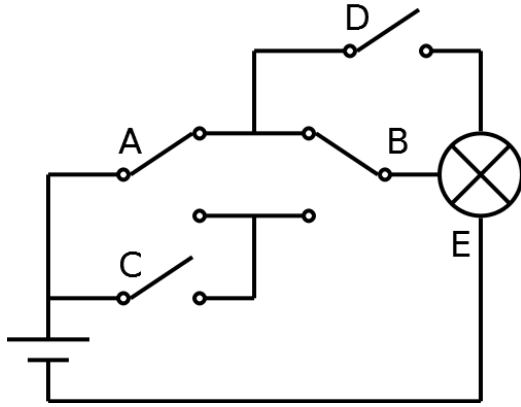


Abbildung 5.2.: Datengenerierende Kausalstruktur zusammen mit der dazugehörigen Konfigurationstabelle $\mathcal{K}_{5,2}$ über $\mathbf{H} = \{R, S, T, U, V\}$. $\mathcal{K}_{5,2}$ ist homogen und maximal divers.



- $A :=$ „Switching-Schalter A zeigt nach oben“,
 $B :=$ „Switching-Schalter B zeigt nach oben“,
 $C :=$ „Schalter C ist offen“,
 $D :=$ „Schalter D ist offen“,
 $E :=$ „Glühlampe E brennt“.

Abbildung 5.3.: Stromkreis, dem die Kausalstruktur aus Abbildung 5.1 zugrunde liegt.

der im Folgenden präsentierten Prozessschritte. Die Beispiele weisen insofern eine genügend hohe Komplexität auf, dass mit ihrer Hilfe gut veranschaulicht werden kann, wie die Algorithmen konkret arbeiten. Wählt man diesbezüglich zu einfache Kausalstrukturen, reduzieren sich die Algorithmen oftmals auf wenige auszuführende Schritte, sodass wesentliche Vorgänge nicht zur Anwendung kommen. Das wiederum erschwert die Nachvollziehbarkeit des Algorithmus in seiner vollen Ausprägung. Nichtsdestotrotz stellt sich möglicherweise gerade hinsichtlich des Beispiels aus Abbildung 5.1 die Frage, ob es sich bei diesem Hypergraphen tatsächlich um eine Kausalstruktur handeln kann. Dass dem so ist, soll der folgende Stromkreis aus Abbildung 5.3 mit den entsprechenden Faktordefinitionen verdeutlichen. Diesem Stromkreis liegen die kausalen Abhängigkeiten aus Abbildung 5.1 zugrunde.

Die beiden Konfigurationstabellen $\mathcal{K}_{5,1}$ über $\mathbf{G} = \{A, B, C, D, E\}$ beziehungsweise $\mathcal{K}_{5,2}$ über $\mathbf{H} = \{R, S, T, U, V\}$ lassen sich durch die folgenden beiden totalen Instanzfunktionen $g : \mathcal{B}^5 \rightarrow \mathcal{B}$ beziehungsweise $h : \mathcal{B}^5 \rightarrow \mathcal{B}$, $\mathcal{B} = \{0, 1\}$, in KDNF darstellen:

$$\begin{aligned}
 g(A, B, C, D, E) = & ABCDE + ABC\bar{D}E + AB\bar{C}DE + AB\bar{C}\bar{D}E + A\bar{B}CDE + A\bar{B}C\bar{D}E \\
 & + A\bar{B}\bar{C}DE + A\bar{B}\bar{C}\bar{D}E + \bar{A}BCDE + \bar{A}BC\bar{D}E + \bar{A}\bar{B}CDE + \bar{A}\bar{B}C\bar{D}E \\
 & + \bar{A}\bar{B}\bar{C}DE + \bar{A}\bar{B}\bar{C}\bar{D}E, \quad (A, B, C, D, E) \in \mathcal{B}^5
 \end{aligned} \quad (5.3)$$

$$\begin{aligned}
 h(R, S, T, U, V) = & RSTUV + R\bar{S}\bar{T}UV + R\bar{S}TUV + R\bar{S}\bar{T}\bar{U}V \\
 & + \bar{R}STUV + \bar{R}\bar{S}\bar{T}UV + \bar{R}\bar{S}TUV + \bar{R}\bar{S}\bar{T}\bar{U}V, \quad (R, S, T, U, V) \in \mathcal{B}^5
 \end{aligned} \quad (5.4)$$

Von diesen KDN-Formen ausgehend, kann (5.3) beziehungsweise (5.4) folgendermassen umgeformt werden:

$$g(A, B, C, D, E) = ABCDE + ABC\bar{D}E + AB\bar{C}DE + AB\bar{C}\bar{D}E + A\bar{B}CDE + A\bar{B}C\bar{D}E + A\bar{B}\bar{C}DE + A\bar{B}\bar{C}\bar{D}E + \bar{A}BCDE + \bar{A}BC\bar{D}E + \bar{A}\bar{B}CDE + \bar{A}\bar{B}C\bar{D}E + \bar{A}\bar{B}\bar{C}DE + \bar{A}\bar{B}\bar{C}\bar{D}E \quad (5.5)$$

$$= ABE + A\bar{D}E + \bar{B}\bar{C}E + \bar{A}\bar{B}E + \bar{A}\bar{B}\bar{E} \quad (5.6)$$

$$+ \bar{A}BC\bar{E} + \bar{A}B\bar{D}\bar{E} + \bar{A}BCD\bar{E} + \bar{A}\bar{B}CD\bar{E} \quad (5.7)$$

$$= ABE + A\bar{D}E + \bar{B}\bar{C}E + \bar{A}\bar{B}E + (\bar{A} + \bar{B})(\bar{A} + D)(B + C)(A + B)\bar{E} \quad (5.8)$$

$$= (AB + A\bar{D} + \bar{B}\bar{C} + \bar{A}\bar{B})E + \overline{(AB + A\bar{D} + \bar{B}\bar{C} + \bar{A}\bar{B})}\bar{E} \quad (5.9)$$

$$= AB + A\bar{D} + \bar{B}\bar{C} + \bar{A}\bar{B} \leftrightarrow E \quad (5.10)$$

$$h(R, S, T, U, V) = RSTUV + RS\bar{T}UV + R\bar{S}TUV + R\bar{S}\bar{T}\bar{U}V + \bar{R}STUV + \bar{R}S\bar{T}UV + \bar{R}\bar{S}TUV + \bar{R}\bar{S}\bar{T}\bar{U}\bar{V} \quad (5.11)$$

$$= RSUV + RTUV + R\bar{S}\bar{T}\bar{U}V + SUV + STUV + \bar{R}\bar{S}TUV + \bar{R}\bar{S}\bar{T}\bar{U}\bar{V} \quad (5.12)$$

$$= (RV + SV + \bar{R}\bar{S}\bar{V}) \cdot (SU + TU + \bar{S}\bar{T}\bar{U}) \quad (5.13)$$

$$= ((R + S)V + \overline{(R + S)}\bar{V}) \cdot ((S + T)U + \overline{(S + T)}\bar{U}) \quad (5.14)$$

$$= (R + S \leftrightarrow V) \cdot (S + T \leftrightarrow U) \quad (5.15)$$

Die beiden Booleschen Ausdrücke (5.10) beziehungsweise (5.15) sind in MT-Form und repräsentieren Minimale Theorien von $\mathcal{K}_{5.1}$ beziehungsweise $\mathcal{K}_{5.2}$. Es gilt zusätzlich zu bemerken, dass (5.15), aufgrund des in beiden einfachen Minimalen Theorien enthaltenen S , komplex ist. Ferner sind beide Minimalen Theorien genau für jene Belegungen wahr, die den Konfigurationen aus der entsprechenden Konfigurationstabelle entsprechen. Diese letzte, auf die angenommene ideale Datenlage zurückgehende Eigenschaft gilt für alle Minimale Theorien von Konfigurationstabellen, die homogen und maximal divers sind. Sie wird dementsprechend für alle Minimalen Theorien dieses Kapitels erfüllt sein.

Interpretiert man die Minimalen Theorien (5.10) beziehungsweise (5.15) kausal, resultieren Kausalmodelle, die mit den entsprechenden Konfigurationen verträglich sind. Aus der einfachen Minimalen Theorie (5.10) ergibt sich ein mit $\mathcal{K}_{5.1}$ verträgliches Kausalmodell ω_g mit der Wirkung E und deren vier Alternativursachen AB , $A\bar{D}$, $\bar{B}\bar{C}$ und $\bar{A}\bar{B}$. Indem das Modell nur direkte Abhängigkeiten aufweist, existieren weder einfache Minimale Theorien, die zur Bestätigung indirekter Zusammenhänge herangezogen werden, noch gibt es Switching-Literale. In Form von (5.1) notiert, resultiert also das folgende Modell:

$$\omega_g := [AB + A\bar{D} + \bar{B}\bar{C} + \bar{A}\bar{B} \leftrightarrow E] \quad (5.16)$$

Die komplexe Minimale Theorie $(R+S \leftrightarrow V) \cdot (S+T \leftrightarrow U)$ beschreibt kausal interpretiert ein mit $\mathcal{K}_{5.2}$ verträgliches Kausalmodell ω_h , das eine gemeinsame Ursache S von V und U aufweist:

$$\omega_h := [(R + S \leftrightarrow V) \cdot (S + T \leftrightarrow U)] \quad (5.17)$$

Graphisch dargestellt entsprechen die beiden Kausalmodelle ω_g und ω_h genau jenen Kausalgraphen aus den Abbildungen 5.1 und 5.2, die den Konfigurationstabellen $\mathcal{K}_{5.1}$ und $\mathcal{K}_{5.2}$ zugrunde gelegt wurden.

Für eine systematische Suche nach allen mit einer Konfigurationstabelle verträglichen Kausalmodellen ist eine derart direkte Herangehensweise wenig sinnvoll. So darf durchaus eingewendet werden, dass die obigen Umformungen wenig zielführend sind, wenn die datengenerierenden Kausalstrukturen gänzlich unbekannt sind und nicht bereits zu Beginn klar ist, wie der aus den Umformungen resultierende Boolesche Ausdruck aussehen soll. Hat man keine Vorstellung davon, wie diese Kausalstrukturen aussehen, wird es auch mit einem sehr gut geschulten Auge extrem schwierig sein, KDN-Formen mittels direkter Verwendung der Booleschen Gesetze in MT-Formen zu bringen, sofern sich eine KDNF überhaupt als MTF darstellen lässt. Erschwerend kommt hinzu, dass es aufgrund der allgemeinen empirischen Unterbestimmtheit oftmals auch mehrere verträgliche Kausalmodelle und damit innerhalb der Konfigurationsalgebra auch mehrere MT-Formen der Instanzfunktion gibt, die das Instantiierungsverhalten der Faktoren aus der entsprechenden Konfigurationstabelle mathematisch zum Ausdruck bringt. Wie sich später noch zeigen wird, liegt etwa gerade für die Konfigurationstabelle $\mathcal{K}_{5.1}$ eine solche Mehrdeutigkeit vor. Nichtsdestotrotz liefert die obige Einbettung des analytischen Moments in die Konfigurationsalgebra ein zentrales Arbeitsmittel bei der Systematisierung des Aufdeckungsverfahrens csCNA+. Unter anderem wird dadurch der Einsatz gängiger Boolescher Verarbeitungsprozesse von Booleschen Ausdrücken gerechtfertigt und eine einheitliche Formalisierung der prinzipiellen Verfahrensschritte ermöglicht. Letzteres erlaubt eine direkte Gegenüberstellung der beiden bereits existierenden Algorithmen von QCA und CNA.

Die grundlegenden Prozessschritte für die Systematisierung des analytischen Moments wurden in Abschnitt 3.2 bereits skizziert. Im Grundsatz werden dabei die Minimalen Theorien schrittweise aufgebaut. Diese Vorgehensweise ist deshalb sinnvoll, weil die aus dem analytischen Moment resultierenden Kausalmodelle kausal interpretierte Minimale Theorien einer Konfigurationstabelle $\mathcal{K}_F$ sind. Die Minimalen Theorien ihrerseits bestehen notwendigerweise aus RDN-Bikonditionalen, die von $\mathcal{K}_F$ impliziert werden. Um also Kausalmodelle zu finden, müssen die Minimalen Theorien bekannt sein und um letztere zu finden, müssen die aus $\mathcal{K}_F$ gefolgerten RDN-Bikonditionale vorliegen. Folgende drei Hauptschritte haben wir definiert, die im Folgenden detailliert diskutiert und systematisiert werden:

- (I) Bestimmung aller RDN-Bikonditionale,

- (II) Bildung von Minimalen Theorien,
- (III) Bildung von Kausalmodellen.

5.2. Bestimmung aller RDN-Bikonditionale

Ausgehend von einer homogenen und maximal diversen Konfigurationstabelle $\mathcal{K}_F$ über die geprüfte Faktormenge F , beginnt man beim analytischen Moment damit, die RDN-Bikonditionale zu bestimmen. Die RDN-Bikonditionale bilden die Grundbausteine der Minimalen Theorien von $\mathcal{K}_F$, die sich aus der redundanzfreien konjunktiven Verknüpfung von ersteren ergeben. Stehen für eine Kausalanalyse alleine die Konfigurationen aus $\mathcal{K}_F$ als Informationen zur Verfügung, ist vorerst davon auszugehen, dass von den Literalen jedes Faktors aus F RDN-Bikonditionale existieren. Wie bereits bei der Skizzierung des analytischen Moments aus Abschnitt 3.2 argumentiert (vgl. S. 102 ff.), muss jedoch nicht zwangsläufig von jedem Literal der Faktoren aus F nach minimal notwendigen Disjunktionen minimal hinreichender Konjunktionsterme gesucht werden. Für jene Faktoren aus F , die die beiden bekannten Ausschlusskriterien (WD1) und (WD2) erfüllen, kann von Beginn weg auf eine Suche nach entsprechenden RDN-Bedingungen verzichtet werden, da es basierend auf der angenommenen idealen Datenlage keine solche RDN-Bedingungen geben kann. Bevor also nach RDN-Bikonditionalen gesucht wird, wird für alle Faktoren geprüft, ob sie die folgenden Bedingungen (WD1) oder (WD2) erfüllen (vgl. S. 105 ff.):

- (WD1) Der Faktor $X_i \in F$ ist in der geprüften Konfigurationstabelle $\mathcal{K}_F$ (a) in allen Konfigurationen anwesend oder (b) in allen Konfigurationen abwesend.
- (WD2) In der geprüften Konfigurationstabelle $\mathcal{K}_F$ sind zwei Konfigurationen enthalten, die sich nur darin unterscheiden, dass X_i in der einen Konfiguration anwesend ist und in der anderen nicht.

Für jeden Faktor aus F , der weder (WD1) noch (WD2) erfüllt, wird nach den RDN-Bikonditionalen gesucht. Für eine bessere Übersichtlichkeit soll die Menge dieser Faktoren mit einem Stern „*“ ausgezeichnet und als Menge F^* notiert werden, wobei $F^* \subseteq F$ gilt. Die Menge F^* , die sich durch den Ausschluss jener Faktoren aus F ergibt, die die beiden Ausschlusskriterien (WD1) oder (WD2) erfüllen, kann mit dem folgenden, in Pseudocode dargestellten Algorithmus 5.1 zusammengefasst werden.

Algorithmus 5.1 Suche nach den Faktoren aus $\mathbf{F}^*$

Input: Homogene und maximal diverse Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ über $\mathbf{F} = \{X_1, \dots, X_n\}$

Output: Menge $\mathbf{F}^*$ von Faktoren, von denen nach RDN-Bedingungen gesucht wird

- 1: **for each** Faktor X_i aus $\mathbf{F}$ **do**
 - 2: **if** X_i erfüllt weder (WD1) noch (WD2) **then**
 - 3: Füge X_i zu $\mathbf{F}^*$ hinzu
 - 4: **return** $\mathbf{F}^*$
-

Wendet man den Algorithmus 5.1 auf die beiden Konfigurationstabellen $\mathcal{K}_{5,1}$ über $\mathbf{G} = \{A, B, C, D, E\}$ und $\mathcal{K}_{5,2}$ über $\mathbf{H} = \{R, S, T, U, V\}$ an, resultieren für erstere der Faktor E und für letztere die Faktoren U und V – kurz: $\mathbf{G}^* = \{E\}$ und $\mathbf{H}^* = \{U, V\}$.

Von den nach Ausschluss mittels (WD) übrigbleibenden Faktoren werden die RDN-Bedingungen ihrer Literale gesucht. Bei diesem Prozessschritt kommt die Konfigurationsalgebra besonders zum tragen. Basierend auf diesem mathematischen Fundament können die Konfigurationen durch Instanzfunktionen ausgedrückt und so direkt für verschiedene Boolesche Analysemethoden zugänglich gemacht werden. Einer solchen Transformation zufolge entspricht die Ermittlung der RDN-Bedingungen eines Literals X_i aus einer Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ über $\mathbf{F} = \{X_1, \dots, X_n\}$ der Suche nach den RDN-Formen einer inneren Instanzfunktion f_{X_i} . Konkret wird für jeden nach Verwendung von (WD) übrigbleibenden Faktor $X_i \in \mathbf{F}^* \subseteq \mathbf{F}$ die innere Instanzfunktion f_{X_i} gebildet, die das Instantiierungsverhalten des Faktors X_i in Abhängigkeit der restlichen Faktoren aus $\mathbf{F}$ beschreibt. Für eine solche innere Instanzfunktion $f_{X_i} : \mathcal{D} \rightarrow \mathcal{B}$, wobei $\mathcal{B} = \{0, 1\}$ und $\mathcal{D} \subseteq \mathcal{B}^{n-1}$, gilt:

$$f_{X_i}(X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n) = \begin{cases} 1, & (X_1, \dots, X_{i-1}, 1, X_{i+1}, \dots, X_n) \in \mathcal{I}(f) \\ 0, & (X_1, \dots, X_{i-1}, 0, X_{i+1}, \dots, X_n) \in \mathcal{I}(f) \end{cases} \quad (5.18)$$

$\mathcal{I}(f)$ entspricht dabei dem Instantiierungsbereich der totalen Instanzfunktion f , die das Instantiierungsverhalten der Faktoren aus $\mathbf{F}$ gemäss $\mathcal{K}_{\mathbf{F}}$ ausdrückt, indem die Belegungen aus $\mathcal{I}(f)$ den Konfigurationen aus $\mathcal{K}_{\mathbf{F}}$ entsprechen. Die Instanzfunktion f und eine innere Instanzfunktion f_{X_i} stehen gemäss Kapitel 4 derart in Beziehung, dass eine RDN-Form ϕ_{X_i} von f_{X_i} als Bestandteil eines RDN-Bikonditionals ($\phi_{X_i} \leftrightarrow X_i$) Teil einer MT-Form von f sein kann, die ihrerseits eine Minimale Theorie von $\mathcal{K}_{\mathbf{F}}$ beschreibt.

Angewendet auf die Konfigurationstabelle $\mathcal{K}_{5,1}$ über $\mathbf{G} = \{A, B, C, D, E\}$ mit $\mathbf{G}^* = \{E\}$, wird die innere Instanzfunktion g_E mit dem abhängigen Faktor E gebildet, die durch die Wertetabelle aus Tabelle 5.4 vollständig beschrieben ist.

A	B	C	D	$E = g_E(A, B, C, D)$
1	1	1	1	1
1	1	1	0	1
1	1	0	1	1
1	1	0	0	1
1	0	1	1	0
1	0	1	0	1
1	0	0	1	1
1	0	0	0	1
0	1	1	1	0
0	1	1	0	0
0	1	0	1	0
0	1	0	0	0
0	0	1	1	1
0	0	1	0	1
0	0	0	1	1
0	0	0	0	1

Tabelle 5.4.: Wertetabelle der inneren Instanzfunktion g_E von g für die Konfigurationstabelle $\mathcal{K}_{5.1}$.

Bei der Konfigurationstabelle $\mathcal{K}_{5.2}$ werden aufgrund der beiden Faktoren U und V aus $\mathbf{H}^*$ auch zwei solche innere Instanzfunktionen h_U und h_V gebildet, deren Wertetabellen in Tabelle 5.5 dargestellt sind.

Im Gegensatz zu einer Instanzfunktion f , die das in einer Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ dargestellte Instantiierungsverhalten der Faktoren aus $\mathbf{F}$ ausdrückt, ist eine innere Instanzfunktion f_{X_i} , die das Instantiierungsverhalten von $X_i \in \mathbf{F}$ in Abhängigkeit der restlichen Faktoren aus $\mathbf{F}$ beschreibt, nicht zwangsläufig total. Ob f_{X_i} total oder partiell ist, ist von den untersuchten Daten abhängig, aus denen f_{X_i} gebildet wird. Wird eine homogene und maximal diverse $\mathcal{K}_{\mathbf{F}}$ über $\mathbf{F} = \{X_1, \dots, X_n\}$ von einer komplexen Kausalstruktur generiert, die relativ zur untersuchten Faktormenge mindestens zwei Wirkungen X_i und X_j beinhaltet, ist nicht zusammen mit allen $2^{(n-1)}$ Kombinationen von an- und abwesenden Faktoren aus $\mathbf{F} \setminus \{X_i\}$ ein dazugehöriges Instantiierungsverhalten des Faktors X_i in $\mathcal{K}_{\mathbf{F}}$ vorhanden.³ Weil es neben X_i mindestens ein Literal X_j eines anderen Faktors aus $\mathbf{F}$ gibt, das kausal von Literalen von Faktoren aus $\mathbf{F}$ abhängig ist, kann der Faktor X_j unter der Voraussetzung der (nicht trivialen) Homogenität nicht derart beliebig variieren, dass er zusam-

³Eine Ausnahme bildet jene Konfigurationstabelle, für die die triviale Homogenität gilt, gemäss der potentielle Störursachen derart variieren, dass alle logisch möglichen Konfigurationen auftreten (vgl. S. 123). Dieser uninteressante Spezialfall, aus dem keine Kausalmodelle gefolgert werden können, wird in diesem Kapitel nicht betrachtet.

R	S	T	V	$U = h_U(R, S, T, V)$	R	S	T	U	$V = h_V(R, S, T, U)$
1	1	1	1	1	1	1	1	1	1
1	1	0	1	1	1	1	0	1	1
1	0	1	1	1	1	0	1	1	1
1	0	0	1	0	1	0	0	0	1
0	1	1	1	1	0	1	1	1	1
0	1	0	1	1	0	1	0	1	1
0	0	1	0	1	0	0	1	1	0
0	0	0	0	0	0	0	0	0	0

(a)
(b)

Tabelle 5.5.: Wertetabellen der beiden inneren Instanzfunktionen h_U und h_V von h für $\mathcal{K}_{5,2}$.

men mit jeder logisch möglichen Kombination von an- und abwesenden Faktoren aus $\mathbf{F} \setminus \{X_i, X_j\}$ sowohl an- als auch abwesend sein kann. Anhand von $\mathcal{K}_{5,2}$ veranschaulicht, zeigt sich beispielsweise für die innere Instanzfunktion h_V , dass sie unter anderem für die Belegung $(0, 0, 0, 1)$ nicht definiert ist. Wäre sie für diese Belegung definiert, würde dies bedeuten, dass h aus (5.4) entweder für $(0, 0, 0, 1, 0)$ oder $(0, 0, 0, 1, 1)$ den Wert 1 annimmt. Weil f das Verhalten der Faktoren in der Konfigurationstabelle $\mathcal{K}_{5,2}$ beschreibt, würde daraus die Interpretation folgen, dass U anwesend sein kann, ohne dass S oder T anwesend sind. Dies ist jedoch vor dem Hintergrund der dahinterliegenden Common-Cause-Struktur sowie der angenommenen Homogenität nicht möglich und würde zu einem Widerspruch führen. Analoges gilt für h_U und beispielsweise dieselbe Belegung $(0, 0, 0, 1)$.

Indem ein RDN-Bikonditional von X_i ausgedrückt wird durch eine innere Instanzfunktion f_{X_i} in RDN-Form, entspricht die Suche nach den RDN-Bikonditionalen von X_i der Suche nach den RDN-Formen einer inneren Instanzfunktion f_{X_i} . Konkret wird gemäss Definition 4.15 nach den aus Primimplikanten bestehenden irredundanten Abdeckungen von f_{X_i} gesucht. Wie sich im folgenden Abschnitt zeigen wird, ermöglicht diese Transformation der Suche nach RDN-Bedingungen von X_i zur Suche nach RDN-Formen der inneren Instanzfunktion f_{X_i} den direkten Einsatz gängiger Methoden aus der Booleschen Algebra, mit deren Hilfe sich die Suche nach den RDN-Bikonditionalen systematisieren lässt.

5.2.1. Systematisierte Suche nach den RDN-Formen der inneren Instanzfunktionen

Ähnlich wie im Fall der manuellen Verwendung der Booleschen Gesetze bei der Suche nach den MT-Formen einer Instanzfunktion ist eine manuelle Verwendung der Booleschen Gesetze für die Suche nach den RDN-Formen einer inneren Instanzfunktion im Prinzip möglich, aber im Allgemeinen nur bedingt für eine Systematisierung geeignet. Zur Veranschaulichung dieser Schwierigkeit wird erneut $\mathcal{K}_{5,1}$ über $\mathbf{G} = \{A, B, C, D, E\}$ betrachtet, für die mit $\mathbf{G}^* = \{E\}$ im vorliegenden Schritt nach den RDN-Bikonditionalen von E gesucht wird. Über deren Minterme kann die totale innere Instanzfunktion g_E folgendermassen eindeutig als KDNF dargestellt werden, wobei die links des Gleichheitszeichens in Klammern stehenden Argumente von g_E aus Platzgründen durch den Platzhalter $(\cdot)$ ersetzt werden:

$$g_E(\cdot) = ABCD + ABC\bar{D} + AB\bar{C}D + AB\bar{C}\bar{D} + A\bar{B}CD + A\bar{B}\bar{C}D + A\bar{B}\bar{C}\bar{D} + \bar{A}\bar{B}CD + \bar{A}\bar{B}\bar{C}D + \bar{A}\bar{B}\bar{C}\bar{D} + \bar{A}\bar{B}\bar{C}\bar{D} \quad (5.19)$$

Für die Suche nach den RDN-Formen der inneren Instanzfunktion g_E wird die KDNF aus (5.19) in einem ersten Schritt zum Ausdruck $AB + A\bar{B}\bar{D} + A\bar{B}\bar{C} + \bar{A}\bar{B}$ vereinfacht:

$$\begin{aligned} g_E(\cdot) &= ABCD + ABC\bar{D} + AB\bar{C}D + AB\bar{C}\bar{D} + A\bar{B}CD + A\bar{B}\bar{C}D + A\bar{B}\bar{C}\bar{D} + \bar{A}\bar{B}CD + \bar{A}\bar{B}\bar{C}D + \bar{A}\bar{B}\bar{C}\bar{D} \\ &\quad + \bar{A}\bar{B}\bar{C}\bar{D} + \bar{A}\bar{B}\bar{C}\bar{D} \\ &= ABC(D + \bar{D}) + AB\bar{C}(D + \bar{D}) + A\bar{B}CD + A\bar{B}\bar{C}(D + \bar{D}) + \bar{A}\bar{B}C(D + \bar{D}) + \bar{A}\bar{B}\bar{C}(D + \bar{D}) \\ &= AB(C + \bar{C}) + A\bar{B}(CD + \bar{C}) + \bar{A}\bar{B}(C + \bar{C}) \\ &= AB + A\bar{B}(CD + \bar{C}D + \bar{C}\bar{D} + \bar{C}\bar{D}) + \bar{A}\bar{B} \\ &= AB + A\bar{B}(\bar{D}(C + \bar{C}) + \bar{C}(D + \bar{D})) + \bar{A}\bar{B} \\ &= AB + A\bar{B}\bar{D} + A\bar{B}\bar{C} + \bar{A}\bar{B}. \end{aligned}$$

Der Ausdruck $AB + A\bar{B}\bar{D} + A\bar{B}\bar{C} + \bar{A}\bar{B}$ lässt sich folgendermassen weiter reduzieren, wobei der mit dem $(*)$ -Symbol gekennzeichnete Umformungsschritt und die damit einhergehende Erweiterung des vorangehenden Ausdrucks bereits anzeigt, dass manuelle Umformungen kaum geeignet sind, um RDN-Formen zu finden:

$$\begin{aligned} g_E(\cdot) &= AB + A\bar{B}\bar{D} + A\bar{B}\bar{C} + \bar{A}\bar{B} \\ &= A(B + \bar{B}\bar{D}) + \bar{B}(A\bar{C} + \bar{A}) \\ &\stackrel{(*)}{=} A(BD + B\bar{D} + B\bar{D} + \bar{B}\bar{D}) + \bar{B}(A\bar{C} + \bar{A}C + \bar{A}\bar{C} + \bar{A}\bar{C}) \\ &= A(B + \bar{D}) + \bar{B}(\bar{C} + \bar{A}) \\ &= AB + A\bar{D} + \bar{B}\bar{C} + \bar{A}\bar{B}. \end{aligned}$$

Die mittels des Booleschen Ausdrucks $AB + A\bar{D} + \bar{B}\bar{C} + \bar{A}\bar{B}$ dargestellte Instanzfunktion g_E ist in RDN-Form, womit ein erstes RDN-Bikonditional von E vorliegt:

$$AB + A\bar{D} + \bar{B}\bar{C} + \bar{A}\bar{B} \leftrightarrow E \quad (5.20)$$

Setzt man ausgehend von $AB + A\bar{B}\bar{D} + A\bar{B}\bar{C} + \bar{A}\bar{B}$ erneut zu einer Reduktion an, ergibt sich eine zusätzliche RDN-Form von g_E . Auch hier wird bei (*) der vorangehende Ausdruck erweitert, um schliesslich zur RDN-Form zu gelangen:

$$\begin{aligned} g_E(\cdot) &= AB + A\bar{B}\bar{D} + A\bar{B}\bar{C} + \bar{A}\bar{B} \\ &= AB + A\bar{B}\bar{C} + A\bar{B}\bar{D} + \bar{A}\bar{B} \\ &= A(B + \bar{B}\bar{C}) + \bar{B}(A\bar{D} + \bar{A}) \\ &\stackrel{(*)}{=} A(BC + \bar{B}\bar{C} + \bar{B}\bar{C} + \bar{B}\bar{C}) + \bar{B}(A\bar{D} + \bar{A}\bar{D} + \bar{A}\bar{D} + \bar{A}\bar{D}) \\ &= A(B + \bar{C}) + \bar{B}(\bar{D} + \bar{A}) \\ &= AB + A\bar{C} + \bar{B}\bar{D} + \bar{A}\bar{B} \end{aligned}$$

Aus dieser RDN-Form kann direkt das folgende RDN-Bikonditional von E abgelesen werden:

$$AB + A\bar{C} + \bar{B}\bar{D} + \bar{A}\bar{B} \leftrightarrow E \quad (5.21)$$

Tatsächlich können durch ähnliche Umformungen noch zwei weitere RDN-Formen von g_E gebildet werden, aus denen sich die folgenden beiden RDN-Bikonditionale von E ergeben⁴:

$$AB + A\bar{D} + A\bar{C} + \bar{A}\bar{B} \leftrightarrow E \quad (5.22)$$

$$AB + \bar{B}\bar{D} + \bar{B}\bar{C} + \bar{A}\bar{B} \leftrightarrow E \quad (5.23)$$

Die bezüglich der RDN-Bikonditionale vorhandene Mehrdeutigkeit der Konfigurationstabelle $\mathcal{K}_{5,1}$ kann nicht direkt aus $\mathcal{K}_{5,1}$ selbst abgelesen werden, geschweige denn kann von Beginn weg angegeben werden, wie viele RDN-Bikonditionale und daraus gebildete Kausalmodelle es gibt. Das Beispiel zeigt, dass eine direkte Verwendung der Booleschen Gesetze für das konfigurationsale kausale Modellieren nicht praktikabel ist und die Analyse von Konfigurationstabellen nach einer Systematisierung der Datenverarbeitung verlangt. Dieser Umstand zeigt sich vor allem auch daran, dass für die Umformungen die Ausdrücke teils zuerst verkompliziert werden müssen, damit man am Ende zu einer RDN-Form gelangt. Bei den oben explizit durch (*) ausgezeichneten Umformungsschritten wird der dem jeweiligen Schritt vorangehende Ausdruck zuerst erweitert, um daraufhin wieder vereinfacht zu werden. Die Vereinfachungen von Booleschen Ausdrücken mittels direkter Verwendung der Booleschen Gesetze sind somit insgesamt

⁴Eine Herleitung der zwei weiteren RDN-Formen von g_E findet man in Abschnitt A.2 des Anhangs (vgl. S. 249).

nicht linear in dem Sinne, dass Ausdrücke mit jedem Schritt stetig kürzer oder immer die gleichen Umformungsschritte iterativ angewendet werden.

Mit csQCA und csCNA existieren zwei Verfahren, die eine solche Systematisierung mittels Algorithmen ermöglichen. Beide Algorithmen werden in der Folge basierend auf dem gemeinsamen Fundament der Konfigurationsalgebra beschrieben und miteinander verglichen. Dazu werden die ihnen zugrunde gelegten Prinzipien präsentiert und die Datenverarbeitungsprozesse im Detail dargelegt. Beide Algorithmen folgen im Rahmen der Suche nach den RDN-Bedingungen eines Literals W , das heisst nach den RDN-Formen einer inneren Instanzfunktion f_W , den folgenden zwei Prozessschritten, die sich direkt aus der Definition einer RDN-Form ergeben (vgl. Def. 4.15, S. 170):

- (a) Suche nach den Primimplikanten von f_W ,
- (b) Suche nach den aus Primimplikanten bestehenden irredundanten Abdeckungen von f_W (RDN-Formen von f_W).

Es gilt dabei voranzuschicken, dass die im Folgenden csQCA zugeschriebenen Algorithmen teils von den Standard-Algorithmen abweichen. Das liegt hauptsächlich daran, dass die standardmässig genutzte Version von csQCA nicht genau dasselbe Suchziel verfolgt wie jenes, das der hier vorliegende kausalthoretische Rahmen nahelegt. Weil die essentiellen Prinzipien, die das csQCA-Verfahren auszeichnen, trotz vorgenommenen Anpassungen jedoch dieselben bleiben, werden wir auch diese adaptierten Versionen als csQCA-Algorithmen bezeichnen. Dies nicht zuletzt auch zwecks einer übersichtlichen Kontrastierung der Verfahren. Wo genau und weshalb die hier adaptierten csQCA-Algorithmen vom Standard abweichen, wird sich im folgenden Abschnitt zeigen.

5.2.2. csQCA und das Quine-McCluskey-Verfahren

Bei QCA ist das Forschungsdesign typischerweise so aufgebaut, dass nach (kausalen) „Erklärungen“ für das Instantiierungsverhalten eines bestimmten Faktors W gesucht wird (vgl. u. a. Berg-Schlosser & Meur (2009), S. 24). Unter Berücksichtigung des kausalthoretischen Rahmens bedeutet dies für das analytische Moment, dass, ausgehend von einer Konfigurationstabelle $\mathcal{K}_F$ über die geprüfte Faktormenge F , $W \in F$, nach den einfachen Minimalen Theorien von W gesucht wird. Angesichts dessen entspricht obiges Beispiel mit der datengenerierenden Kausalstruktur aus Abbildung 5.1, die das einzelne, von Literalen der restlichen Faktoren der geprüften Faktormenge G kausal abhängige Literal E aufweist, genau jenem Forschungsdesign, für das

csQCA hauptsächlich vorgesehen ist. Dabei liegen mit der homogenen und maximal diversen Konfigurationstabelle $\mathcal{K}_{5,1}$ ideale Daten vor.

Allgemein werden für die Suche nach den RDN-Bedingungen eines Literals W , beziehungsweise den RDN-Formen der entsprechenden inneren Instanzfunktion f_W , von csQCA Methoden eingesetzt, die in anderen Anwendungsgebieten wie der Schaltalgebra bereits rege zur Minimierung von Booleschen Ausdrücken genutzt werden und sich ebenfalls für den Einsatz im vorliegenden Kontext eignen.⁵ Als prominentes Minimierungsverfahren ist etwa das Karnaugh-Veitch-Diagramme (KV-Diagramme) nutzende Verfahren zu nennen, das graphische Elemente zur Gruppierung gleicher Belegungen nutzt, die die Auslesung einer RDN-Form wesentlich vereinfachen (vgl. u. a. Veitch (1952); Karnaugh (1953); Nelson et al. (1995), S. 175 ff.). Ein weiteres prominentes Verfahren, das für den vorliegenden Zweck besonders nützlich ist, ist das sogenannte *Quine-McCluskey-Verfahren* (QMC) (vgl. u. a. Quine (1952); McCluskey (1956); Nelson et al. (1995), S. 211 ff.). Dieses Verfahren entspricht dem Hauptverfahren von csQCA für die Suche nach den RDN-Formen einer inneren Instanzfunktion (vgl. bspw. Schneider & Wagemann (2012), S. 104 ff.).

Beim QMC-Verfahren werden die verschiedenen Minterme einer KDNF paarweise verglichen und schrittweise Literale weggestrichen. Das Prinzip des QMC-Verfahrens beruht auf dem Distributivgesetz, dem Komplement sowie dem neutralen Element der Booleschen Algebra, mit deren Hilfe zwei Konjunktionsterme, die sich lediglich in der Negation einer Variable X_i unterscheiden, um diese reduziert und verschmolzen werden können:

$$\begin{aligned} X_1 \cdot \dots \cdot X_i \cdot \dots \cdot X_n + X_1 \cdot \dots \cdot \overline{X_i} \cdot \dots \cdot X_n &= X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \cdot (X_i + \overline{X_i}) \\ &= X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \cdot 1 \\ &= X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \end{aligned}$$

Diese Abfolge von Umformungsschritten, auch *Resolutionsregel* (RES) genannt, wurde bei den obigen detailliert angegebenen Umformungsschritten für g_E bereits ausgiebig angewendet. Das QMC-Verfahren zeichnet sich nun vor allem durch die Darstellung aus, die es ermöglicht alle RDN-Formen einer inneren Instanzfunktion f_W systematisch zu finden. Wie bereits einleitend erwähnt, wird das Verfahren dazu in zwei Teile unterteilt, die im Folgenden beschrieben werden und auf welche die Definition der RDN-Form abgestimmt ist: (a) die Suche nach den Primimplikanten von f_W und (b) die Suche nach den irredundanten Abdeckungen von f_W .

⁵Eine Auswahl bekannter graphischer und algorithmischer Verfahren ist u. a. zu finden bei Denis-Papin et al. (1974), Kapitel 9.

(a) Suche nach den Primimplikanten

Im hier vorliegenden Kontext kann die obige Resolutionsregel folgendermassen genutzt und interpretiert werden: Seien $X_1 \cdot \dots \cdot X_i \cdot \dots \cdot X_n$ und $X_1 \cdot \dots \cdot \bar{X}_i \cdot \dots \cdot X_n$ zwei hinreichende Bedingungen von W . Die beiden hinreichenden Bedingungen unterscheiden sich lediglich bezüglich des Faktors X_i , der in einer Bedingung nicht-negiert in Form des positiven Literals X_i und in der anderen Bedingung negiert in Form des negativen Literals $\bar{X}_i$ enthalten ist. Offensichtlich spielt es damit für diese Bedingungen bezüglich der Suffizienz für W keine Rolle, ob der Faktor X_i an- oder abwesend ist. X_i und $\bar{X}_i$ können daher aus Gründen der Redundanz aus der Bedingung entfernt werden. Mittels der Konfigurationsalgebra lässt sich dies folgendermassen mathematisch ausdrücken⁶:

$$\begin{aligned} & (X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W) \cdot (X_1 \cdot \dots \cdot X_{i-1} \cdot \bar{X}_i \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W) \\ & = X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W \end{aligned} \quad (\text{RES})$$

Für die systematische Umsetzung ist es sinnvoll, die Resolutionsregel mithilfe der inneren Instanzfunktion f_W auszudrücken. Dabei gilt basierend auf Definition 4.14, dass $X_1 \cdot \dots \cdot X_i \cdot \dots \cdot X_n$ und $X_1 \cdot \dots \cdot \bar{X}_i \cdot \dots \cdot X_n$ zwei Implikanten einer inneren Instanzfunktion f_W sind. Ist $X_1 \cdot \dots \cdot X_i \cdot \dots \cdot X_n$ oder $X_1 \cdot \dots \cdot \bar{X}_i \cdot \dots \cdot X_n$ gegeben, nimmt f_W den Funktionswert 1 an. Die Resolutionsregel (RES) kann demnach folgendermassen reformuliert werden:

- (1) $X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n$ ist ein Implikant von f_W
 - (2) $X_1 \cdot \dots \cdot X_{i-1} \cdot \bar{X}_i \cdot X_{i+1} \cdot \dots \cdot X_n$ ist ein Implikant von f_W
-
- $X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n$ ist ein Implikant von f_W

Hinsichtlich des Ziels die Primimplikanten einer inneren Instanzfunktion f_W zu finden, wird die Resolutionsregel nun so lange auf die Implikanten von f_W angewendet, bis es keine zwei Implikanten mehr gibt, die via (RES) verschmolzen werden können. Bei den Implikanten, die sich nicht weiter verschmelzen lassen, handelt es sich um Primimplikanten. Dazu werden in einem ersten Schritt die Minterme der inneren Instanzfunktion so zusammengestellt, dass solche mit der gleichen Anzahl negativer Literale eine Gruppe bilden und die Gruppen ihrerseits mit aufsteigender Anzahl Negationen aufgelistet werden. Für obige innere Instanzfunktion g_E (vgl. S. 177) entspricht diese Liste jener aus Tabelle 5.6, wobei jeder der Minterme durch ein $\mathbf{m}_j$, $j \in \mathbb{N}$, eindeutig gekennzeichnet wird. Die Dezimalzahl des Indexes eines Minterms entspricht dabei der Binärzahl, die sich aus der Aneinanderreihung jener Werte ergibt, die dazu führen, dass der Minterm den Wert 1 annimmt. Konkret anhand von $\mathbf{m}_{15}$ aus Tabelle 5.6 veranschaulicht, nimmt dieser Minterm genau dann den Wert 1 an, wenn alle Faktoren A , B , C und D den Wert 1 annehmen. Aneinandergereiht ergibt dies 1111, was seinerseits die Binärzahl für 15 des Dezi-

⁶Ein Nachweis der Resolutionsregel findet man in Abschnitt A.2 des Anhangs (vgl. S. 250).

$\mathbf{m}_{15}$	A	B	C	D
$\mathbf{m}_{14}$	A	B	C	$\bar{D}$
$\mathbf{m}_{13}$	A	B	$\bar{C}$	D
$\mathbf{m}_{12}$	A	B	$\bar{C}$	$\bar{D}$
$\mathbf{m}_{10}$	A	$\bar{B}$	C	$\bar{D}$
$\mathbf{m}_9$	A	$\bar{B}$	$\bar{C}$	D
$\mathbf{m}_3$	$\bar{A}$	$\bar{B}$	C	D
$\mathbf{m}_8$	A	$\bar{B}$	$\bar{C}$	$\bar{D}$
$\mathbf{m}_2$	$\bar{A}$	$\bar{B}$	C	$\bar{D}$
$\mathbf{m}_1$	$\bar{A}$	$\bar{B}$	$\bar{C}$	D
$\mathbf{m}_0$	$\bar{A}$	$\bar{B}$	$\bar{C}$	$\bar{D}$

Tabelle 5.6.: Für die Verwendung des Verfahrens nach Quine und McCluskey aufbereitete Mintermtabelle der inneren Instanzfunktion g_E .

malsystems ist. Diese in der Digitaltechnik oftmals verwendete Indexierung ermöglicht es, ohne Querverweise alleine basierend auf dem Index einen bestimmten Minterm zu identifizieren.

Mithilfe dieser Auflistung können nun in einem zweiten Schritt die einzelnen Minterme von benachbarten Gruppen miteinander verglichen werden. Findet man zwei Terme, die bis auf ein Literal, das im einen Term in Form des positiven und im anderen in Form des negativen Literals enthalten ist, identisch sind, werden die Terme gemäss der Resolutionsregel verschmolzen und das wegreduzierte Literal durch einen Stern ersetzt. Dabei kann im gleichen Schritt ein Minterm mehrmals benutzt werden. Aus der Mintermtabelle 5.6 resultiert so Tabelle 5.7a, die alle nach der ersten Verschmelzung entstandenen Implikanten auflistet. Resultieren dabei identische Implikanten mehrmals, werden diese gemäss dem Idempotenzgesetz ($S + S = S$ bzw. $S \cdot S = S$) bis auf einen gestrichen. Kann ein Term nicht weiter reduziert werden, indem sich kein anderer Term mehr finden lässt, mit dem ersterer via der Resolutionsregel verschmolzen werden kann, handelt es sich um einen Primimplikanten. Dieser Umstand wird durch das Kürzel „prim“ in einer zusätzlichen Spalte am Ende der Tabelle gekennzeichnet. Für alle anderen Implikanten wird obiger Schritt solange wiederholt, bis die Resolutionsregel nicht mehr angewendet werden kann. Sterne können dabei ebenfalls als Literale angesehen werden, sodass auch diese bei einer Verschmelzung an identischer Stelle positioniert sein müssen. Für g_E ergibt ein weiterer Verschmelzungsdurchgang die finale Tabelle 5.7b.

Sobald die Resolutionsregel nicht mehr angewendet werden kann und somit nur noch Primimplikanten übrig sind, kann das Reduktionsverfahren abgebrochen werden. Für g_E resultieren die sechs Primimplikanten AB , $A\bar{D}$, $A\bar{C}$, $\bar{B}\bar{D}$, $\bar{B}\bar{C}$ sowie $\bar{A}\bar{B}$. Dass aus dieser Vorgehensweise tatsächlich alle Primimplikanten von g_E und damit alle minimal hinreichenden Bedingungen von

Tabelle 5.7.: Auflistung 5.7a aller Implikanten von g_E , die nach einem ersten Verschmelzungsdurchgang aus der Mintermtabelle 5.6 entstehen und Auflistung 5.7b aller Primimplikanten von g_E .

Eine wichtige Voraussetzung dafür, dass mit dem QMC-Verfahren alle Primimplikanten der inneren Instanzfunktion g_E gefunden werden ist, dass g_E ihrerseits total ist. Im zweiten Beispiel aus Abbildung 5.2 und der damit verbundenen Suche nach den Primimplikanten der inneren Instanzfunktionen h_U und h_V wurde bereits festgestellt, dass diese beiden Instanzfunktionen partiell sind. Für solche Instanzfunktionen muss das QMC-Verfahren für diesen ersten Teilschritt angepasst werden, damit es auch tatsächlich Primimplikanten findet. Konkret werden im Fall einer partiellen inneren Instanzfunktion f_W mit abhängigem Faktor W nicht nur die Minterme von f_W für die Reduktion verwendet, sondern die aus den Mintermen der Erweiterung f_W^+ bestehende Obermenge davon (vgl. u. a. Nelson et al. (1995), S. 119 f.). Würde man etwa für die Bestimmung der Primimplikanten von h_U lediglich die acht Minterme dieser partiellen Instanzfunktion für die Verschmelzungsschritte verwenden, liessen sich mittels der Resolutionsregel nicht alle redundanten Literale aus den Implikanten von h_U wegreduzieren. Aus der entsprechenden in Tabelle 5.8a dargestellten Auflistung der Minterme würden die vermeintlichen, in Tabelle 5.8b

$\mathbf{m}_{15}$	R	S	T	V	
$\mathbf{m}_{13}$	R	S	$\bar{T}$	V	
$\mathbf{m}_{11}$	R	$\bar{S}$	T	V	
$\mathbf{m}_7$	$\bar{R}$	S	T	V	
$\mathbf{m}_5$	$\bar{R}$	S	$\bar{T}$	V	
$\mathbf{m}_2$	$\bar{R}$	$\bar{S}$	T	$\bar{V}$	
(a)					
$\mathbf{m}_5, \mathbf{m}_7, \mathbf{m}_{13}, \mathbf{m}_{15}$					$* \quad S \quad * \quad V \quad \text{prim}$
$\mathbf{m}_5, \mathbf{m}_7$					$\bar{R} \quad S \quad * \quad V \quad \text{prim}$
$\mathbf{m}_2$					$\bar{R} \quad \bar{S} \quad T \quad \bar{V} \quad \text{prim}$
(b)					

Tabelle 5.8.: Ausgehend von den Mintermen der partiellen Instanzfunktion h_U , die in Tabelle 5.8a aufgelistet sind, ergeben sich via QMC die Terme der Auflistung 5.8b. Diese Terme sind keine Primimplikanten von h_U .

dargestellten Primimplikanten folgen. Berücksichtigt man hingegen zur Ermittlung der Primimplikanten die Erweiterung h_U^+ , ergibt sich aus der Ausgangstabelle 5.9a via wiederholter Verwendung der Resolutionsregel gemäss obigem Vorgehen die Liste 5.9b von Primimplikanten.

$\mathbf{m}_{15}$	R	S	T	V	
$\mathbf{m}_{14}$	R	S	T	$\bar{V}$	
$\mathbf{m}_{13}$	R	S	$\bar{T}$	V	
$\mathbf{m}_{11}$	R	$\bar{S}$	T	V	
$\mathbf{m}_7$	$\bar{R}$	S	T	V	
$\mathbf{m}_{12}$	R	S	$\bar{T}$	$\bar{V}$	
$\mathbf{m}_{10}$	R	$\bar{S}$	T	$\bar{V}$	
$\mathbf{m}_6$	$\bar{R}$	S	T	$\bar{V}$	
$\mathbf{m}_5$	$\bar{R}$	S	$\bar{T}$	V	
$\mathbf{m}_3$	$\bar{R}$	$\bar{S}$	T	V	
$\mathbf{m}_8$	R	$\bar{S}$	$\bar{T}$	$\bar{V}$	
$\mathbf{m}_4$	$\bar{R}$	S	$\bar{T}$	$\bar{V}$	
$\mathbf{m}_2$	$\bar{R}$	$\bar{S}$	T	$\bar{V}$	$\mathbf{m}_4, \mathbf{m}_5, \mathbf{m}_6, \mathbf{m}_7, \mathbf{m}_{12}, \mathbf{m}_{13}, \mathbf{m}_{14}, \mathbf{m}_{15} \quad * \quad S \quad * \quad * \quad \text{prim}$
$\mathbf{m}_1$	$\bar{R}$	$\bar{S}$	$\bar{T}$	V	$\mathbf{m}_2, \mathbf{m}_3, \mathbf{m}_6, \mathbf{m}_7, \mathbf{m}_{10}, \mathbf{m}_{11}, \mathbf{m}_{14}, \mathbf{m}_{15} \quad * \quad * \quad T \quad * \quad \text{prim}$
(a)					
(b)					

Tabelle 5.9.: Ausgehend von den Mintermen der Erweiterung h_U^+ , die in Tabelle 5.9a aufgelistet sind, resultieren mittels QMC die Primimplikanten von h_U aus Tabelle 5.9b.

Indem für die Ermittlung der Primimplikanten von einer Erweiterung f_W^+ ausgegangen wird, können auch Primimplikanten folgen, die zwar einen Minterm von f_W^+ , nicht aber einen Minterm der partiellen Instanzfunktion f_W abdecken. Dementsprechend werden bei partiellen Instanzfunk-

tionen in einem letzten Vorgang dieses Teilschrittes des QMC-Verfahrens jene Primimplikanten herausgefiltert, die tatsächlich auch Minterme von f_W abdecken. Für diese Primimplikanten gilt, dass sie in den Mintermen von f_W enthalten sind. Alle restlichen Primimplikanten werden für den weiteren Verlauf nicht mehr berücksichtigt. Für h_U resultieren so insgesamt die Primimplikanten S und T . Analog ergeben sich für h_V die Primimplikanten R und S .

Ist $\mathcal{K}_F$ die untersuchte Konfigurationstabelle mit $W \in \mathbf{F}^*$, kann das obige Vorgehen der systematischen Suche nach allen Primimplikanten einer inneren Instanzfunktion f_W insgesamt beispielsweise mit dem in Pseudocode dargestellten Algorithmus 5.2 zusammengefasst werden. Dabei sollen zur besseren Veranschaulichung die in Fettschrift notierten Kleinbuchstaben $\mathbf{r}$ sowie $\mathbf{s}$ Konjunktionsterme und $\mathbf{I}_k$ beziehungsweise $\mathbf{J}_k$ Mengen von Konjunktionstermen mit k Konjunkten beschreiben. Das heisst, $\mathbf{r}$ könnte etwa die aus Literalen bestehende Konjunktion X_1X_2 oder $X_1\overline{X_2}$ repräsentieren und diese beiden Konjunktionsterme wären dann beispielsweise zwei Elemente der Menge $\mathbf{I}_2$.

Algorithmus 5.2 Suche nach Primimplikanten von f_W gemäss csQCA (QMC-Verfahren)

Input: Menge $\mathbf{I}_n$ der Minterme von f_W^+ , Menge $\mathbf{M}$ der Minterme von f_W

Output: Menge $\mathbf{PI}(f_W)$ der Primimplikanten von f_W

```

1: repeat for  $k = n, n - 1, \dots$ 
2:    $\mathbf{I}_{k-1} := \emptyset$ 
3:    $\mathbf{J}_k := \emptyset$ 
4:   for each Element  $\mathbf{r}$  aus  $\mathbf{I}_k$  do
5:     if  $\nexists X_j$  aus  $\mathbf{r} = \mathbf{s}X_j$ , sodass  $\mathbf{s}X_j$  und  $\mathbf{s}\overline{X_j}$  in  $\mathbf{I}_k$  enthalten sind then
6:       Füge  $\mathbf{r}$  zu  $\mathbf{J}_k$  hinzu
7:     else
8:       while  $\exists X_j$  aus  $\mathbf{r} = \mathbf{s}X_j$ , wobei  $\mathbf{s}X_j$  und  $\mathbf{s}\overline{X_j}$  in  $\mathbf{I}_k$  enthalten sind do
9:         Füge  $\mathbf{s}$  zu  $\mathbf{I}_{k-1}$  hinzu
10:    Entferne alle identischen Elemente aus  $\mathbf{I}_{k-1}$  bis auf eines
11:  until  $\mathbf{I}_{k-1} = \emptyset$ 
12:  Bilde  $\mathbf{PI}(f_W^+) = \mathbf{J}_n \cup \mathbf{J}_{n-1} \cup \dots$ 
13:  for each Element  $\mathbf{r}$  aus  $\mathbf{PI}(f_W^+)$  do
14:    if  $\mathbf{r}$  deckt ein Minterm aus  $\mathbf{M}$  ab then
15:      Füge  $\mathbf{r}$  zu  $\mathbf{PI}(f_W)$  hinzu
16: return  $\mathbf{PI}(f_W)$ 

```

(b) Suche nach den RDN-Formen

Die aus Teil (a) gefolgerten Primimplikanten einer Instanzfunktion f_W entsprechen den Konjunktionstermen, die in einem zweiten Teil so disjunktiv miteinander verknüpft werden müssen, dass der gesamte Boolesche Ausdruck Bedingung (b) der Definition 4.15 erfüllt. Dazu wird bei QMC standardmässig die sogenannte *Primimplikantentabelle* erstellt, bei der die Primimplikanten der Instanzfunktion den Mintermen dieser Funktion gegenübergestellt werden, um die irredundanten Abdeckungen der Instanzfunktion zu ermitteln.

Ein Primimplikant deckt genau dann einen Minterm ab, wenn immer wenn der Minterm gegeben ist auch der Primimplikant gegeben ist. Die Minterme $\mathbf{m}_j$, $1 \leq j \leq n$, einer Instanzfunktion f in KDNF sind ihrerseits Implikanten von f für die gilt, dass $\mathbf{m}_1 + \mathbf{m}_2 + \dots + \mathbf{m}_n$ eine Abdeckung von f ist. Dementsprechend bilden Primimplikanten $\mathbf{p}_i$, $1 \leq i \leq m$, der Instanzfunktion f genau dann eine irredundante Abdeckung $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_m$ von f , wenn (1) $\mathbf{m}_1 + \mathbf{m}_2 + \dots + \mathbf{m}_n$ gegeben ist, die Disjunktion $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_m$ gegeben ist und (2) kein $\mathbf{p}_i$ aus $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_m$ derart entfernt werden kann, dass (1) weiterhin gilt. Anders formuliert bildet $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_m$ genau dann eine irredundante Abdeckung von f , wenn es zu jedem dieser Primimplikanten $\mathbf{p}_i$ mindestens einen Minterm $\mathbf{m}_j$ gibt, der relativ zu den restlichen Primimplikanten aus $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_m$ nur von $\mathbf{p}_i$ abgedeckt wird. Dabei deckt $\mathbf{p}_i$ den Minterm $\mathbf{m}_j$ genau dann ab, wenn $\mathbf{p}_i$ in $\mathbf{m}_j$ enthalten ist (vgl. S. 152).

Die Primimplikantentabelle für die in Tabelle 5.6 aufgelisteten Minterme von g_E ist in Tabelle 5.10 dargestellt. Das x-Symbol soll ausdrücken, dass der Minterm der entsprechenden Spalte durch den Primimplikanten der entsprechenden Zeile abgedeckt wird. Deckt ein Primimplikant einen Minterm nicht ab, bleibt die entsprechende Zelle leer.

Durch die tabellarische Darstellung ist relativ einfach nachvollziehbar, was für die Bildung von irredundanten Abdeckungen genau gesucht wird: Gesucht werden Kombinationen von Primim-

	$\mathbf{m}_0$	$\mathbf{m}_1$	$\mathbf{m}_2$	$\mathbf{m}_3$	$\mathbf{m}_8$	$\mathbf{m}_9$	$\mathbf{m}_{10}$	$\mathbf{m}_{12}$	$\mathbf{m}_{13}$	$\mathbf{m}_{14}$	$\mathbf{m}_{15}$
AB								x	x	x	x
$A\bar{D}$					x		x	x		x	
$A\bar{C}$					x	x		x	x		
$\bar{B}\bar{D}$	x	x	x		x		x				
$\bar{B}\bar{C}$	x	x			x	x					
$\bar{A}\bar{B}$	x	x	x	x							

Tabelle 5.10.: Primimplikantentabelle der inneren Instanzfunktion g_E .

plikanten, sodass bei jedem Minterm ein x-Symbol steht und jeder Primimplikant dieser Kombination bei mindestens einem Minterm als einziger ein x-Symbol aufweist. Verknüpft man die Primimplikanten einer solchen Kombination disjunktiv, liegt eine irredundante Abdeckung der Instanzfunktion vor und damit insgesamt ein Ausdruck in RDN-Form.

Es fällt auf, dass im Falle der Primimplikantentabelle 5.10 sowohl AB als auch $\bar{A}\bar{B}$ Teil jeder RDN-Form sein müssen, da es jeweils einen Minterm gibt, der nur durch diese beiden abgedeckt werden kann (vgl. die Minterme $\mathbf{m}_3$ und $\mathbf{m}_{15}$). Diese Primimplikanten sind Kernprimimplikanten (vgl. Def. 4.9, S. 152). Kernprimimplikanten einer inneren Instanzfunktion f_W entsprechen minimal hinreichenden Bedingungen von W , die in jeder RDN-Bedingung von W enthalten sind. Sofern also aus Konfigurationen einfache Minimale Theorien abgeleitet und damit Kausalmodelle gebildet werden können, kann von solchen minimal hinreichenden Bedingungen gesagt werden, dass sie in jedem Kausalmodell der untersuchten Konfigurationstabelle als Ursache von W auftreten. Zur Verdeutlichung wird nochmals der Minterm $\mathbf{m}_{15}$ aus der Primimplikantentabelle 5.10 betrachtet. Mit $ABCD$ als Minterm der Instanzfunktion g_E ist bekannt, dass die Konjunktion $ABCD$ hinreichend ist für E . Weiter hat sich mittels wiederholter Anwendung der Resolutionsregel gezeigt, dass sich alle hinreichenden Bedingungen zu den minimal hinreichenden Bedingungen AB , $\bar{A}\bar{D}$, $\bar{A}\bar{C}$, $\bar{B}\bar{D}$, $\bar{B}\bar{C}$ sowie $\bar{A}\bar{B}$ reduzieren lassen. Da bis auf die minimal hinreichende Bedingung AB jede andere minimal hinreichende Bedingung mindestens ein negatives Literal beinhaltet, kann für die Instantiierung von E in der Konfiguration aus $\mathcal{K}_{5,1}$, bei der alle Faktoren anwesend sind, nur die minimal hinreichende Bedingung AB verantwortlich gemacht werden. Eine minimal notwendige Disjunktion von minimal hinreichenden Bedingungen muss somit AB als Disjunkt enthalten, da eine Disjunktion ohne diese minimal hinreichende Bedingung nicht notwendig für E und damit keine einfache Minimale Theorie von E bilden kann. Im Gegensatz dazu kann beispielsweise im Zusammenhang mit der für E hinreichenden Bedingung $\bar{A}\bar{B}\bar{C}\bar{D}$ (vgl. Minterm $\mathbf{m}_9$) ohne zusätzliches Wissen nicht gesagt werden, ob nun $\bar{A}\bar{C}$ oder $\bar{B}\bar{C}$ die minimal hinreichende Bedingung ist, die für die Suffizienz von $\bar{A}\bar{B}\bar{C}\bar{D}$ für E kausal verantwortlich ist. Insofern sind solche minimal hinreichenden Bedingungen eines Literals W , die Kernprimimplikanten der inneren Instanzfunktion f_W darstellen, besonders interessant. Aus Sicht der reinen Ausführung muss man zur Bestimmung solcher minimal hinreichenden Bedingungen einer Wirkung W , die Teil jeder RDN-Bedingung von W sind, lediglich überprüfen, ob innerhalb der Primimplikantentabelle von f_W Spalten von Mintermen existieren, die nur von einem Primimplikanten und damit von einem Kernprimimplikanten abgedeckt werden (vgl. bspw. Scarbata (2001), S. 91 f.).

Zur systematischen Ermittlung einer irredundanten Abdeckung werden gemäss dem QMC-Verfahren zuerst diese Kernprimimplikanten bestimmt. Enthält eine Primimplikantentabelle einen Kernprimimplikanten, ist bekannt, dass dieser Term ein Teil von jedem Ausdruck in RDN-Form

	$\mathbf{m}_8$	$\mathbf{m}_9$	$\mathbf{m}_{10}$
$A\bar{D}$	x		x
$A\bar{C}$	x	x	
$\bar{B}\bar{D}$	x		x
$\bar{B}\bar{C}$	x	x	

(a)

	$\mathbf{m}_9$	$\mathbf{m}_{10}$
$A\bar{D}$		x
$A\bar{C}$	x	

(b)

Tabelle 5.11.: Reduzierte Primimplikantentabellen der inneren Instanzfunktion g_E für die Suche nach einer irredundanten Abdeckung.

ist. Dementsprechend kann jede Mintermspalte, die sowohl von diesem Kernprimimplikanten als auch von weiteren Primimplikanten abgedeckt wird, für die Kompletzierung einer Abdeckung aus der Tabelle entfernt werden. Der Grund für die Streichung liegt darin, dass die anderen Primimplikanten, die zusammen mit dem Kernprimimplikanten einen Minterm abdecken, bezüglich dieses Minterms in jeder RDN-Form redundant sind, da bereits der unverzichtbare Kernprimimplikant die Aufgabe des Überdeckers übernimmt. Als Beispiel wird nochmals die Primimplikantentabelle 5.10 mit den Kernprimimplikanten AB und $\bar{A}\bar{B}$ betrachtet. Der Kernprimimplikant AB deckt neben $\mathbf{m}_{15}$ auch die Minterme $\mathbf{m}_{12}$, $\mathbf{m}_{13}$ und $\mathbf{m}_{14}$ ab. Aus diesem Grund kann beispielsweise $A\bar{D}$ kein notwendiger Term für die Abdeckung von $\mathbf{m}_{14}$ sein. Dasselbe gilt für $A\bar{D}$ in $\mathbf{m}_{12}$ beziehungsweise für $A\bar{C}$ in $\mathbf{m}_{12}$ und $\mathbf{m}_{13}$, oder allgemein formuliert für alle Primimplikanten, die zusammen mit einem Kernprimimplikanten eine Mintermspalte abdecken. Auf analoge Weise können aufgrund des Kernprimimplikanten $\bar{A}\bar{B}$ die drei Spalten $\mathbf{m}_0$, $\mathbf{m}_1$ und $\mathbf{m}_2$ entfernt werden. Insgesamt resultiert daraus die vereinfachte Primimplikantentabelle 5.11a, die nur noch die Spalten mit jenen Mintermen aufweist, die von AB und $\bar{A}\bar{B}$ nicht abgedeckt werden. Auch die Zeilen wurden entsprechend soweit reduziert, dass nur noch jene Primimplikanten aufgelistet sind, die diese noch nicht abgedeckten Minterme abdecken.

Über die sogenannte *Dominanzprüfung* wird die reduzierte Primimplikantentabelle anschließend weiter schrittweise vereinfacht und die Kernprimimplikanten, sofern es Kernprimimplikanten gibt und sie disjunktiv verknüpft noch keine Abdeckung der inneren Instanzfunktion bilden, parallel dazu um weitere Primimplikanten ergänzt, bis alle Minterme durch Primimplikanten abgedeckt sind. Mit dieser Methode kann eine Primimplikantentabelle unter Umständen stark vereinfacht und damit die Suche nach einer irredundanten Abdeckung effizienter gestaltet werden. Konkret werden die folgenden beiden Regeln auf die (reduzierte) Primimplikantentabelle angewendet (vgl. Nelson et al. (1995), S. 215 ff.):

1. Dominiert eine Zeile i eine andere Zeile j , kann letztere aus der Primimplikantentabelle gestrichen werden. Dabei dominiert eine Zeile i eine andere Zeile j genau dann, wenn

erstere in mindestens all jenen Spalten ein x-Symbol aufweist, in denen letztere ein x-Symbol hat. Gibt es Zeilen, die sich gegenseitig dominieren, können diese Zeilen alle bis auf eine gestrichen werden.

2. Dominiert eine Spalte k eine andere Spalte l , kann erstere aus der Primimplikantentabelle gestrichen werden. Dabei dominiert eine Spalte k eine andere Spalte l genau dann, wenn erstere mindestens in all jenen Zeilen ein x-Symbol aufweist, in denen letztere ein x-Symbol hat. Gibt es Spalten, die sich gegenseitig dominieren, können diese Spalten alle bis auf eine gestrichen werden.

Die Idee hinter der ersten Regel ist, dass wenn der Primimplikant $\mathbf{p}_i$ aus Zeile i in eine Abdeckung aufgenommen wird, $\mathbf{p}_i$ auch gleich alle Minterme abdeckt, die vom Primimplikanten $\mathbf{p}_j$ aus Zeile j abgedeckt werden. Konkret anhand der Primimplikantentabelle 5.12 veranschaulicht, dominiert Zeile 3 Zeile 1, weshalb Zeile 1 gemäss der ersten Regel gestrichen wird. Der Primimplikant $\mathbf{p}_3$ deckt neben den Mintermen, die vom Primimplikanten $\mathbf{p}_1$ abgedeckt werden, zusätzlich auch noch den Minterm $\mathbf{m}_3$ ab. $\mathbf{p}_3$ deckt also mindestens die Minterme ab, die von $\mathbf{p}_1$ abgedeckt werden, weshalb in diesem Beispiel $\mathbf{p}_1$ zugunsten von $\mathbf{p}_3$ vernachlässigt werden kann. Weil Zeile 3 auch Zeile 2 dominiert, kann analog für die Streichung von Zeile 2 argumentiert werden. Die zweite Regel folgt aus der Überlegung, dass wenn man einen Primimplikanten hat, der den Minterm $\mathbf{m}_l$ aus Spalte l abdeckt, auch gleich einen Primimplikanten hat, der den Minterm $\mathbf{m}_k$ aus Spalte k abdeckt. Es soll darauf hingewiesen werden, dass es sich damit bezüglich der Dominanz und Streichung bei den Spalten genau umgekehrt verhält wie bei den Zeilen: Dominiert eine Zeile eine andere, wird letztere gestrichen. Dominiert eine Spalte eine andere, wird erstere gestrichen. Wiederum konkret anhand der Primimplikantentabelle 5.12 veranschaulicht, dominiert Spalte 1 Spalte 2, weshalb gemäss der zweiten Regel Spalte 1 gestrichen wird. Denn wählt man einen der Primimplikanten $\mathbf{p}_1$ oder $\mathbf{p}_3$, die den Minterm $\mathbf{m}_2$ abdecken, hat man auch gleich einen Primimplikanten, der $\mathbf{m}_1$ abdeckt. Analoges gilt zwischen Spalte 1 und der von Spalte 1 dominierten Spalte 3.

		1	2	3
		$\mathbf{m}_1$	$\mathbf{m}_2$	$\mathbf{m}_3$
1	$\mathbf{p}_1$	x	x	
2	$\mathbf{p}_2$	x		x
3	$\mathbf{p}_3$	x	x	x

Tabelle 5.12.: Primimplikantentabelle, bei der gemäss der Dominanzprüfung die Zeilen 1 und 2 sowie die Spalte 1 gestrichen werden können.

Wendet man diese beiden Regeln auf die obige reduzierte Primimplikantentabelle 5.11a an, sieht man, dass die erste Zeile die dritte Zeile dominiert (und umgekehrt) und die zweite Zeile die vierte Zeile dominiert (und umgekehrt). Gemäss der ersten Regel können demnach beispielsweise die dritte und die vierte Zeile gestrichen werden.⁷ Weiter zeigt sich, dass die erste Spalte die zweite und dritte Spalte dominiert, weshalb gemäss der zweiten Regel auch diese erste Spalte mit dem Minterm m_8 gestrichen werden kann. Aus der Primimplikantentabelle 5.11a resultiert über diese Dominanzprüfung die wesentlich vereinfachte Primimplikantentabelle 5.11b. Von dieser reduzierten Primimplikantentabelle werden nun erneut jene Primimplikanten gewählt, die relativ zu dieser Tabelle als einzige einen Minterm abdecken. Die Zeilen der ausgewählten Primimplikanten sowie die Spalten mit jenen Mintermen, die durch diese Primimplikanten abgedeckt werden, werden erneut aus der reduzierten Primimplikantentabelle entfernt. Dieses Reduktionsverfahren einer Primimplikantentabelle mittels Primimplikantenwahl sowie der Anwendung der Regeln 1 und 2 wird iterativ so lange wiederholt, bis nach der Streichung der Zeilen und Spalten eine leere Primimplikantentabelle zurückbleibt. Sobald dies der Fall ist, liegt mit der disjunktiven Verknüpfung der bis dahin gewählten Primimplikanten eine irredundante Abdeckung vor.

Eine Primimplikantentabelle, ob bereits vereinfacht oder nicht, kann auch *zyklisch* sein (vgl. Nelson et al. (1995), S. 217). Das bedeutet, dass diese Tabelle weder Primimplikanten enthält, die als einzige einen Minterm abdecken, noch die obigen beiden Regeln angewendet werden können. In diesem Fall wird ein beliebiger Primimplikant gewählt und die entsprechende Zeile sowie die dazugehörigen Spalten, bei denen die Zeile ein x-Symbol aufweist, gestrichen. Nach McCluskey fällt dabei die Wahl standardmässig auf einen Primimplikanten, der maximal viele Minterme abdeckt (vgl. McCluskey (1956), S. 1426). Die resultierende reduzierte Primimplikantentabelle kann wieder zyklisch sein, womit erneut eine beliebige Wahl getroffen wird. Andernfalls werden von Neuem die obigen beiden Regeln angewendet und jene Primimplikanten gewählt, die als einzige einen Minterm abdecken.

Angewendet auf die Primimplikantentabelle 5.11b zeigt sich, dass mit $A\bar{D}$ und $A\bar{C}$ zwei weitere Primimplikanten vorliegen, die jeweils als einzige einen Minterm abdecken. Die beiden Primimplikanten werden also gewählt, die entsprechenden Zeilen gestrichen und alle Spalten entfernt, die bei diesen Zeilen ein x-Symbol aufweisen. Zumal aus dieser erneuten Streichung eine leere Primimplikantentabelle resultiert, kann das Verfahren abgebrochen werden. Mit der disjunktiven Verknüpfung dieser beiden letzten Primimplikanten $A\bar{D}$ und $A\bar{C}$ zusammen mit den zu Beginn bestimmten Kernprimimplikanten liegt eine irredundante Abdeckung $AB + \bar{A}\bar{B} + A\bar{D} + A\bar{C}$ vor.

⁷Es könnten auch andere gestrichen werden, wie beispielsweise die erste und zweite Zeile.

Aus dem Teil des Verfahrens nach Quine und McCluskey, der nach aus Primimplikanten bestehenden irredundanten Abdeckungen sucht, folgt standardmässig nur genau eine irredundante Abdeckung, die eine disjunktive Minimalform darstellt. Allenfalls werden bei QMC etwa mit heuristischen Methoden noch weitere disjunktive Minimalformen gesucht.⁸ Im Zusammenhang mit der Suche nach nur einer Lösung zeigt etwa gerade die obige am Beispiel von g_E veranschaulichte iterative Verwendung der Dominanzprüfung, dass damit zwar eine korrekte irredundante Abdeckung von g_E gefunden wurde. Es gibt aber bekanntlich noch drei weitere, unter denen sich unter anderem auch die tatsächlich gesuchte Struktur befindet (vgl. (5.20), (5.22) und (5.23), S. 186). Auch eine erweiterte Suche nach allen irredundanten Abdeckungen, die ebenfalls disjunktive Minimalformen darstellen, greift zu kurz. Im Fall von g_E würde man damit zwar tatsächlich die vier irredundanten Abdeckungen finden, zumal es sich bei allen vier auch um disjunktive Minimalformen handelt. Es hat sich jedoch bereits im Zusammenhang mit der Definition 4.15 einer Instanzfunktion in RDNF gezeigt, dass im kausaltheorietischen Kontext teils auch eine irredundante Abdeckung die Ursachen einer Wirkung einfangen kann, die nicht jener Disjunktion von Konjunktionstermen entspricht, die minimal viele Disjunkte aufweist und in deren Konjunktionsterme minimal viele Literale enthalten sind (vgl. S. 168). Alleine basierend auf den Konfigurationen kann die Berücksichtigung solcher rein quantitativen Eigenschaften nicht die Rechtfertigung dafür liefern, weshalb man eine sich allenfalls aus einer irredundanten Abdeckung ergebende einfache Minimale Theorie einer anderen einfachen Minimalen Theorie desselben Literals für die Bildung einer Kausalhypothese vorziehen sollte. Folglich müssen beim konfiguralen kausalen Modellieren jeweils alle irredundanten Abdeckungen gesucht werden, zumal bei jeder dieser Abdeckungen die Möglichkeit besteht, dass sich daraus ein Kausalmodell ergibt, das der tatsächlichen Kausalstruktur entspricht.

Eine iterative Dominanzprüfung ist somit im hier vorliegenden Kontext des konfiguralen kausalen Modellierens wenig sinnvoll. Tatsächlich wird aber bei QCA oftmals genau diese Standardvariante von QMC mit iterativer Dominanzprüfung verwendet.⁹ Damit läuft man mit QCA Gefahr, auch bei idealen Daten einen kausalen Fehlschluss zu produzieren, indem die sich aus

⁸Solche heuristischen Methoden helfen dabei, mit dem aus der Komplexitätstheorie bekannten Überdeckungsproblem umzugehen, das zur Klasse der NP-vollständigen Probleme gehört (vgl. u. a. Karp (1972)). Denn gemäss der Komplexitätstheorie ist die Suche nach einer irredundanten Abdeckung von Mintermen durch Primimplikanten algorithmisch (vermutlich) nicht effizient lösbar. Konkret würde das etwa bedeuten, dass u. a. der zeitliche Rechenaufwand eines Algorithmus, der jeweils exakt alle irredundanten Abdeckungen sucht, mit zunehmender Anzahl an Mintermen und Primimplikanten extrem schnell sowie stark ansteigen würde und zwar unabhängig davon, welche Rechenschritte genau ausgeführt werden (vgl. u. a. Garey & Johnson (1979)). Heuristische Methoden können dabei helfen, den Rechenaufwand zu reduzieren.

⁹Die meistverwendete Software dürfte die von Charles Ragin und Sean Davey entwickelte „fs/QCA“ sein (vgl. Ragin & Davey (2009)), deren Marktanteil in den vergangenen Jahren bei über 80% lag (vgl. Thiem & Dusa (2013), S. 506). Führt man innerhalb dieses Programms eine Standardanalyse für den crisp-set Fall aus, kommt obiges standard QMC-Verfahren zum Einsatz.

einer Analyse ergebenden Kausalhypothesen jenes Kausalmodell, das die tatsächliche Kausalstruktur einfängt, nicht enthalten. Um im Rahmen des konfiguralen kausalen Modellierens die Korrektheit des Aufdeckungsverfahrens zu wahren, muss QMC derart angepasst werden, dass systematisch alle irredundanten Abdeckungen gesucht werden. Eine Methode, die sich für die Suche nach allen irredundanten Abdeckungen besonders anbietet, ist das *Petrick-Verfahren* (vgl. u.a. Nelson et al. (1995), S. 222 ff.; Roth & Kinney (2010), S. 171 ff.). Obwohl das Petrick-Verfahren strenggenommen nicht als integraler Teil des QMC-Verfahrens angesehen wird, wird es doch oft in einem Atemzug mit QMC für die systematische Suche nach allen Abdeckungen genannt. Da das Petrick-Verfahren weiter wesentliche Teile, namentlich die Primimplikantentabelle, des QMC-Verfahrens nutzt und nicht zuletzt auch bei McCluskey selbst ein in diese Richtung gehender Ansatz zu finden ist (vgl. McCluskey (1956), S. 1437 ff.), wird es in der vorliegenden Arbeit als erweiterter Teil des QMC-Verfahrens angesehen und entsprechend auch letzterem zugeordnet. Es soll jedoch nochmals erwähnt sein, dass QCA typischerweise mit dem oben beschriebenen Standard-Quine-McCluskey-Verfahren mit nur einer resultierenden Abdeckung (in disjunktiver Minimalform) arbeitet, weshalb das Petrick-Verfahren zumindest nicht als charakteristischer Standard von csQCA angesehen werden kann. Da das Petrick-Verfahren jedoch eng mit dem QMC-Verfahren verbandelt ist, das seinerseits der Kern von csQCA darstellt, wird aus Gründen einer übersichtlichen Kontrastierung der Aufdeckungsverfahren csQCA und csCNA das Petrick-Verfahren bei csQCA angesiedelt.

Das Verfahren von Petrick nutzt ausgehend von der Primimplikantentabelle einen algebraischen Ansatz zur Bestimmung aller irredundanten Abdeckungen. Dazu werden wie beim QMC-Verfahren zuerst die Kernprimimplikanten bestimmt und die Primimplikantentabelle entsprechend reduziert, sodass nur noch die Spalten jener Minterme aufgeführt sind, die von den Kernprimimplikanten nicht abgedeckt werden, und nur noch die Primimplikanten-Zeilen vorhanden sind, die bei mindestens einer dieser Minterm-Spalten ein x-Symbol aufweisen. Weiter werden die übrigbleibenden n Primimplikanten mit einem Label, wie etwa P_i , $1 \leq i \leq n$, versehen. Dieses Label ist deshalb wichtig, weil ein entsprechendes P_i in der Folge als Boolesche Variable interpretiert wird und die damit gebildeten Ausdrücke so für die Booleschen Gesetze zugänglich gemacht werden. Dies ist auch der Grund, weshalb hier für einen Primimplikanten nicht der bereits teils verwendete fette Kleinbuchstabe $\mathbf{p}_i$ benutzt wird. Während letzterer ein Kurzbezeichner für einen beliebigen Primimplikanten darstellt, der lediglich für eine bessere Lesbarkeit sorgen soll, soll mit dem Label P_i angedeutet werden, dass tatsächlich mit diesem Label als Variable operiert wird und es sich dabei um einen bestimmten Primimplikanten handelt.

Ausgehend von der Primimplikantentabelle 5.10 ergibt sich für die innere Instanzfunktion g_E nach der Bestimmung der Kernprimimplikanten die bereits bekannte reduzierte Primimplikan-

		m₈	m₉	m₁₀
P_1	$A\bar{D}$	x		x
P_2	$A\bar{C}$	x	x	
P_3	$\bar{B}\bar{D}$	x		x
P_4	$\bar{B}\bar{C}$	x	x	

Tabelle 5.13.: Reduzierte Primimplikantentabelle der inneren Instanzfunktion g_E mit gelabelten Primimplikanten.

tentabelle 5.13, wobei neu die übriggebliebenen Primimplikanten mit einem Label markiert werden.

Anhand dieser Primimplikantentabelle 5.13 veranschaulicht, wird mithilfe der gelabelten Primimplikanten nun folgendermassen ein Term gebildet, der sich wie ein Boolescher Ausdruck behandeln lässt: Tabelle 5.13 zeigt, dass zur Bildung einer irredundanten Abdeckung von g_E die bereits bestimmten Kernprimimplikanten noch derart um Primimplikanten ergänzt werden müssen, dass auch die Minterme **m₈**, **m₉** und **m₁₀** abgedeckt sind. Der Minterm **m₈** ist genau dann abgedeckt, wenn P_1 oder P_2 oder P_3 oder P_4 in die Abdeckung aufgenommen werden. Mit der Aufnahme von P_2 oder P_4 in die Abdeckung gilt dasselbe für den Minterm **m₉**. Schliesslich ist der Minterm **m₁₀** genau dann abgedeckt, wenn P_1 oder P_3 in der Abdeckung enthalten sind. Dieser Umstand lässt sich nun mit einem Booleschen Ausdruck repräsentieren. So ist **m₈** genau dann abgedeckt, wenn $(P_1 + P_2 + P_3 + P_4)$ wahr ist, **m₉** genau dann abgedeckt, wenn $(P_2 + P_4)$ wahr ist und **m₁₀** genau dann abgedeckt, wenn $(P_1 + P_3)$ wahr ist. Insgesamt sind demnach alle von den Kernprimimplikanten noch nicht abgedeckten Minterme genau dann abgedeckt, wenn $(P_1 + P_2 + P_3 + P_4)$ und $(P_2 + P_4)$ und $(P_1 + P_3)$ wahr sind beziehungsweise wenn

$$(P_1 + P_2 + P_3 + P_4) \cdot (P_2 + P_4) \cdot (P_1 + P_3) \quad (5.24)$$

wahr ist.

Der Ausdruck (5.24) wird nun mittels der Booleschen Gesetze in disjunktive Normalform gebracht und minimiert. Konkret wird (5.24) durch die Distributivgesetze ausmultipliziert und via den Idempotenzgesetzen ($X \cdot X = X$ bzw. $X + X = X$) und Absorbtionsgesetzen ($X \cdot (X + Y) = X$ bzw. $X + XY = X$) so weit wie möglich vereinfacht, bis letztere nicht mehr angewendet werden können. Dieser Vorgang ist deshalb relativ einfach und führt zu einer eindeutigen Lösung, weil die gebildeten Booleschen Ausdrücke am Ausgangspunkt der Umformungen keine Komplemente enthalten. Für (5.24) resultiert über dieses Ausmultiplizieren und Minimieren der Ausdruck (5.26):

$$(P_1 + P_2 + P_3 + P_4) \cdot (P_2 + P_4) \cdot (P_1 + P_3) = (P_2 + P_4) \cdot (P_1 + P_3) \quad (5.25)$$

$$= P_1P_2 + P_2P_3 + P_1P_4 + P_3P_4 \quad (5.26)$$

Der Boolesche Ausdruck (5.26) ist äquivalent zu (5.24) und bringt zum Ausdruck, dass die Minterme $\mathbf{m}_8$, $\mathbf{m}_9$ und $\mathbf{m}_{10}$ genau dann abgedeckt sind, wenn die Primimplikanten P_1 und P_2 oder P_2 und P_3 oder P_1 und P_4 oder P_3 und P_4 in die Abdeckung aufgenommen werden. Ersetzt man das jeweilige Label durch den entsprechenden Primimplikanten und verknüpft diese disjunktiv mit den bereits bestimmten Kernprimimplikanten, ergeben sich daraus alle vier irredundanten Abdeckungen beziehungsweise RDN-Formen der inneren Instanzfunktion g_E :

$$g_E(\cdot) = AB + \bar{A}\bar{B} + A\bar{D} + A\bar{C} \quad (5.27)$$

$$g_E(\cdot) = AB + \bar{A}\bar{B} + A\bar{C} + \bar{B}\bar{D} \quad (5.28)$$

$$g_E(\cdot) = AB + \bar{A}\bar{B} + A\bar{D} + \bar{B}\bar{C} \quad (5.29)$$

$$g_E(\cdot) = AB + \bar{A}\bar{B} + \bar{B}\bar{D} + \bar{B}\bar{C} \quad (5.30)$$

Vor dem Hintergrund des Zusammenhangs, dass eine RDN-Form einer inneren Instanzfunktion f_W eine RDN-Bedingung von W ist, ergeben sich aus den vier obigen RDN-Formen direkt die vier RDN-Bikonditionale von E , die aus $\mathcal{K}_{5,1}$ folgen:

$$AB + \bar{A}\bar{B} + A\bar{D} + A\bar{C} \leftrightarrow E \quad (5.31)$$

$$AB + \bar{A}\bar{B} + A\bar{C} + \bar{B}\bar{D} \leftrightarrow E \quad (5.32)$$

$$AB + \bar{A}\bar{B} + A\bar{D} + \bar{B}\bar{C} \leftrightarrow E \quad (5.33)$$

$$AB + \bar{A}\bar{B} + \bar{B}\bar{D} + \bar{B}\bar{C} \leftrightarrow E \quad (5.34)$$

Das Verfahren für die Suche nach irredundanten Abdeckungen mittels Primimplikantentabellen und Petrick-Verfahren bleibt dasselbe für partielle Instanzfunktionen. Für h_U beispielsweise ergibt sich die Primimplikantentabelle 5.14. Analog kann die Primimplikantentabelle von h_V gebildet werden, sodass sich insgesamt für h_U und h_V folgende RDN-Formen ergeben:

$$h_U(\cdot) = S + T \quad (5.35)$$

$$h_V(\cdot) = R + S \quad (5.36)$$

Im Gegensatz zu den RDN-Formen von g_E bestehen die RDN-Formen von h_U und h_V ausschliesslich aus Kernprimimplikanten. Die Primimplikantentabelle 5.14 etwa zeigt, dass es sowohl für S als auch für T Minterme gibt, die nur von S (vgl. $\mathbf{m}_5$ und $\mathbf{m}_{13}$ aus Tabelle 5.14) und T (vgl. $\mathbf{m}_2$ und $\mathbf{m}_{11}$) abgedeckt werden. Analoges gilt für h_V . Insgesamt beschränkt sich

	$\mathbf{m}_2$	$\mathbf{m}_5$	$\mathbf{m}_7$	$\mathbf{m}_{11}$	$\mathbf{m}_{13}$	$\mathbf{m}_{15}$
S		x	x		x	x
T	x		x	x		x

Tabelle 5.14.: Primimplikantentabelle der inneren Instanzfunktion h_U .

das Petrick-Verfahren in diesem Fall also auf den ersten Teilschritt der Bestimmung der Kernprimimplikanten und der zweite Teilschritt der Bildung und Umformung eines Booleschen Ausdrucks kommt nicht zur Anwendung. Aus den zwei obigen RDN-Formen ergeben sich die folgenden zwei RDN-Bikonditionale von U und V :

$$S + T \leftrightarrow U \quad (5.37)$$

$$R + S \leftrightarrow V \quad (5.38)$$

Abschliessend soll die operationale Umsetzung der obigen Vorgehensweise für Teil (b) erneut mittels Pseudocode zusammengefasst werden. Algorithmus 5.3 zeigt eine mögliche Beschreibung für die Suche nach allen irredundanten Abdeckungen einer Instanzfunktion gemäss dem Petrick-Verfahren. Dabei seien δ und ρ Boolesche Ausdrücke. Des Weiteren sind $\mathbf{p}_i$ beziehungsweise $\mathbf{m}_j$, $i, j \in \mathbb{N}$ – wie in dieser Arbeit üblicherweise verwendet – Kurzbezeichner für Minterme beziehungsweise Primimplikanten. Im Gegensatz dazu ist P_i ein Label, das zwar ebenfalls den Primimplikanten $\mathbf{p}_i$ repräsentiert, aber zusätzlich auch für sich als Variable eingesetzt wird.

Algorithmus 5.3 Suche nach RDN-Formen von f_W gemäss dem Petrick-Verfahren

Input: Menge $\mathbf{M}$ der Minterme von f_W , Menge $\mathbf{PI}(f_W)$ der Primimplikanten von f_W

Output: Menge $\mathbf{A}(f_W)$ aller RDN-Formen von f_W

- 1: Bilde die Disjunktion δ^* aller Kernprimimplikanten aus $\mathbf{PI}(f_W)$
 - 2: **if** δ^* deckt alle Minterme aus $\mathbf{M}$ ab **then**
 - 3: $\mathbf{A}(f_W) = \{\delta^*\}$
 - 4: **else**
 - 5: Bilde zu jedem Primimplikanten $\mathbf{p}_i$, $1 \leq i \leq k$, ein eindeutiges Label P_i , $1 \leq i \leq k$
 - 6: **for each** Minterm $\mathbf{m}_j$, $1 \leq j \leq l$, der nicht durch δ^* abgedeckt wird **do**
 - 7: Bilde die Disjunktion $\delta_{\mathbf{m}_j}$ aller Label der Primimplikanten, die $\mathbf{m}_j$ abdecken
 - 8: Bilde die Konjunktion $\delta := \delta_{\mathbf{m}_1} \cdots \delta_{\mathbf{m}_j} \cdots \delta_{\mathbf{m}_l}$
 - 9: Forme δ mit den Distributivgesetzen zu einer DNF um und minimiere sie zu δ' mit (a) den Idempotenz- und (b) den Absorbtionsgesetzen, bis (a) und (b) nicht mehr anwendbar sind
 - 10: **for each** Konjunktionsterm $P_1 P_2 \cdots P_m$, $1 \leq m \leq k$, aus δ' **do**
 - 11: Transformiere $P_1 P_2 \cdots P_m$ zu $\mathbf{p}_1 + \mathbf{p}_2 + \cdots + \mathbf{p}_m$ und füge $\rho := \delta^* + \mathbf{p}_1 + \mathbf{p}_2 + \cdots + \mathbf{p}_m$ zu $\mathbf{A}(f_W)$ hinzu
 - 12: **return** $\mathbf{A}(f_W)$
-

5.2.3. csCNA und das Baumgartner-Verfahren

Bei der von Michael Baumgartner entwickelten Coincidence Analysis (CNA) ist das Forschungsdesign typischerweise so aufgebaut, dass relativ zur geprüften Faktormenge $\mathbf{F}$ alle Kausalmodelle beliebiger Komplexität gesucht werden, die verträglich mit der untersuchten Konfigurationstabelle $\mathcal{K}_{\mathbf{F}}$ sind (vgl. u. a. Baumgartner (2006); Baumgartner (2009); Baumgartner (2013a)). Im Gegensatz zu QCA hat CNA somit den Anspruch auch Modelle mit mehreren Wirkungen und indirekten Kausalzusammenhängen aufzudecken, wie beispielsweise Common-Cause-Strukturen oder Kausalketten, die über komplexe Minimale Theorien gebildet werden.

Ähnlich wie bei QCA wird bei CNA für die Suche nach den RDN-Bikonditionalen eines Literals W , beziehungsweise den RDN-Formen der inneren Instanzfunktion f_W , ein Minimierungsverfahren eingesetzt, das in die beiden Teile (a) der Suche nach den Primimplikanten von f_W und (b) der Suche nach den irredundanten Abdeckungen von f_W zerlegt werden kann. Im Gegensatz zu QCA bedient sich Baumgartner jedoch nicht den gängigen Verfahren aus der Booleschen Algebra, sondern entwickelt einen Algorithmus, der aus dem Begriff der RDN-Bedingung gemäss Definition 2.3 entspringt und sich direkt an den Definitionen 2.1 beziehungsweise 2.2 von minimal hinreichenden beziehungsweise minimal notwendigen Bedingungen von W orientiert (vgl. S. 33 f.). Aus Gründen einer übersichtlichen Kontrastierung soll dieses Minimierungsverfahren unter dem Namen „Verfahren nach Baumgartner“ geführt und die von Baumgartner ursprünglich genutzten Verfahrensschritte sowie die Syntax an das Verfahren aus Abschnitt 5.2.2 angepasst werden. Die grundlegenden Prinzipien, die die Arbeitsweise des csCNA-Verfahrens charakterisieren, bleiben davon unberührt.

(a) Suche nach den Primimplikanten

Gemäss der Definition 2.1 ist ein Konjunktionsterm genau dann minimal hinreichend für ein Literal W , wenn kein echter Teil dieses Terms hinreichend ist. Ein Konjunktionsterm ist dabei genau dann hinreichend für W , wenn immer wenn dieser Konjunktionsterm gegeben ist, auch W gegeben ist. Daraus kann direkt gefolgert werden, dass ein Konjunktionsterm genau dann minimal hinreichend für W ist, wenn er hinreichend ist und aus diesem Term kein Literal entfernt werden kann, ohne dass in der Konfigurationstabelle mindestens eine Konfiguration vorhanden ist, bei der zwar der um ein Literal reduzierte Konjunktionsterm gegeben ist, der Faktor W jedoch abwesend ist. Umgekehrt ist ein Literal X_i redundant und kann aus einer hinreichenden Bedingung $X_1 \cdots X_{i-1} X_i X_{i+1} \cdots X_n$ von W entfernt werden, wenn letzteres nicht gilt. Das heisst, X_i ist redundant, wenn es keine Konfigurationen gibt, bei denen $X_1 \cdots X_{i-1} X_{i+1} \cdots X_n$ gegeben ist und W nicht gegeben ist. Die daraus ableitbare Regel der Reduktion wird in der vorliegenden

Arbeit als *Redundanzregel* (RED) bezeichnet und lässt sich analog zur Resolutionsregel unter Verwendung der Konfigurationsalgebra folgendermassen mathematisch ausgedrücken¹⁰:

$$\begin{aligned} & (X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W) \cdot \overline{(X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \cdot \bar{W})} \\ & = X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W \end{aligned} \quad (\text{RED})$$

Für eine systematische Verwendung der Redundanzregel ist es erneut sinnvoll, sie mittels der inneren Instanzfunktion zu beschreiben. Demnach zeigt sich die Redundanz eines Literals X_i gemäss obigen Überlegungen derart, dass (1) $X_1 \cdots X_{i-1} X_i X_{i+1} \cdots X_n$ ein Implikant von f_W ist und (2) nicht gilt, dass $X_1 \cdots X_{i-1} X_{i+1} \cdots X_n$ Teil eines Implikanten von $\bar{f}_W$ ist. Für $\bar{f}_W$ gilt dabei, dass sie für dieselben Belegungen wie f_W definiert ist und genau dann den Funktionswert 0 beziehungsweise 1 annimmt, wenn f_W den Funktionswert 1 beziehungsweise 0 annimmt (vgl. S. 148). Im vorliegenden Kontext ist also ein Implikant von $\bar{f}_W$ nichts anderes als eine hinreichende Bedingung von $\bar{W}$. Sind (1) und (2) erfüllt kann gefolgert werden, dass $X_1 \cdots X_{i-1} X_{i+1} \cdots X_n$ ein Implikant von f_W ist. Zusammengefasst kann die Redundanzregel (RED) folgendermassen reformuliert werden:

- (1) $X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n$ ist ein Implikant von f_W
 - (2) $X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n$ ist nicht Teil eines Implikanten von $\bar{f}_W$
-
- $X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n$ ist ein Implikant von f_W

Aus (1) folgt, dass $\bar{f}_W$ für eine Belegung definiert ist, für die die Konjunktion $X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n$ den Wert 1 annimmt. Zusammen mit (2) folgt daraus, dass wenn $X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n$ gegeben ist, $\bar{f}_W$ den Wert 0 annimmt.¹¹ Weiter gilt, dass $\bar{f}_W$ genau dann den Wert 0 annimmt, wenn f_W den Wert 1 annimmt. Daraus folgt, dass wenn $X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n$ gegeben ist, f_W den Wert 1 annimmt und dementsprechend $X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n$ ein Implikant von f_W ist.

Es soll an dieser Stelle bemerkt werden, dass die Konklusion der Redundanzregel bei einer totalen Instanzfunktion f_W alleine aus Prämisse (2) gefolgert werden kann. Tatsächlich reicht diese Prämisse bei totalen Instanzfunktionen aus um Primimplikanten zu finden, da solche alleine durch den Instantiierungsbereich beziehungsweise den Non-Instantiierungsbereich der Funktion vollständig definiert sind. Das heisst, ist $\mathbf{m}_i$ ein Implikant einer beliebigen totalen Instanzfunktion f , nimmt $\bar{f}$ für die entsprechende Belegung den Wert 0 an. Das gleiche gilt umgekehrt für f bei den Implikanten von $\bar{f}$. Bei einer totalen inneren Instanzfunktion f_W könnte man demnach alleine die Minterme von $\bar{f}_W$ betrachten, daraus eine kleinstmögliche Konjunktion von Faktoren herausfiltern, die in keinem dieser Implikanten enthalten ist und direkt darauf schliessen, dass es

¹⁰Ein Nachweis der Redundanzregel findet man in Abschnitt A.2 des Anhangs (vgl. S. 250).

¹¹Die alternative Möglichkeit, wonach $\bar{f}_W$ für die Belegung, für die die Konjunktion $X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n$ wahr ist, nicht definiert ist, wird durch (1) ausgeschlossen.

sich dabei um einen Primimplikanten von f_W handelt. Sei f_C eine solche innere Instanzfunktion mit abhängigem Faktor C , die total und in KDNF von der Form $AB + A\bar{B} + \bar{A}B$ ist. Indem f_C total ist kann alleine dadurch gerechtfertigt werden, dass A und B Primimplikanten von f_C sind, weil A und B kleinste Teile sind, die nicht im einzigen Implikanten $\bar{A}\bar{B}$ von $\bar{f}_C(A, B) = \bar{A}\bar{B}$ enthalten sind. Wie sich im Verlauf des Kapitels noch zeigen wird, lassen sich die Primimplikanten einer inneren Instanzfunktionen f_W im Allgemeinen jedoch nicht alleine basierend auf der Prämisse (2) der Redundanzregel bestimmen, da die innere Instanzfunktion partiell sein kann. Damit besteht die Möglichkeit, dass eine Konjunktion $X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n$ zwar nicht Teil eines Implikants von $\bar{f}_W$ ist, f_W für die entsprechende Belegung jedoch nicht definiert ist.

Hinsichtlich der Suche nach den Primimplikanten der inneren Instanzfunktion wendet das Verfahren nach Baumgartner die Redundanzregel ausgehend von den Mintermen der Instanzfunktion f_W im Vergleich mit den Mintermen von $\bar{f}_W$ solange an, bis alle Primimplikanten ermittelt sind. Ein Primimplikant ist dann ermittelt, wenn aus diesem Term kein Faktor entfernt werden kann, ohne dass Prämisse (2) der Redundanzregel verletzt ist. Eine entsprechende operationale Umsetzung kann folgendermassen aussehen, wobei $X_1X_2 \cdots X_n$ ein Minterm von f_W ist (vgl. u. a. Baumgartner (2009), S. 84): Für jedes X_i aus dem Minterm $X_1X_2 \cdots X_n$, $1 \leq i \leq n$, und jede Permutation von $X_1X_2 \cdots X_n$ wird X_i aus dem Term entfernt und überprüft, ob $X_1X_2 \cdots X_{i-1}X_{i+1} \cdots X_n$ Teil eines Implikants von $\bar{f}_W$ ist. Ist dies der Fall, fügt man X_i wieder zu $X_1X_2 \cdots X_{i-1}X_{i+1} \cdots X_n$ hinzu und macht mit der Eliminierung von X_{i+1} weiter. Ist dies nicht der Fall, wird der Vorgang mit X_{i+1} fortgeführt ohne X_i wieder in $X_1X_2 \cdots X_{i-1}X_{i+1} \cdots X_n$ aufzunehmen. Diese Schritte werden für jeden Minterm und jede Permutation davon wiederholt.

Indem man in umgekehrter Richtung verfährt lässt sich demonstrieren, dass mit diesem Algorithmus tatsächlich alle Primimplikanten einer Instanzfunktion gefunden werden: Es wird davon ausgegangen, dass $X_1X_2 \cdots X_k$, $k \leq n$, ein beliebiger Primimplikant der Instanzfunktion f_W ist. Für den Term $X_1X_2 \cdots X_k$ gilt, dass er ein Implikant von f_W ist. Daraus folgt, dass auch die Erweiterungen von $X_1X_2 \cdots X_k$ bis hin zu $X_1X_2 \cdots X_kX_{k+1} \cdots X_n$ Implikanten von f_W sind, wobei es sich bei $X_1X_2 \cdots X_kX_{k+1} \cdots X_n$ um einen Minterm von f_W handelt. Da der obige Algorithmus auf alle Minterme von f_W und alle Permutationen davon die Redundanzregel wiederholt anwendet, wird unter anderem auch der Konjunktionsterm $X_{k+1} \cdots X_nX_1X_2 \cdots X_k$ analysiert. Mit X_{k+1} beginnend und von links nach rechts abarbeitend wird über die wiederholte Anwendung der Redundanzregel daraus der Primterm $X_1X_2 \cdots X_k$ resultieren.

Die systematische Vorgehensweise soll kurz anhand der Ermittlung des RDN-Bikonditionals $A + B \leftrightarrow C$ veranschaulicht werden, das eine (einfache) Minimale Theorie von $\mathcal{K}_{5,15}$ darstellt. Gesucht wird die RDN-Form der inneren Instanzfunktion f_C mit der Wertetabelle 5.15b. Aus der Wertetabelle ergibt sich einerseits, dass die Instanzfunktion f_C die drei Minterme AB , $A\bar{B}$

$\mathcal{K}_{5,15}$	A	B	C		A	B	f_C	$\overline{f_C}$
c_1	1	1	1		1	1	1	0
c_2	1	0	1		1	0	1	0
c_3	0	1	1		0	1	1	0
c_4	0	0	0		0	0	0	1
(a)					(b)			

Tabelle 5.15.: Konfigurationstabelle $\mathcal{K}_{5,15}$ mit $C \in \mathbf{F}^*$ und der sich daraus ergebenden totalen inneren Instanzfunktionen f_C und $\overline{f_C}$.

sowie $\overline{AB}$ aufweist und sich in KDNF schreiben lässt als $f_C(A, B) = AB + \overline{A}\overline{B} + \overline{A}B$. Andererseits hat $\overline{f_C}$ den Minterm $\overline{A}\overline{B}$ und kann mittels $\overline{f_C}(A, B) = \overline{A}\overline{B}$ in KDNF repräsentiert werden.

Beginnend mit dem Minterm AB von f_C wird A entfernt und überprüft, ob B Teil des Implikants $\overline{A}\overline{B}$ von $\overline{f_C}$ ist. Da dies nicht der Fall ist, wird A aus Gründen der Redundanz permanent weggelassen und mit dem reduzierten Term B fortgefahren. Der Term B kann nicht weiter reduziert werden, da sonst aufgrund von $B = B \cdot 1$ folgen würde, dass das neutrale Element 1 ein Implikant von f_C ist und somit $f_C(A, B) = 1$ gilt. Letzteres ist offensichtlich nicht der Fall. Insgesamt resultiert daraus ein erster Primimplikant B von f_C . Der gleiche Vorgang wird nun für alle Permutationen von AB wiederholt, damit sich jeder in AB enthaltene Primimplikant ermitteln lässt. Es wird somit BA gewählt, daraus B entfernt und überprüft, ob A Teil des Implikants $\overline{A}\overline{B}$ von $\overline{f_C}$ ist. Dies ist nicht der Fall, womit B aus BA entfernt wird. Der Term A kann analog zu obiger Argumentation für den Primimplikanten B nicht weiter reduziert werden und es wird gefolgert, dass A ein Primimplikant von f_C ist. Weitere Permutationen von AB gibt es nicht und es kann mit dem zweiten Minterm $\overline{A}\overline{B}$ fortgefahren werden. Nach Anwendung derselben Schritte ergibt sich aus dem Minterm $\overline{A}\overline{B}$ beziehungsweise dessen Permutation $\overline{B}\overline{A}$ der Primimplikant A . Schliesslich auch für den Minterm $\overline{A}B$ ausgeführt, ergeben sich alles in allem die beiden Primimplikanten A und B von f_C .

Es ist unschwer zu erkennen, dass diese systematische Vorgehensweise aufgrund der darin enthaltenen zahlreichen Minterm-Permutationen und der daraus folgenden wiederholten Analyse gleicher Terme nicht recheneffizient ist. Vom QMC-Verfahren für die Suche nach Primimplikanten gemäss Algorithmus 5.2 inspiriert, soll das Baumgartner-Verfahren entsprechend angepasst werden. Im Wesentlichen soll die Darstellungsform des Quine-McCluskey-Verfahrens mittels Tabellen übernommen werden, ohne dabei das Grundprinzip der Redundanzregel zu verletzen. Wie sich gleich zeigen wird, liegt der wesentliche Vorteil dieser Darstellung darin, dass die vielen Permutationsschritte ersetzt werden durch die Weitergabe der reduzierten Terme in neue Tabellen und so wiederholte Analysen gleicher Terme wegfallen. Des Weiteren bietet diese

Übernahme eine gute Basis für die direkte Kontrastierung des QMC- und des Baumgartner-Algorithmus für die Suche nach den Primimplikanten. Zu diesem Zweck wird in der Folge dieselbe Konfigurationstabelle $\mathcal{K}_{5,1}$ über $\mathbf{H} = \{A, B, C, D, E\}$ mit der inneren Instanzfunktion g_E untersucht, die bereits sowohl manuell über die direkte Verwendung der Booleschen Gesetze als auch systematisch mittels des Quine-McCluskey-Verfahrens analysiert wurde (vgl. S. 177). Gesucht werden also erneut die RDN-Formen der inneren Instanzfunktion g_E , die total ist und folgendermassen als KDNF dargestellt werden kann:

$$g_E(\cdot) = ABCD + ABC\bar{D} + AB\bar{C}D + AB\bar{C}\bar{D} + A\bar{B}C\bar{D} + A\bar{B}\bar{C}D + A\bar{B}\bar{C}\bar{D} + \bar{A}\bar{B}CD + \bar{A}\bar{B}C\bar{D} + \bar{A}\bar{B}\bar{C}D + \bar{A}\bar{B}\bar{C}\bar{D} \quad (5.39)$$

Damit die Redundanzregel angewendet werden kann, ist weiter die innere Instanzfunktion $\bar{g}_E$ erforderlich, die sich ebenfalls direkt aus der Konfigurationstabelle ablesen lässt. Die innere Instanzfunktion $\bar{g}_E$ ist auch total und via Minterme durch folgende KDNF repräsentierbar:

$$\bar{g}_E(\cdot) = \bar{A}\bar{B}CD + \bar{A}BCD + \bar{A}BC\bar{D} + \bar{A}B\bar{C}D + \bar{A}B\bar{C}\bar{D} \quad (5.40)$$

Ähnlich dem QMC-Verfahren werden für die Suche nach den Primimplikanten von g_E in einem ersten Schritt die Minterme aufgelistet. Tabelle 5.16 zeigt die entsprechende Ausgangslage.

g₁₁	<i>A</i>	$\bar{B}$	<i>C</i>	<i>D</i>
g₇	$\bar{A}$	<i>B</i>	<i>C</i>	<i>D</i>
g₆	$\bar{A}$	<i>B</i>	<i>C</i>	$\bar{D}$
g₅	$\bar{A}$	<i>B</i>	$\bar{C}$	<i>D</i>
g₄	$\bar{A}$	<i>B</i>	$\bar{C}$	$\bar{D}$
m₁₅	<i>A</i>	<i>B</i>	<i>C</i>	<i>D</i>
m₁₄	<i>A</i>	<i>B</i>	<i>C</i>	$\bar{D}$
m₁₃	<i>A</i>	<i>B</i>	$\bar{C}$	<i>D</i>
m₁₂	<i>A</i>	<i>B</i>	$\bar{C}$	$\bar{D}$
m₁₀	<i>A</i>	$\bar{B}$	<i>C</i>	$\bar{D}$
m₉	<i>A</i>	$\bar{B}$	$\bar{C}$	<i>D</i>
m₈	<i>A</i>	$\bar{B}$	$\bar{C}$	$\bar{D}$
m₃	$\bar{A}$	$\bar{B}$	<i>C</i>	<i>D</i>
m₂	$\bar{A}$	$\bar{B}$	<i>C</i>	$\bar{D}$
m₁	$\bar{A}$	$\bar{B}$	$\bar{C}$	<i>D</i>
m₀	$\bar{A}$	$\bar{B}$	$\bar{C}$	$\bar{D}$

Tabelle 5.16.: Für die Verwendung des Verfahrens nach Baumgartner aufbereitete Mintermtabelle mit Mintermen $\mathbf{g}_j$, $j \in \{4, 5, 6, 7, 11\}$, von $\bar{g}_E$ und Mintermen $\mathbf{m}_i$, $i \in \{0, 1, 2, 3, 8, 9, 10, 12, 13, 14, 15\}$, von g_E .

Anstelle der Gruppierung aller Minterme von g_E anhand der Anzahl negativer Literale, werden in Tabelle 5.16 jedoch zwei andere Gruppierungen gegenübergestellt: In der ersten Gruppe befinden sich die Minterme $\mathbf{g}_j$, $j \in \{4, 5, 6, 7, 11\}$, der Instanzfunktion $\overline{g_E}$ und in der zweiten Gruppe die Minterme $\mathbf{m}_i$, $i \in \{0, 1, 2, 3, 8, 9, 10, 12, 13, 14, 15\}$, der Instanzfunktion g_E . Der Index eines Minterms ist wiederum über die Binärzahl festgelegt, die über jene Wertezuordnung bestimmt wird, für die der Minterm den Wert 1 annimmt (vgl. S. 189).

Mithilfe dieser Auflistung wird nun jeder Minterm $\mathbf{m}_i$ von g_E nacheinander betrachtet und für jedes darin enthaltene Literal geprüft, ob der um dieses Literal reduzierte Term Teil eines Minterms $\mathbf{g}_j$ von $\overline{g_E}$ ist. Ist dies nicht der Fall, wird der entsprechende, um ein Literal reduzierte Term in eine neue Tabelle geschrieben, sonst nicht. Für den Minterm $\mathbf{m}_{15}$ ausgeführt bedeutet dies konkret: Aus $ABCD$ wird A entfernt und geprüft, ob BCD Teil eines $\mathbf{g}_j$ ist. Mit $\mathbf{g}_7$ ist ein solcher Term vorhanden, womit BCD nicht als Implikant in eine neue Tabelle übernommen werden kann. Weiter mit B und dem reduzierten Term ACD kommt man aufgrund von $\mathbf{g}_{11}$ zum selben Schluss. Im Gegensatz dazu kann sowohl der um C reduzierte Term ABD als auch der um D reduzierte Term ABC als Implikant in einer neuen Tabelle übernommen werden, da kein $\mathbf{g}_j$ existiert, von dem ABD oder ABC ein Teil ist. Wiederholt für jeden Minterm und jedes darin enthaltene Literal resultiert daraus die Tabelle 5.17a, wobei gleiche Terme bis auf einen aus der Tabelle entfernt werden können. Kann aus einem Term kein Literal entfernt werden, ohne dass der reduzierte Term jeweils Teil eines Minterms $\mathbf{g}_j$ von $\overline{g_E}$ ist, handelt es sich um einen Primimplikanten von g_E . Für alle anderen Implikanten wird eine Reduktion mittels Redundanzregel solange wiederholt, bis sich kein reduzierter Term in eine neue Tabelle übertragen lässt. Aus einem weiteren Durchgang resultiert Tabelle 5.17b, die nur aus Primimplikanten besteht. Das Reduktionsverfahren kann damit abgebrochen und die mit „prim“ gekennzeichneten Terme als Resultat ausgegeben werden.

Im Unterschied zum QMC-Verfahren bleibt das Verfahren nach Baumgartner bei partiellen Instanzfunktionen unverändert und verwendet genau die gleichen Informationen wie jene, die bei totalen inneren Instanzfunktionen zum Einsatz kommen. In Tabelle 5.18 sind die Ausgangstabelle 5.18a sowie die daraus resultierende finale, die Primimplikanten der partiellen inneren Instanzfunktion h_U beinhaltende Tabelle 5.18b dargestellt. Tabelle 5.18b entsteht erneut durch wiederholtes Anwenden der Redundanzregel. Analog kann für h_V verfahren werden, woraus sich die beiden Primimplikanten R und S ergeben. Alles in allem ergeben sich somit auch mit diesem Verfahren die gesuchten Primimplikanten.

$\mathbf{g}_{11}$	A	$\bar{B}$	C	D	
$\mathbf{g}_7$	$\bar{A}$	B	C	D	
$\mathbf{g}_6$	$\bar{A}$	B	C	$\bar{D}$	
$\mathbf{g}_5$	$\bar{A}$	B	$\bar{C}$	D	
$\mathbf{g}_4$	$\bar{A}$	B	$\bar{C}$	$\bar{D}$	
$\mathbf{m}_{15}, \mathbf{m}_{13}$	A	B	$*$	D	
$\mathbf{m}_{15}, \mathbf{m}_{14}$	A	B	C	$*$	
$\mathbf{m}_{14}, \mathbf{m}_{10}$	A	$*$	C	$\bar{D}$	
$\mathbf{m}_{14}, \mathbf{m}_{12}$	A	B	$*$	$\bar{D}$	
$\mathbf{m}_{13}, \mathbf{m}_9$	A	$*$	$\bar{C}$	D	
$\mathbf{m}_{13}, \mathbf{m}_{12}$	A	B	$\bar{C}$	$*$	
$\mathbf{m}_{12}, \mathbf{m}_8$	A	$*$	$\bar{C}$	$\bar{D}$	
$\mathbf{m}_{10}, \mathbf{m}_2$	$*$	$\bar{B}$	C	$\bar{D}$	
$\mathbf{m}_{10}, \mathbf{m}_8$	A	$\bar{B}$	$*$	$\bar{D}$	
$\mathbf{m}_9, \mathbf{m}_1$	$*$	$\bar{B}$	$\bar{C}$	D	
$\mathbf{m}_9, \mathbf{m}_8$	A	$\bar{B}$	$\bar{C}$	$*$	
$\mathbf{m}_8, \mathbf{m}_0$	$*$	$\bar{B}$	$\bar{C}$	$\bar{D}$	
$\mathbf{m}_3, \mathbf{m}_1$	$\bar{A}$	$\bar{B}$	$*$	D	
$\mathbf{m}_3, \mathbf{m}_2$	$\bar{A}$	$\bar{B}$	C	$*$	
$\mathbf{m}_2, \mathbf{m}_0$	$\bar{A}$	$\bar{B}$	$*$	$\bar{D}$	
$\mathbf{m}_1, \mathbf{m}_0$	$\bar{A}$	$\bar{B}$	$\bar{C}$	$*$	

(a)

$\mathbf{g}_{11}$	A	$\bar{B}$	C	D	
$\mathbf{g}_7$	$\bar{A}$	B	C	D	
$\mathbf{g}_6$	$\bar{A}$	B	C	$\bar{D}$	
$\mathbf{g}_5$	$\bar{A}$	B	$\bar{C}$	D	
$\mathbf{g}_4$	$\bar{A}$	B	$\bar{C}$	$\bar{D}$	
$\mathbf{m}_{15}, \mathbf{m}_{14}, \mathbf{m}_{13}, \mathbf{m}_{12}$	A	B	$*$	$*$	prim
$\mathbf{m}_{14}, \mathbf{m}_{12}, \mathbf{m}_{10}, \mathbf{m}_8$	A	$*$	$*$	$\bar{D}$	prim
$\mathbf{m}_{13}, \mathbf{m}_{12}, \mathbf{m}_9, \mathbf{m}_8$	A	$*$	$\bar{C}$	$*$	prim
$\mathbf{m}_{10}, \mathbf{m}_8, \mathbf{m}_2, \mathbf{m}_0$	$*$	$\bar{B}$	$*$	$\bar{D}$	prim
$\mathbf{m}_9, \mathbf{m}_8, \mathbf{m}_1, \mathbf{m}_0$	$*$	$\bar{B}$	$\bar{C}$		prim
$\mathbf{m}_3, \mathbf{m}_2, \mathbf{m}_1, \mathbf{m}_0$	$\bar{A}$	$\bar{B}$	$*$	$*$	prim

(b)

Tabelle 5.17.: Auflistung 5.17a aller Implikanten von g_E , die nach einem ersten Durchgang mittels Redundanzregel aus der Tabelle 5.16 entstehen und Auflistung 5.17b aller Primimplikanten von g_E .

$\mathbf{g}_9$	R	$\bar{S}$	$\bar{T}$	V	
$\mathbf{g}_0$	$\bar{R}$	$\bar{S}$	$\bar{T}$	$\bar{V}$	
$\mathbf{m}_{15}$	R	S	T	V	
$\mathbf{m}_{13}$	R	S	$\bar{T}$	V	
$\mathbf{m}_{11}$	R	$\bar{S}$	T	V	
$\mathbf{m}_7$	$\bar{R}$	S	T	V	
$\mathbf{m}_5$	$\bar{R}$	S	$\bar{T}$	V	
$\mathbf{m}_2$	$\bar{R}$	$\bar{S}$	T	$\bar{V}$	

(a)

$\mathbf{g}_9$	R	$\bar{S}$	$\bar{T}$	V	
$\mathbf{g}_0$	$\bar{R}$	$\bar{S}$	$\bar{T}$	$\bar{V}$	
$\mathbf{m}_{15}, \mathbf{m}_{13}, \mathbf{m}_7, \mathbf{m}_5$	$*$	S	$*$	$*$	prim
$\mathbf{m}_{15}, \mathbf{m}_{11}, \mathbf{m}_7, \mathbf{m}_2$	$*$	$*$	T	$*$	prim

(b)

Tabelle 5.18.: Mintermtabelle 5.18a mit Mintermen $\mathbf{g}_0$ und $\mathbf{g}_9$ von $\bar{h}_U$ sowie Mintermen $\mathbf{m}_i$, $i \in \{2, 5, 7, 11, 13, 15\}$, von h_U und nach Beendigung der wiederholten Verwendung der Redundanzregel resultierende Auflistung 5.18b mit den Primimplikanten S und T von h_U .

Die allgemeine Vorgehensweise dieses ersten Teils ist im Algorithmus 5.4 zusammengefasst. Die Schreibweise ist dabei identisch zu interpretieren wie jene des Algorithmus 5.2. Das heisst, die in Fettschrift notierten Kleinbuchstaben $\mathbf{r}$ und $\mathbf{s}$ stellen Konjunktionsterme dar und $\mathbf{I}_k$ beziehungsweise $\mathbf{J}_k$ beschreiben Mengen von Konjunktionstermen mit k Konjunkten.

Algorithmus 5.4 Suche nach Primimplikanten von f_W gemäss CNA (Baumgartner-Verfahren)

Input: Menge $\mathbf{I}_n$ der Minterme von f_W , Menge $\mathbf{G}$ der Minterme von $\overline{f_W}$

Output: Menge $\mathbf{PI}(f_W)$ der Primimplikanten von f_W

```

1: repeat for  $k = n, n - 1, \dots$ 
2:    $\mathbf{I}_{k-1} := \emptyset$ 
3:    $\mathbf{J}_k := \emptyset$ 
4:   for each Element  $\mathbf{r}$  aus  $\mathbf{I}_k$  do
5:     if  $\nexists X_j$  aus  $\mathbf{r} = \mathbf{s}X_j$ , sodass  $\mathbf{s}$  nicht Teil eines  $\mathbf{g} \in \mathbf{G}$  ist then
6:       Füge  $\mathbf{r}$  zu  $\mathbf{J}_k$  hinzu
7:     else
8:       while  $\exists X_j$  aus  $\mathbf{r} = \mathbf{s}X_j$ , wobei  $\mathbf{s}$  nicht Teil eines  $\mathbf{g} \in \mathbf{G}$  ist do
9:         Füge  $\mathbf{s}$  zu  $\mathbf{I}_{k-1}$  hinzu
10:  Entferne alle identischen Elemente aus  $\mathbf{I}_{k-1}$  bis auf jeweils eines
11: until  $\mathbf{I}_{k-1} = \emptyset$ 
12: return  $\mathbf{PI}(f_W) = \mathbf{J}_n \cup \mathbf{J}_{n-1} \cup \dots$ 

```

(b) Suche nach den RDN-Formen

Im zweiten Teil (b) werden die aus Teil (a) gefolgerten Primimplikanten einer Instanzfunktion f_W erneut so miteinander disjunktiv verknüpft, dass der gesamte Boolesche Ausdruck Bedingung (b) der Definition 4.15 erfüllt. Dazu nutzt das Baumgartner-Verfahren wiederum ein sich direkt aus der Definition 2.2 ergebendes Prinzip: Sei $\mathbf{p}_i$, $1 \leq i \leq m$, ein beliebiger aus Literalen der geprüften Faktoren aus $\mathbf{F}$ bestehender minimal hinreichender Konjunktionsterm von W , wobei der Faktor W ebenfalls in $\mathbf{F}$ enthalten ist. Eine notwendige Bedingung $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_m$ von W ist genau dann minimal notwendig, wenn für jede darin enthaltene Konjunktion $\mathbf{p}_i$, $1 \leq i \leq m$, gilt, dass wenn sie aus $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_m$ entfernt wird, der um $\mathbf{p}_i$ reduzierte Term nicht notwendig für W ist. Die um $\mathbf{p}_i$ reduzierte Disjunktion ist genau dann nicht notwendig für W , wenn es Konfigurationen gibt, bei denen zwar W gegeben ist, aber kein Disjunkt der um $\mathbf{p}_i$ reduzierten Disjunktion. Umgekehrt ist ein $\mathbf{p}_i$, $1 \leq i \leq m$, aus $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_i + \dots + \mathbf{p}_m$ genau dann redundant und kann aus der notwendigen Bedingung entfernt werden, wenn immer wenn W gegeben ist, mindestens ein Disjunkt aus $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_{i-1} + \mathbf{p}_{i+1} + \dots + \mathbf{p}_m$ gegeben ist.

Mithilfe der Konfigurationsalgebra ausgedrückt kann ein $\mathbf{p}_i$ aus der Abdeckung $\mathbf{p}_1 + \dots + \mathbf{p}_i + \dots + \mathbf{p}_m$ von f_W entfernt werden, wenn $\mathbf{p}_1 + \dots + \mathbf{p}_{i-1} + \mathbf{p}_{i+1} + \dots + \mathbf{p}_m$ eine Abdeckung von f_W bleibt. Mit Einbezug der Minterme von f_W kann $\mathbf{p}_i$ somit dann aus der Disjunktion entfernt werden, wenn daraus keine Abdeckung resultiert, die mindestens einen Minterm von f_W nicht abdeckt. Kann kein Primimplikant aus der Disjunktion entfernt werden, ohne dass mindestens ein Minterm von f_W nicht abgedeckt wird, handelt es sich bei dieser Disjunktion um eine irredundante Abdeckung von f_W .

Um basierend auf dieser Überlegung systematisch alle irredundanten Abdeckungen von f_W zu finden, dient als Ausgangspunkt die Disjunktion $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_n$ aller Primimplikanten, die aus Teil (a) resultieren. Dass es sich dabei jeweils um eine Abdeckung von f_W handelt ist trivial, da es basierend auf Teil (a) zu jedem Minterm von f_W mindestens einen Primimplikanten gibt, der diesen Minterm abdeckt. Gemäss dem Baumgartner-Verfahren kann nun für jedes $\mathbf{p}_i$ aus $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_n$ und jede Permutation von $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_n$ folgendermassen vorgegangen werden (vgl. u. a. Baumgartner (2009), S. 87): Der Primimplikant $\mathbf{p}_i$ wird aus der Disjunktion entfernt und überprüft, ob es einen Minterm gibt, der von keinem Primimplikanten aus $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_{i-1} + \mathbf{p}_{i+1} + \dots + \mathbf{p}_n$ abgedeckt wird. Ist dies der Fall, fügt man $\mathbf{p}_i$ wieder zu $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_{i-1} + \mathbf{p}_{i+1} + \dots + \mathbf{p}_n$ hinzu und macht mit der Eliminierung von $\mathbf{p}_{i+1}$ weiter. Andernfalls wird der Vorgang mit $\mathbf{p}_{i+1}$ fortgeführt ohne $\mathbf{p}_i$ wieder in $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_{i-1} + \mathbf{p}_{i+1} + \dots + \mathbf{p}_n$ aufzunehmen. Sobald man mit diesem Vorgang keinen Primimplikanten mehr permanent aus einer Disjunktion entfernen kann, handelt es sich um eine irredundante Abdeckung von f_W .

Obwohl ein solches Vorgehen bisweilen auch als Teil des Quine-McCluskey-Verfahrens angesehen wird (vgl. bspw. Ihringer (1994), S. 201 f.), soll es an dieser Stelle zwecks der Kontrastierung von csQCA und csCNA eindeutig dem Baumgartner-Verfahren zugeordnet werden. Der Nachweis dafür, dass aus diesem Verfahren wiederum alle irredundanten Abdeckungen von f_W resultieren, kann auf analoge Weise geführt werden wie jener der zeigt, dass innerhalb der Suche nach den Primimplikanten mittels der wiederholten Anwendung der Redundanzregel alle Primimplikanten gefunden werden. Anstelle von konjunktiv verknüpften Literalen werden hier disjunktiv verknüpfte Konjunktionsterme betrachtet: Ist $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_k$ eine beliebige irredundante Abdeckung von f_W mit k Primimplikanten, ist auch die Hinzunahme jedes weiteren Primimplikanten bis hin zu $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_k + \mathbf{p}_{k+1} + \dots + \mathbf{p}_m$ eine Abdeckung von f_W . Da die obige Redundanzprüfung für jeden Primimplikanten jeder Permutation von $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_k + \mathbf{p}_{k+1} + \dots + \mathbf{p}_n$ durchgeführt wird, wird mit bei $\mathbf{p}_{k+1}$ beginnender Abarbeitung von $\mathbf{p}_{k+1} + \dots + \mathbf{p}_n + \mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_k$ die irredundante Abdeckung $\mathbf{p}_1 + \mathbf{p}_2 + \dots + \mathbf{p}_k$ von f_W folgen.

Zur detaillierten Veranschaulichung der einzelnen Schritte wird erneut das Beispiel aus Tabelle 5.15 mit der inneren Instanzfunktion $f_C = AB + \bar{A}\bar{B} + \bar{A}B$ betrachtet (vgl. 207). Mit Teil (a) des

	m_0	m_1	m_2	m_3	m_8	m_9	m_{10}	m_{12}	m_{13}	m_{14}	m_{15}
AB								x	x	x	x
$A\bar{D}$					x		x	x		x	
$A\bar{C}$					x	x		x	x		
$\bar{B}\bar{D}$	x		x		x		x				
$\bar{B}\bar{C}$	x	x			x	x					
$\bar{A}\bar{B}$	x	x	x	x							

Tabelle 5.19.: Primimplikantentabelle der inneren Instanzfunktion g_E .

Baumgartner-Verfahrens wurden aus $\mathcal{K}_{5,15}$ die beiden Primimplikanten A und B von f_C ermittelt. Diese werden zur Abdeckung $A + B$ verknüpft. Aus $A + B$ wird nun A entfernt und überprüft, ob es einen Minterm von f_C gibt, der von B nicht abgedeckt wird. Dies gilt genau dann, wenn ein Minterm $\bar{B}$ aufweist. Letzteres ist offensichtlich der Fall, womit A nicht redundant ist und dementsprechend wieder in die Abdeckung integriert wird. Weiter mit dem nächsten Primimplikanten wird B aus der Abdeckung entfernt und analog überprüft, ob es einen Minterm von f_C gibt, der $\bar{A}$ enthält. Auch das ist der Fall, womit B wieder in die Abdeckung aufgenommen wird. Es kann gefolgert werden, dass die Disjunktion $A + B$ eine irredundante Abdeckung von f_W ist. Die gleichen Schritte werden für die Permutation $B + A$ durchgeführt, aus der dieselbe irredundante Abdeckung folgt. Insgesamt resultiert somit aus Teil (b) des Baumgartner-Verfahrens eine einzige, aus A und B bestehende irredundante Abdeckung von f_C , aus der die eindeutige RDN-Form $f_C(A, B) = A + B$ folgt.

Für die Kontrastierung von csQCA und csCNA wird wiederum die Konfigurationstabelle 5.1 betrachtet, aus der die Primimplikanten AB , AD , $A\bar{C}$, $\bar{B}\bar{D}$, $\bar{B}\bar{C}$ sowie $\bar{A}\bar{B}$ folgen (vgl. S. 177 ff.). Dazu wird auch hier in Anlehnung an das QMC-Verfahren eine Primimplikantentabelle benutzt, die sich zusätzlich zur besseren Vergleichbarkeit der beiden Verfahren aufgrund der übersichtlichen Darstellung anbietet. In Tabelle 5.19 ist diese nochmals dargestellt.

Zur Vermeidung der vielen Permutationen und der damit verbundenen wiederholten Analyse gleicher Terme, wird der Vorgang der Suche nach allen irredundanten Abdeckungen mittels des Baumgartner-Verfahrens erneut durch die Verwendung von Tabellen etwas angepasst und damit recheneffizienter gestaltet. Tabelle 5.20 zeigt eine solche, die durch wiederholte Verwendung der Redundanzprüfung jedes Primimplikanten und dessen allfällige Streichung aus der Abdeckung entsteht. Diese Tabelle soll als *Abdeckungstabelle* bezeichnet werden.

Sei $\mathbf{D}_i$ eine Menge von Disjunktionen, die jeweils aus i Primimplikanten bestehen. Ausgangspunkt der Abdeckungstabelle bildet die aus einer Disjunktion bestehende Menge $\mathbf{D}_6$. Die in $\mathbf{D}_6$ enthaltene Disjunktion ist jene, die aus der disjunktiven Verknüpfung aller Primimplikanten

D₆	(i)	AB	+	$A\bar{D}$	+	$A\bar{C}$	+	$\bar{B}\bar{D}$	+	$\bar{B}\bar{C}$	+	$\bar{A}\bar{B}$	
D₅	(i)	AB	+	*	+	$A\bar{C}$	+	$\bar{B}\bar{D}$	+	$\bar{B}\bar{C}$	+	$\bar{A}\bar{B}$	
	(ii)	AB	+	$A\bar{D}$	+	*	+	$\bar{B}\bar{D}$	+	$\bar{B}\bar{C}$	+	$\bar{A}\bar{B}$	
	(iii)	AB	+	$A\bar{D}$	+	$A\bar{C}$	+	*	+	$\bar{B}\bar{C}$	+	$\bar{A}\bar{B}$	
	(iv)	AB	+	$A\bar{D}$	+	$A\bar{C}$	+	$\bar{B}\bar{D}$	+	*	+	$\bar{A}\bar{B}$	
D₄	(5i),(5ii)	AB	+	*	+	*	+	$\bar{B}\bar{D}$	+	$\bar{B}\bar{C}$	+	$\bar{A}\bar{B}$	irred
	(5i),(5iv)	AB	+	*	+	$A\bar{C}$	+	$\bar{B}\bar{D}$	+	*	+	$\bar{A}\bar{B}$	irred
	(5ii),(5iii)	AB	+	$A\bar{D}$	+	*	+	*	+	$\bar{B}\bar{C}$	+	$\bar{A}\bar{B}$	irred
	(5iii),(5iv)	AB	+	$A\bar{D}$	+	$A\bar{C}$	+	*	+	*	+	$\bar{A}\bar{B}$	irred

Tabelle 5.20.: Abdeckungstabelle, die beginnend bei **D₆** durch wiederholtes Prüfen der Redundanz der Primimplikanten die irredundanten Abdeckungen aus **D₄** liefert.

von g_E aus Teil (a) entsteht. Aus dieser Disjunktion wird nun Primimplikant für Primimplikant entfernt und jeweils geprüft, ob die resultierende Disjunktion immer noch eine Abdeckung aller Minterme darstellt oder nicht. Dies kann leicht anhand der Primimplikantentabelle 5.19 überprüft werden: Entfernt man aus der Disjunktion einen Primimplikanten $\mathbf{p}_i$, wird die Zeile desselben Primimplikanten aus der Primimplikantentabelle gestrichen. Resultiert daraus eine reduzierte Primimplikantentabelle, die eine Spalte enthält, die kein x-Symbol aufweist, deckt die um $\mathbf{p}_i$ reduzierte Disjunktion nicht mehr alle Minterme ab. Enthält die reduzierte Primimplikantentabelle umgekehrt nach wie vor in jeder Spalte mindestens ein x-Symbol, stellt die reduzierte Disjunktion eine Abdeckung aller Minterme dar. Jede um einen Primimplikanten reduzierte Disjunktion, die eine Abdeckung aller Minterme bleibt, wird in die Menge **D₅** übertragen, wobei der wegreduzierte Primimplikant durch ein *-Symbol ersetzt wird.¹² Dieser Vorgang wird für jeden Primimplikanten jeder Disjunktion aus **D₅** wiederholt und die reduzierten Disjunktionen in **D₄** übertragen. Dabei werden identische Disjunktionen bis auf eine weggestrichen. Lässt sich aus einer Disjunktion kein Primimplikant entfernen, ohne dass die um diesen Primimplikanten reduzierte Disjunktion keine Abdeckung aller Minterme mehr darstellt, handelt es sich um eine irredundante Abdeckung. Eine solche Disjunktion wird mit „irred“ markiert. Ist keine Disjunktion mehr reduzierbar, kann der Vorgang abgebrochen werden. Bei den markierten Disjunktionen handelt es sich schliesslich um alle aus Primimplikanten bestehenden irredundanten Abdeckungen, aus denen sich die bereits bekannten vier RDN-Formen von g_E beziehungsweise die vier RDN-Bikonditionale von E abgelesen lassen.

Die Abdeckungstabellen der partiellen Instanzfunktionen h_U und h_V sehen vergleichsweise einfach aus, weil alle Primimplikanten beider Funktionen Kernprimimplikanten sind. Am Beispiel

¹²Das *-Symbol hat hier keine weitere Funktion und dient lediglich dazu, die Tabelle etwas übersichtlicher zu gestalten.

$\mathbf{D}_2$	(i)	S	+	T	
$\mathbf{D}_2$	(i)	S	+	T	irred

Tabelle 5.21.: Abdeckungstabelle von h_U , die beginnend bei $\mathbf{D}_2$ durch Prüfen der Redundanz der Primimplikanten die irredundanten Abdeckung von h_U liefert.

von h_U veranschaulicht sieht man dementsprechend, dass bereits die aus zwei Disjunkten bestehende Disjunktion am Ausgang der Abdeckungstabelle 5.21 eine irredundante Abdeckung darstellt. Analoges gilt für h_V , womit insgesamt erneut die beiden RDN-Formen $S + T$ von h_U und $R + S$ von h_V folgen beziehungsweise die beiden RDN-Bikonditionale $S + T \leftrightarrow U$ und $R + S \leftrightarrow V$.

Insgesamt zeigt sich anhand der Abdeckungstabelle deutlich, dass das Baumgartner-Verfahren die redundanten Teile aus der aus allen Primimplikanten bestehenden Abdeckung abbaut.

Der Abschnitt wird wiederum abgeschlossen mit einem in Pseudocode dargestellten Algorithmus 5.5, der die obige Vorgehensweise zusammengefasst. Dabei sind δ sowie ρ Boolesche Ausdrücke in DNF und $\mathbf{D}_k$ beziehungsweise $\mathbf{A}_k$ Mengen von solchen Termen in DNF mit k Disjunkten. Das heisst, δ könnte etwa $X_1X_2X_3 + X_4\overline{X_2}$ oder $X_3\overline{X_4} + X_1X_2$ repräsentieren und diese beiden Disjunktionen wären dann beispielsweise Elemente der Menge $\mathbf{A}_2$.

Algorithmus 5.5 Suche nach RDN-Formen von f_W gemäss CNA (Baumgartner-Verfahren)

Input: Menge $\mathbf{M}$ der Minterme von f_W , Menge $\mathbf{PI}(f_W)$ der Primimplikanten von f_W

Output: Menge $\mathbf{A}(f_W)$ aller RDN-Formen von f_W

- 1: Bilde die einelementige Menge $\mathbf{D}_n$ bestehend aus der Disjunktion aller n Primimplikanten
 - 2: **repeat for** $k = n, n - 1, \dots$
 - 3: $\mathbf{D}_{k-1} := \emptyset$
 - 4: $\mathbf{A}_k := \emptyset$
 - 5: **for each** Disjunktion ρ aus $\mathbf{D}_k$ **do**
 - 6: **if** $\nexists \mathbf{p}_j$ aus $\rho = \delta + \mathbf{p}_j$, sodass δ alle Minterme aus $\mathbf{M}$ abdeckt **then**
 - 7: Füge ρ zu $\mathbf{A}_k$ hinzu
 - 8: **else**
 - 9: **while** $\exists \mathbf{p}_j$ aus $\rho = \delta + \mathbf{p}_j$, wobei δ alle Minterme aus $\mathbf{M}$ abdeckt **do**
 - 10: Füge δ zu $\mathbf{D}_{k-1}$ hinzu
 - 11: Entferne alle identischen Elemente aus $\mathbf{D}_{k-1}$ bis auf jeweils eines
 - 12: **until** $\mathbf{D}_{k-1} = \emptyset$
 - 13: **return** $\mathbf{A}(f_W) = \mathbf{A}_n \cup \mathbf{A}_{n-1} \cup \dots$
-

5.2.4. csQCA-Algorithmus und csCNA-Algorithmus im direkten Vergleich

Nachdem basierend auf der Konfigurationsalgebra aufgezeigt wurde, welche Prinzipien die Algorithmen von csQCA und csCNA für die Suche nach RDN-Formen nutzen und wie diese operational umgesetzt werden können, werden die Suchalgorithmen abschliessend nochmals direkt gegenübergestellt. Dazu werden in aller Kürze die wesentlichen Verfahrensschritte dargestellt und die wichtigsten Unterschiede zwischen den Vorgehensweisen von csQCA und csCNA hervorgehoben. Ausgangspunkt soll dabei eine Konfigurationstabelle $\mathcal{K}_F$ mit $W \in \mathbf{F}^*$ sein, für die die RDN-Bikonditionale von W gesucht werden, die aus $\mathcal{K}_F$ folgen. Mathematisch ausgedrückt entspricht die Suche nach diesen RDN-Bikonditionalen von W der Suche nach den RDN-Formen der inneren Instanzfunktion f_W von f mit abhängigem Faktor W , wobei f total ist und $\mathcal{I}(f) = \mathcal{K}_F$ gilt. Obwohl sich die beiden Algorithmen bei der Verarbeitung der vorhandenen Daten unterscheiden, wird mit dieser Transformation der Problemstellung sichtbar, dass beide Algorithmen dieselbe Grundstruktur aufweisen. Beide suchen in einem ersten Teil (a) nach den Primimplikanten von f_W , die in einem zweiten Teil (b) zu irredundanten Abdeckungen von f_W verknüpft werden. Jede so resultierende irredundante Abdeckung von f_W entspricht durch den Äquivalenzoperator mit W verknüpft einem RDN-Bikonditional von W . Die Funktionsweisen und die wesentlichen Unterschiede sind in Tabelle 5.22 wiedergegeben. Dabei wird zusätzlich auf die Algorithmen referenziert, die die jeweiligen Vorgänge detailliert beschreiben. Hierbei soll nochmals gesagt sein, dass das unter csQCA geführte Petrick-Verfahren vom Standard des QCM-Verfahrens abweicht, das csQCA üblicherweise einsetzt. Nichtsdestotrotz wird es in der vorliegenden Arbeit bei csQCA angesiedelt, da das Verfahren von Petrick eng mit QMC verknüpft ist. Die Kennzeichnung mittels Stern soll jedoch die Abweichung vom Standard verdeutlichen.

Im Kontext des konfiguralen kausalen Modellierens und der damit einhergehenden Interpretation von Instanzfunktionen hat das Baumgartner-Verfahren einen Vorteil bei der Suche nach den Primimplikanten. Wie sich mit dem vorherigen Abschnitt gezeigt hat, benötigt das Baumgartner-Verfahren für die Ermittlung der Primimplikanten aus Teil (a) jeweils nur genau jene Informationen, die in der untersuchten Konfigurationstabelle enthalten sind. Im Gegensatz dazu trifft das QMC-Verfahren aufgrund der Verwendung von Erweiterungen oftmals zusätzliche Annahmen, die so in der Konfigurationstabelle nicht enthalten sind. So wurde beispielsweise zu Beginn des Abschnitts 5.2 im Zusammenhang mit der Definition der inneren Instanzfunktion h_V gezeigt, dass sie deshalb nicht für die Belegung $(0,0,0,1)$ definiert ist, weil daraus folgen würde, dass U dann gegeben ist, wenn S und T nicht gegeben sind (vgl. S. 184). Gemäss der Konfigurationstabelle gibt es entsprechende Konfigurationen aber nicht. Indem QMC mit der Erweiterung h_V^+ arbeitet, nimmt sie jedoch genau das an. Es ist dabei wichtig zu bemerken, dass daraus nicht folgt, dass das QMC-Verfahren damit falsche Resultate liefert. Betrachtet man die

csQCA	csCNA
<i>(a) Suche nach Primimplikanten:</i>	
→ Quine-McCluskey-Verfahren Input: Minterme von f_W und f_W^+ . Wiederholte Verwendung der Resolutionsregel (vgl. Algorithmus 5.2).	→ Baumgartner-Verfahren Input: Minterme von f_W und $\overline{f_W}$. Wiederholte Verwendung der Redundanzregel (vgl. Algorithmus 5.4).
<i>(b) Suche nach RDN-Formen:</i>	
→ Petrick-Verfahren* Kernprimimplikanten-Bestimmung und anschließende Vervollständigung via Verfahren nach Petrick (vgl. Algorithmus 5.3).	→ Baumgartner-Verfahren Entfernen von Primimplikanten aus der aus allen Primimplikanten bestehenden Abdeckung (vgl. Algorithmus 5.5).

Tabelle 5.22.: Gegenüberstellung des csQCA- und csCNA-Verfahrens für die Suche nach den RDN-Formen von f_W .
Der Stern soll verdeutlichen, dass der damit gekennzeichnete Teil etwas vom üblicherweise verwendeten Standardverfahren abweicht.

Verfahren also als reine Rechenschritte, bei denen nur die Resultate als Instantiierungsverhalten von Faktoren interpretiert werden können, ist die Verwendung von Algorithmus 5.2 angemessen. Möchte man jedoch jeweils auch die Zwischenschritte als reales Instantiierungsverhalten interpretieren und die Verwendung der Booleschen Gesetze als Imitation des konfigurationsalen kausalen Modellierens verstehen, eignet sich der Algorithmus 5.2 dafür im Allgemeinen nicht. Mit diesem Algorithmus werden über allfällige Erweiterungen Instantiierungsverhalten ausgedrückt, die so nicht in der Konfigurationstabelle enthalten sind oder den darin enthaltenen Konfigurationen widersprechen (vgl. dazu u. a. auch Baumgartner & Eppe (2014)). Das QCA-Verfahren, dessen primäres Aufdeckungsverfahren QMC darstellt, geht mit der Schwierigkeit solcher zusätzlicher Annahmen teilweise derart um, dass es sich mit „intermediate solutions“ begnügt (vgl. u.a. Schneider & Wagemann (2012), Abschnitt 6.4): Indem nicht alle Minterme der entsprechenden Erweiterung betrachtet werden, um nicht den Daten widersprechendes oder gemäss kausalem Vorwissen unmögliches Instantiierungsverhalten annehmen zu müssen, können mit der Primimplikantensuche via QCA Ausdrücke resultieren, die Implikanten nicht vollständig von allen Redundanzen befreien (vgl. dazu Tabelle 5.8, S. 192). Mit dieser Handhabung mag zwar das Problem dieser zusätzlichen Annahmen umgangen werden, aber sie generiert ein viel schwerwiegenderes Problem. Denn Kapitel 2 hat gezeigt, dass nur solche aus hinreichenden Bedingungen bestehende notwendige Bedingungen eines Literals als kausal relevant für dieses Literal interpretiert werden können, die frei von jeglichen Redundanzen sind. Um weder zusätzliche Annahmen betreffend das Instaniierungsverhalten einführen zu müssen noch Redundanzen zuzulassen, wird künftig für die Ermittlung der Primimplikanten der Algorithmus 5.4 von Baumgartner jenem von Quine und McCluskey vorgezogen.

Bei der Suche nach irredundanten Abdeckungen verhält es sich aus Gründen der Recheneffizienz umgekehrt. Das in Algorithmus 5.3 dargestellte Petrick-Verfahren hat gegenüber dem Algorithmus 5.5 den Vorteil, dass sich mit der vorgängigen Bestimmung der Kernprimimplikanten die Analyse teils erheblich vereinfachen lässt. Je mehr Kernprimimplikanten es im Verhältnis zu allen Primimplikanten gibt, desto einfacher gestaltet sich die Suche nach den irredundanten Abdeckungen. Im Extremfall besteht eine irredundante Abdeckung nur aus Kernprimimplikanten, womit das Verfahren direkt nach deren Bestimmung abgebrochen werden kann.

Alles in allem wird deshalb künftig für die Ermittlung der RDN-Formen von inneren Instanzfunktionen für Teil (a), die Suche nach den Primimplikanten, mit dem Baumgartner-Algorithmus 5.4 gearbeitet, der auf dem Redundanzprinzip beruht. Für Teil (b), die Suche nach den irredundanten Abdeckungen, wird das in Algorithmus 5.3 dargestellte Petrick-Verfahren verwendet.

5.3. Bildung von Minimalen Theorien

Im zweiten der drei auf Seite 5.1 dargestellten Hauptschritte des analytischen Moments werden die Minimalen Theorien der untersuchten Konfigurationstabelle $\mathcal{K}_F$ gebildet. Gemäss Definition 2.8 ist eine Konjunktion Ψ von RDN-Bikonditionalen genau dann eine Minimale Theorie von $\mathcal{K}_F$, wenn (a) Ψ äquivalent zur Konjunktion Φ aller RDN-Bikonditionale ist, die aus $\mathcal{K}_F$ folgen, (b) kein echter Teil von Ψ äquivalent ist zu Φ und (c) keine zwei RDN-Bikonditionale von Literalen desselben Faktors in Ψ enthalten sind.

Bei den bisher betrachteten zwei Beispielen aus den Abbildungen 5.1 und 5.2 ist dieser Schritt relativ trivial. Im Beispiel aus Abbildung 5.1 liegen mit den vier aus dem vorherigen Hauptschritt resultierenden RDN-Bikonditionalen vier Ausdrücke vor, die obige drei Bedingungen einer Minimalen Theorie erfüllen und demnach einfache Minimale Theorien von $\mathcal{K}_{5.1}$ sind. Für jeden der vier Ausdrücke (5.31) bis (5.34) gilt, dass er äquivalent zur Konjunktion aller vier RDN-Bikonditionale ist. Weiter erfüllen alle vier Ausdrücke auch Bedingung (b). Letzteres gilt allgemein für Boolesche Ausdrücke, die nur aus genau einem RDN-Bikonditional β bestehen. Für diesen trivialen Sonderfall gilt gemäss dem Gesetz des neutralen Elements $\beta = \beta \cdot 1$, womit der einzige echte Teil von β das ausgezeichnete Element 1 ist. Zu 1 kann ein RDN-Bikonditional β jedoch nicht äquivalent sein, weil β damit für alle logisch möglichen Belegungen wahr sein müsste. Aus einer entsprechenden Konfigurationstabelle $\mathcal{K}_F$ lässt sich kein RDN-Bikonditional ableiten, weil für alle Faktoren aus F (WD2) gilt. Schliesslich ist für alle vier RDN-Bikonditionale auch Bedingung (c) aus Definition 2.8 erfüllt. Für dieses Beispiel kann die Analyse mit der kausalen Interpretation der vier einfachen Minimalen Theorien abgeschlossen

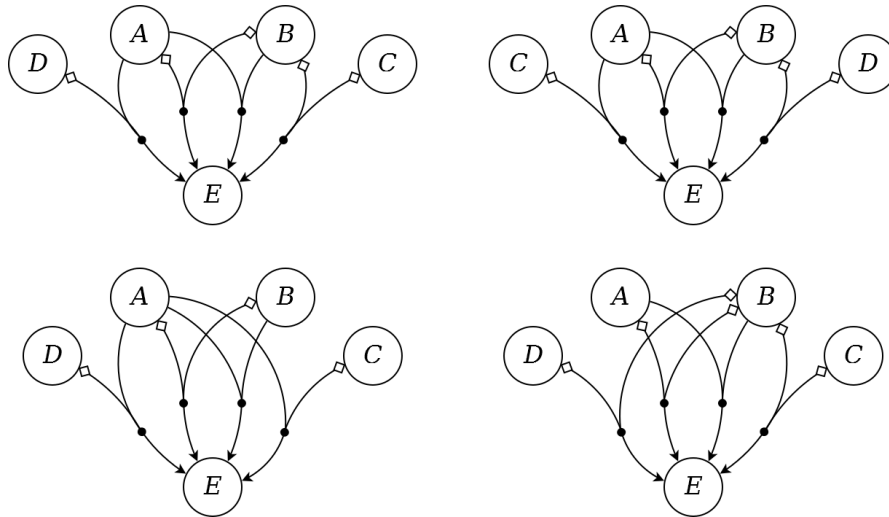


Abbildung 5.23.: Die vier mit $\mathcal{K}_{5,1}$ verträglichen Kausalmodelle, wobei das erste der datengenerierenden Kausalstruktur entspricht.

werden, woraus sich die vier in Abbildung 5.23 dargestellten Kausalmodelle ergeben. Das erste dieser vier Modelle entspricht der datengenerierenden Kausalstruktur aus Abbildung 5.1 (vgl. S. 177). Aus der Analyse der Konfigurationstabelle $\mathcal{K}_{5,1}$ resultieren somit Mehrdeutigkeiten, die über zusätzliche Analysen mit erweiterten Faktormengen abgebaut werden müssen.

Im Beispiel aus Abbildung 5.2 gibt es mit $S + T \leftrightarrow U$ und $R + S \leftrightarrow V$ genau zwei RDN-Bikonditionale, die aus $\mathcal{K}_{5,2}$ folgen. Konjunktiv verknüpft ergibt sich eine Konjunktion von RDN-Bikonditionalen von Literalen paarweise verschiedener Faktoren. Indem weder $S + T \leftrightarrow U$ noch $R + S \leftrightarrow V$ alleine äquivalent zu $(S + T \leftrightarrow U) \cdot (R + S \leftrightarrow V)$ ist, resultiert mit $(S + T \leftrightarrow U) \cdot (R + S \leftrightarrow V)$ insgesamt genau eine Minimale Theorie von $\mathcal{K}_{5,2}$. Auch für dieses Beispiel soll die Analyse an dieser Stelle durch die kausale Interpretation der Minimalen Theorie direkt abgeschlossen werden. Es lässt sich leicht sehen, dass sich damit das eine Common-Cause-Struktur beschreibende Kausalmodell ergibt, das der datengenerierenden Kausalstruktur aus 5.2 entspricht (vgl. S. 177).

Zumal es sich bei den obigen beiden Beispielen betreffend die Systematisierung des zweiten und dritten Hauptschrittes um wenig interessante Szenarien handelt, werden diese in der Folge nicht weiter betrachtet und stattdessen mit einer komplexeren Kausalstruktur weitergearbeitet. Konkret soll für die Entwicklung, Diskussion und Veranschaulichung der Systematisierung der restlichen Schritte des analytischen Moments die Konfigurationstabelle $\mathcal{K}_{5,24}$ aus Abbildung 5.24 betrachtet werden, die von der Kausalstruktur aus derselben Abbildung generiert wird. Neben kausalen Verkettungen und Common-Cause-Strukturen weist diese Kausalstruktur mit

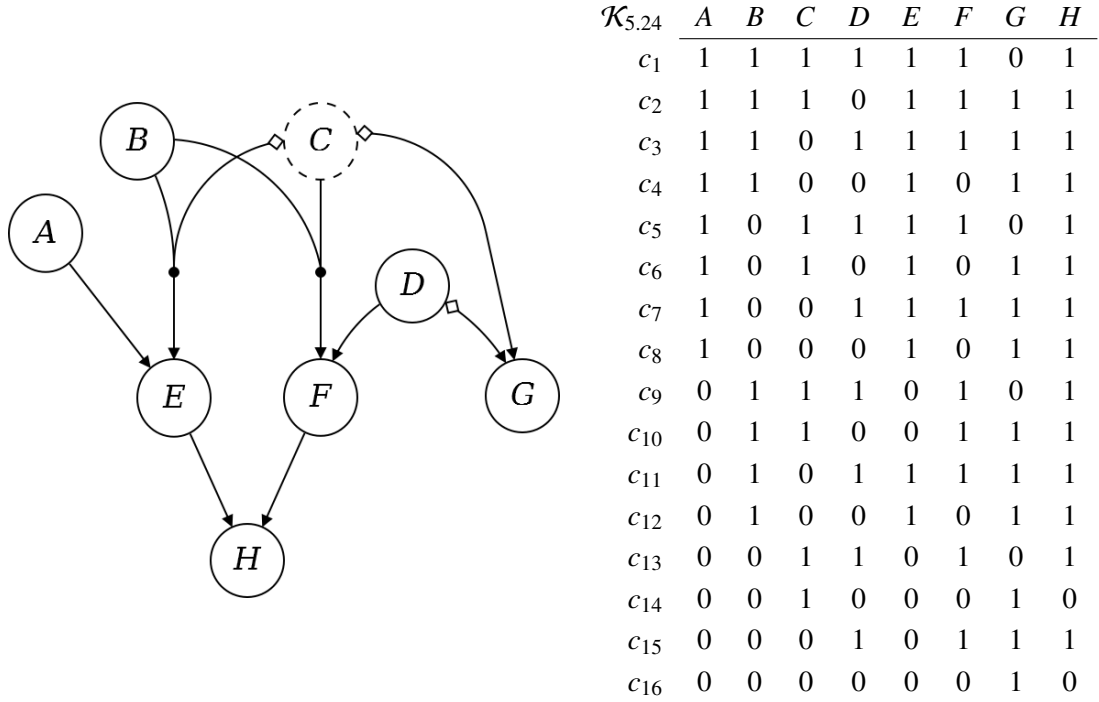


Abbildung 5.24.: Datengenerierende Kausalstruktur zusammen mit der dazugehörigen Konfigurationstabelle $\mathcal{K}_{5,24}$ über $\mathbf{M} = \{A, B, C, D, E, F, G, H\}$. $\mathcal{K}_{5,24}$ ist homogen und maximal divers.

C und $\bar{C}$ auch Switching-Literale auf, sodass zwar C direkt kausal relevant für F und $\bar{C}$ direkt kausal relevant für E sind, keines der beiden Literale aber indirekt kausal relevant für H ist. Im Kausalgraph ist dies entsprechend durch die gestrichelte Kreislinie bei C gekennzeichnet.

Folgt man bei der Analyse von $\mathcal{K}_{5,24}$ dem bis zu diesem Zeitpunkt definierten Aufdeckungsverfahren via den Algorithmen 5.1, 5.4 und 5.3, ergibt sich das folgende System von RDN-Bikonditionalen relativ zu $\mathcal{K}_{5,24}$, das den Ausgangspunkt für die Bildung der Minimalen Theorien von $\mathcal{K}_{5,24}$ darstellt:

$$\bar{G} + \bar{C}F \leftrightarrow D \quad (5.41)$$

$$A + B\bar{C} \leftrightarrow E \quad (5.42)$$

$$A + B\bar{F} + BDG \leftrightarrow E \quad (5.43)$$

$$A + \bar{F}H + BDG \leftrightarrow E \quad (5.44)$$

$$A + BDG + \bar{C}\bar{D}H \leftrightarrow E \quad (5.45)$$

$$D + BC \leftrightarrow F \quad (5.46)$$

$$\bar{C} + \bar{D} \leftrightarrow G \quad (5.47)$$

$$E + F \leftrightarrow H \quad (5.48)$$

$$A + B + D \leftrightarrow H \quad (5.49)$$

$$A + B + F \leftrightarrow H \quad (5.50)$$

$$B + D + E \leftrightarrow H \quad (5.51)$$

5.3.1. Bildung aller Konjunktionen von RDN-Bikonditionalen von Literalen paarweise verschiedener Faktoren

Die systematische Bildung der Minimalen Theorien ergibt sich mehr oder weniger direkt aus der Definition der Minimalen Theorie. Der hier zur Anwendung kommende Ansatz ist, dass die aus Schritt (I) resultierenden RDN-Bikonditionale $\beta_1, \dots, \beta_n$ konjunktiv verknüpft und schrittweise von (strukturellen) Redundanzen befreit werden. Am Ausgangspunkt steht dabei aber nicht die Konjunktion $\Phi = \beta_1 \cdot \dots \cdot \beta_n$ aller RDN-Bikonditionale. Denn gemäss Bedingung (c) der Definition 2.8 kann eine Konjunktion von RDN-Bikonditionalen nur dann eine Minimale Theorie einer Konfigurationstabelle $\mathcal{K}_F$ sein, wenn die Faktoren, von deren Literalen RDN-Bikonditionale in der Konjunktion enthalten sind, paarweise verschieden sind. Als Konsequenz dieser notwendigen Bedingung werden in einem ersten Schritt alle m , $m \geq 1$, kombinatorisch möglichen verschiedenen Konjunktionen $\Gamma_1, \dots, \Gamma_m$ von RDN-Bikonditionalen gebildet, sodass für jede Konjunktion Γ_i , $1 \leq i \leq m$, gilt, dass von jedem Faktor, für dessen Literale es RDN-Bikonditionale gibt, genau ein RDN-Bikonditional in Γ_i enthalten ist.

Im Fall der untersuchten Konfigurationstabelle $\mathcal{K}_{5,24}$ lassen sich die RDN-Bikonditionale (5.41) bis (5.51) zu den folgenden 16 verschiedenen Booleschen Ausdrücken konjunktiv verknüpfen, sodass von H , E , D , F und G genau ein RDN-Bikonditional in der jeweiligen Konjunktion enthalten ist:

$$(E + F \leftrightarrow H) \cdot (A + \bar{B}\bar{C} \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.52)$$

$$(E + F \leftrightarrow H) \cdot (A + \bar{B}\bar{F} + BDG \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.53)$$

$$(E + F \leftrightarrow H) \cdot (A + \bar{F}H + BDG \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.54)$$

$$(E + F \leftrightarrow H) \cdot (A + BDG + \bar{C}\bar{D}H \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.55)$$

$$(A + B + D \leftrightarrow H) \cdot (A + \bar{B}\bar{C} \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.56)$$

$$(A + B + D \leftrightarrow H) \cdot (A + \bar{B}\bar{F} + BDG \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.57)$$

$$(A + B + D \leftrightarrow H) \cdot (A + \bar{F}H + BDG \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.58)$$

$$(A + B + D \leftrightarrow H) \cdot (A + BDG + \bar{C}\bar{D}H \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.59)$$

$$(A + B + F \leftrightarrow H) \cdot (A + \bar{B}\bar{C} \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.60)$$

$$(A + B + F \leftrightarrow H) \cdot (A + \bar{B}\bar{F} + BDG \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.61)$$

$$(A + B + F \leftrightarrow H) \cdot (A + \bar{F}H + BDG \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.62)$$

$$(A + B + F \leftrightarrow H) \cdot (A + BDG + \bar{C}\bar{D}H \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.63)$$

$$(B + D + E \leftrightarrow H) \cdot (A + B\bar{C} \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.64)$$

$$(B + D + E \leftrightarrow H) \cdot (A + B\bar{F} + BDG \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.65)$$

$$(B + D + E \leftrightarrow H) \cdot (A + \bar{F}H + BDG \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.66)$$

$$(B + D + E \leftrightarrow H) \cdot (A + BDG + \bar{C}\bar{D}H \leftrightarrow E) \cdot (\bar{G} + \bar{C}F \leftrightarrow D) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.67)$$

5.3.2. Prüfung der Äquivalenz zur Konjunktion aller RDN-Bikonditionale

Indem die aus dem ersten Teilschritt gebildeten Konjunktionen $\Gamma_1, \dots, \Gamma_m$ möglicherweise jeweils nicht mehr alle aus $\mathcal{K}_F$ gefolgerten RDN-Bikonditionale enthalten, wird in einem zweiten Teilschritt für jede Konjunktion Γ_i geprüft, ob sie äquivalent zur Konjunktion Φ aller RDN-Bikonditionale ist. Ist dem nicht so, genügt Γ_i der Bedingung (a) der Definition 2.8 nicht und weder diese Konjunktion Γ_i noch allfällige äquivalente und redundanzfreie Teile davon können Minimale Theorien von $\mathcal{K}_F$ sein. Genügt Γ_i auch Bedingung (a), werden in einem abschliessenden dritten Teilschritt über ein Reduktionsverfahren alle strukturell minimalen Konjunktionen, die äquivalent zu Γ_i sind, ausgearbeitet. Alle aus diesem dritten Teilschritt der Reduktion hervorgehenden Ausdrücke erfüllen auch Bedingung (b) der Definition 2.8 und entsprechen demnach den Minimalen Theorien von $\mathcal{K}_F$.

Im Beispiel aus Abbildung 5.24 wird also gemäss dem zweiten Teilschritt zur Bildung der Minimalen Theorien geprüft, ob die Konjunktionen (5.52) bis (5.67) äquivalent zur Konjunktion aller RDN-Bikonditionale (5.41) bis (5.51) sind. Nur jene Ausdrücke, für die das gilt, erfüllen neben Bedingung (c) auch Bedingung (a) der Definition 2.8, weshalb nur diese oder echte Teile davon letztendlich auch als Minimale Theorien von $\mathcal{K}_{5.24}$ infrage kommen können. Tatsächlich gilt die Äquivalenz zur Konjunktion aller RDN-Bikonditionale nicht für alle 16 Konjunktionen. Nicht äquivalent zur Konjunktion aller RDN-Bikonditionale sind (5.54), (5.55), (5.66) und (5.67): Alle vier RDN-Bikonditionale sind zusätzlich zu den Belegungen, die den Konfigurationen aus $\mathcal{K}_{5.24}$ entsprechen, noch für weitere Belegungen wahr, für die die Konjunktion aller RDN-Bikonditionale falsch ist. Unter anderem sind sie auch für jene Belegung wahr, bei der E , G und H die Werte 1 und die restlichen Faktoren die Werte 0 annehmen. Eine entsprechende Konfiguration ist in $\mathcal{K}_{5.24}$ nicht enthalten. Dementsprechend können weder diese vier Ausdrücke (5.54), (5.55), (5.66) und (5.67), noch allfällige äquivalente und redundanzfreie Teile davon Minimale Theorien von $\mathcal{K}_{5.24}$ sein. Die übrigbleibenden zwölf Konjunktionen von RDN-Bikonditionalen werden in einem dritten und letzten Teilschritt von allfälligen Redundanzen befreit, damit sie strukturell minimal sind.

5.3.3. Ausarbeitung aller strukturell minimalen Konjunktionen von RDN-Bikonditionalen

Die systematische Ausarbeitung aller strukturell minimalen Teile aller nach Teilschritt 2 übrigbleibenden Konjunktionen von RDN-Bikonditionalen ist von der Vorgehensweise aus Algorithmus 5.5 inspiriert, die für die Entfernung von redundanten minimal hinreichenden Bedingungen (bzw. Primimplikanten von inneren Instanzfunktionen) aus notwendigen Bedingungen (bzw. Abdeckungen von inneren Instanzfunktion) eingesetzt wird. Konkret wird ausgehend von der Menge $\mathbf{B}_n$ aller Konjunktionen von n RDN-Bikonditionalen, die nach dem vorhergehenden Schritt aus Abschnitt 5.3.2 übrigbleiben, folgendermassen verfahren:

Wiederhole die folgenden beiden Schritte für jedes $k = n, n - 1, \dots$ so lange, bis $\mathbf{B}_{k-1}$ leer ist:

1. Für jede Konjunktion aus $\mathbf{B}_k$ und jedes in der Konjunktion enthaltene RDN-Bikonditional β_j , $1 \leq j \leq k$, wird geprüft, ob der um β_j reduzierte Ausdruck äquivalent zur Konjunktion bleibt. Ist er äquivalent, wird der um dieses RDN-Bikonditional reduzierte Term in eine neue Tabelle $\mathbf{B}_{k-1}$ geschrieben. Gibt es keinen reduzierten Ausdruck, der äquivalent ist, wird der Prozess für die entsprechende Konjunktion aus $\mathbf{B}_k$ abgebrochen und die Konjunktion mit einem für Minimale Theorie stehenden Bezeichner „mt“ markiert.
2. Aus der Tabelle $\mathbf{B}_{k-1}$ werden alle identischen Konjunktionen bis auf eine entfernt.

Alle aus der wiederholten Verwendung der obigen beiden Schritte folgenden, mit „mt“ markierten Konjunktionen sind Minimale Theorien von $\mathcal{K}_F$. Weiter entspricht die Menge aller RDN-Bikonditionale aller Minimalen Theorien von $\mathcal{K}_F$ der Menge aller einfachen Minimalen Theorien relativ zu $\mathcal{K}_F$.

Am Beispiel aus Abbildung 5.24 veranschaulicht, entsprechen jene zwölf Konjunktionen von RDN-Bikonditionalen (5.52) bis (5.67), die äquivalent zur Konjunktion aller RDN-Bikonditionale sind, den Elementen der Menge $\mathbf{B}_5$. Für jedes dieser Elemente wird nun der folgende Reduktionsprozess vorgenommen, der am Beispiel der Konjunktion (5.52) im Detail ausgeführt wird. Aus (5.52) wird das erste RDN-Bikonditional $(E + F \leftrightarrow H)$ entfernt und geprüft, ob der resultierende Term äquivalent zu (5.52) ist. Dies ist nicht der Fall. Es wird das zweite RDN-Bikonditional $(A + \bar{B}\bar{C} \leftrightarrow E)$ aus (5.52) entfernt und geprüft, ob dieser daraus resultierende reduzierte Ausdruck äquivalent zu (5.52) ist. Auch dies ist nicht der Fall, weshalb man ohne weitere Aktion mit dem Ausdruck $(\bar{G} + \bar{C}F \leftrightarrow D)$ weiterfährt. Hier stellt sich heraus, dass der um $(\bar{G} + \bar{C}F \leftrightarrow D)$ reduzierte Term äquivalent ist zu (5.52). Entsprechend wird der reduzierte Term in eine neue, als $\mathbf{B}_4$ bezeichnete Tabelle übertragen. Die beiden weiteren RDN-Bikonditionale

aus (5.52) lassen sich ebenfalls nicht entfernen, ohne dass der jeweils resultierende Ausdruck seine Äquivalenz zu (5.52) verliert. Der gleiche Vorgang wird für jede andere Konjunktion aus $\mathbf{B}_5$ durchgeführt, woraus sich insgesamt die zwölf verschiedenen Konjunktionen (5.68) bis (5.79) ergeben, die die Tabelle $\mathbf{B}_4$ bilden. Die Äquivalenzprüfung mit allfälligem Übertrag des reduzierten Ausdrucks in eine neue Tabelle wird für jede Konjunktion aus $\mathbf{B}_4$ wiederholt, wobei sich bei allen Konjunktionen herausstellt, dass sich kein RDN-Bikonditional mehr entfernen lässt, ohne dass der resultierende reduzierte Ausdruck seine Äquivalenz zur entsprechenden Konjunktion aus $\mathbf{B}_4$ verliert. Folglich kann der Prozess für alle Konjunktionen aus $\mathbf{B}_4$ abgebrochen werden und alle mit einem „mt“ als Minimale Theorien von $\mathcal{K}_{5,24}$ markiert werden.

$$(E + F \leftrightarrow H) \cdot (A + B\bar{C} \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.68)$$

$$(E + F \leftrightarrow H) \cdot (A + B\bar{F} + BDG \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.69)$$

$$(A + B + D \leftrightarrow H) \cdot (A + B\bar{C} \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.70)$$

$$(A + B + D \leftrightarrow H) \cdot (A + B\bar{F} + BDG \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.71)$$

$$(A + B + D \leftrightarrow H) \cdot (A + \bar{F}H + BDG \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.72)$$

$$(A + B + D \leftrightarrow H) \cdot (A + BDG + \bar{C}\bar{D}H \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.73)$$

$$(A + B + F \leftrightarrow H) \cdot (A + B\bar{C} \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.74)$$

$$(A + B + F \leftrightarrow H) \cdot (A + B\bar{F} + BDG \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.75)$$

$$(A + B + F \leftrightarrow H) \cdot (A + \bar{F}H + BDG \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.76)$$

$$(A + B + F \leftrightarrow H) \cdot (A + BDG + \bar{C}\bar{D}H \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.77)$$

$$(B + D + E \leftrightarrow H) \cdot (A + B\bar{C} \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.78)$$

$$(B + D + E \leftrightarrow H) \cdot (A + B\bar{F} + BDG \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \quad (5.79)$$

Für jede der Konjunktionen (5.68) bis (5.79) gilt, dass

- (a) sie äquivalent zur Konjunktion aller aus $\mathcal{K}_{5,24}$ gefolgerten RDN-Bikonditionale ist,
- (b) keine echten Teile davon äquivalent zur Konjunktion aller RDN-Bikonditionale sind,
- (c) keine zwei RDN-Bikonditionale von Literalen desselben Faktors in der Konjunktion enthalten sind.

Bei (5.68) bis (5.79) handelt es sich also um Minimale Theorien von $\mathcal{K}_{5,24}$. Es gilt zusätzlich, dass sie genau für jene Belegungen wahr sind, die den Konfigurationen aus $\mathcal{K}_{5,24}$ entsprechen. Dieser Umstand ist nicht generell der Fall und ergibt sich in diesem Beispiel wiederum deshalb, weil die Konfigurationstabelle $\mathcal{K}_{5,24}$ maximal divers ist. Ist die Konfigurationstabelle nicht maximal divers, stellen die darin enthaltenen Konfigurationen typischerweise eine Teilmenge aller

Belegungen dar, für die die Minimalen Theorien wahr sind. So folgen beispielsweise dieselben Minimalen Theorien (5.68) bis (5.79) aus jener nicht maximal diversen Konfigurationstabelle, die die Konfigurationen c_2 bis c_{16} aus $\mathcal{K}_{5,24}$ enthält, nicht aber c_1 .

Die allgemeine Vorgehensweise zur Bildung von Minimalen Theorien aus den RDN-Bikonditionalen ist im Algorithmus 5.6 zusammengefasst. Dabei stellen die griechischen Grossbuchstaben Γ und Γ' Konjunktionen von RDN-Bikonditionalen dar und $\mathbf{B}_k$ beziehungsweise $\mathbf{C}_k$ beschreiben Mengen von Konjunktionen von k RDN-Bikonditionalen.

Algorithmus 5.6 Bildung von Minimalen Theorien

Input: Menge $\mathbf{A}$ aller aus $\mathcal{K}_F$ gefolgerten RDN-Bikonditionale $\beta_1, \beta_2, \dots$ von Literalen von n verschiedenen Faktoren

Output: Menge $\mathbf{T}$ aller Minimalen Theorien von $\mathcal{K}_F$

- 1: Bilde die Menge $\mathbf{B}_n$ bestehend aus allen Konjunktionen von n RDN-Bikonditionalen, wobei für jede Konjunktion aus $\mathbf{B}_n$ gilt, dass keine zwei RDN-Bikonditionale von Literalen desselben Faktors in derselben Konjunktion enthalten sind
 - 2: **for each** Konjunktion (von RDN-Bikonditionalen) Γ aus $\mathbf{B}_n$ **do**
 - 3: **if** Γ ist nicht äquivalent zur Konjunktion aller RDN-Bikonditionale $\Phi = \beta_1 \cdot \beta_2 \cdot \dots$ **then**
 - 4: Entferne Γ aus $\mathbf{B}_n$
 - 5: **repeat for** $k = n, n - 1, \dots$
 - 6: $\mathbf{B}_{k-1} := \emptyset$
 - 7: $\mathbf{C}_k := \emptyset$
 - 8: **for each** Konjunktion (von RDN-Bikonditionalen) Γ aus $\mathbf{B}_k$ **do**
 - 9: **if** $\nexists \beta_j$ aus $\Gamma = \Gamma' \cdot \beta_j$, sodass Γ' äquivalent ist zu Γ **then**
 - 10: Füge Γ zu $\mathbf{C}_k$ hinzu
 - 11: **else**
 - 12: **while** $\exists \beta_j$ aus $\Gamma = \Gamma' \cdot \beta_j$, wobei Γ' äquivalent ist zu Γ **do**
 - 13: Füge Γ' zu $\mathbf{B}_{k-1}$ hinzu
 - 14: Entferne alle identischen Elemente aus $\mathbf{B}_{k-1}$ bis auf jeweils eines
 - 15: **until** $\mathbf{B}_{k-1} = \emptyset$
 - 16: **return** $\mathbf{T} = \mathbf{C}_n \cup \mathbf{C}_{n-1} \cup \dots$
-

5.4. Bildung von Kausalmodellen

Im dritten Hauptschritt werden die Kausalmodelle $\omega_1, \dots, \omega_k$ gebildet, die mit der untersuchten homogenen und maximal diversen Konfigurationstabelle verträglich sind. Jede aus dem vorher-

gehenden Schritt resultierende Minimale Theorie Ψ_i ist Teil eines solchen Kausalmodells ω_i . Um ein Kausalmodell ω_i vollständig zu beschreiben, muss die jeweilige Minimale Theorie Ψ_i von $\mathcal{K}_F$ allenfalls noch um einfache Minimale Theorien zu einem System von Minimalen Theorien ergänzt werden. Dies ist dann notwendig, wenn Ψ_i Folgen von Literalen aufweist, die auf eine kausale Verkettung hindeuten. In diesem Fall müssen die indirekten Kausalzusammenhänge aufgrund der geltenden Intransitivität bestätigt werden oder entsprechende Switching-Literale als solche ausgezeichnet werden.

Für jede aus Hauptschritt (II) resultierende Minimale Theorie Ψ_i der untersuchten Konfigurationstabelle $\mathcal{K}_F$ wird in einem ersten Teilschritt untersucht, ob sie komplex ist oder komplexe Minimale Theorien enthält.¹³ Gibt es komplexe Minimale Theorien, werden in einem nächsten Teilschritt allfällige indirekte Zusammenhänge bestätigt oder Switching-Literale ermittelt. Die abschliessende kausale Interpretation eines resultierenden Systems $\omega_i = [\Psi_i, \beta_1, \dots, \beta_l, \mathbf{F}^s]$, bestehend aus einer Minimalen Theorie Ψ_i von $\mathcal{K}_F$ und allenfalls aus indirekte Beziehungen bestätigenden einfachen Minimalen Theorien $\beta_1, \dots, \beta_l$, $l \geq 0$, sowie der Menge der Switching-Literale $\mathbf{F}^s$, entspricht einem Kausalmodell, das mit $\mathcal{K}_F$ verträglich ist. Ein solches Kausalmodell wird abschliessend üblicherweise mittels kausal wohlgeformten Hypergraphen dargestellt.

5.4.1. Ermittlung komplexer Minimaler Theorien

Für jede aus Schritt (II) resultierende Minimale Theorie Ψ_i wird geprüft, ob darin einfache Minimale Theorien $\beta_1, \dots, \beta_k$ enthalten sind, die Literale vom gleichen Faktor aufweisen. Ist dem so, werden die entsprechenden Konjunktionen von einfachen Minimalen Theorien zur Kennzeichnung mit einem eckigen Klammerpaar umrahmt. Gibt es zwei Minimale Theorien, die beide Literale vom gleichen Faktor enthalten, wird verkürzt gesagt, dass sie diesen Faktor gemeinsam haben.

Ist $\Psi_i = \beta_1 \cdot \dots \cdot \beta_k$ eine aus k RDN-Bikonditionalen bestehende Minimale Theorie der untersuchten Konfigurationstabelle $\mathcal{K}_F$, werden die in Ψ_i enthaltenen komplexen Minimalen Theorien rekursiv über die einfachen Minimalen Theorien $\beta_1, \dots, \beta_k$, die alle in Ψ_i enthalten sind, aufgebaut. Dazu bildet man zunächst die Menge $\mathbf{S}$ aller einfachen Minimalen Theorien aus Ψ_i . Aus $\mathbf{S}$ wird eine beliebige einfache Minimale Theorie β_i , $1 \leq i \leq k$, ausgewählt, aus $\mathbf{S}$ entfernt und

¹³Zur Erinnerung: Eine komplexe Minimale Theorie ist eine Konjunktion von einfachen Minimalen Theorien, die Faktorüberschneidungen aufweisen (vgl. Def. 2.16, S. 68). Eine Minimale Theorie von $\mathcal{K}_F$ muss daher nicht unbedingt als Ganzes komplex sein. Es kann auch nur ein Teil aller darin enthaltenen einfachen Minimalen Theorien Faktorüberschneidungen aufweisen. Ist letzteres der Fall wird gesagt, dass die Minimale Theorie von $\mathcal{K}_F$ eine komplexe Minimale Theorie enthält (vgl. S. 68).

überprüft, ob es eine andere einfache Minimale Theorie β_j aus $S_1 = S \setminus \{\beta_i\}$ gibt, mit der β_i mindestens einen Faktor gemeinsam hat. Findet man eine solche einfache Minimale Theorie β_j , wird diese mit β_i innerhalb eines eckigen Klammerpaars konjunktiv zu $[\beta_i \cdot \beta_j]$ verknüpft und ebenfalls aus S entfernt. Für den Ausdruck $[\beta_i \cdot \beta_j]$ wird der Schritt wiederholt, indem geprüft wird, ob es eine einfache Minimale Theorie in $S \setminus \{\beta_i, \beta_j\}$ gibt, die mindestens einen Faktor mit $[\beta_i \cdot \beta_j]$ gemeinsam hat. Letzteres gilt genau dann, wenn die einfache Minimale Theorie mit β_i oder β_j (oder beiden) einen Faktor gemeinsam hat. Sobald man eine solche einfache Minimale Theorie β_l findet, wird der Ausdruck $[\beta_i \cdot \beta_j]$ innerhalb der eckigen Klammern konjunktiv um β_l zu $[\beta_i \cdot \beta_j \cdot \beta_l]$ erweitert und auch β_l aus S entfernt. Der Schritt wird so lange wiederholt, bis (a) es keine einfache Minimale Theorie mehr gibt, die mit der von eckigen Klammern umrahmten Konjunktion von einfachen Minimalen Theorien einen Faktor gemeinsam hat oder (b) keine einfache Minimale Theorie mehr in der stetig kleiner werdenden Menge S enthalten ist. Gilt nur (a), sind also nach wie vor einfache Minimale Theorien in der reduzierten Menge enthalten, wählt man wiederum eine beliebige einfache Minimale Theorie aus dieser reduzierten Menge und startet den obigen rekursiven Prozess von Neuem. Sobald (b) gilt, wird der Prozess abgeschlossen und alle mit einem eckigen Klammerpaar umrahmten Konjunktionen von einfachen Minimalen Theorien sowie die allenfalls einfachen Minimalen Theorien, die keine Faktoren mit anderen einfachen Minimalen Theorien gemeinsam haben, konjunktiv miteinander verknüpft.

Dieser rekursive Prozess soll nochmals konkret anhand der Minimalen Theorie (5.68) von $\mathcal{K}_{5,24}$ veranschaulicht werden. Zuerst wird die aus den einfachen Minimalen Theorien bestehende Menge

$$S = \{(E + F \leftrightarrow H), (A + \bar{B}\bar{C} \leftrightarrow E), (D + BC \leftrightarrow F), (\bar{C} + \bar{D} \leftrightarrow G)\} \quad (5.80)$$

gebildet. Es wird die einfache Minimale Theorie $(E + F \leftrightarrow H)$ gewählt, aus S entfernt und überprüft, ob es eine einfache Minimale Theorie aus $S \setminus \{(E + F \leftrightarrow H)\}$ gibt, die einen Faktor mit $(E + F \leftrightarrow H)$ gemeinsam hat. Mit $(A + \bar{B}\bar{C} \leftrightarrow E)$ findet man eine solche, die ebenfalls E enthält. Auch diese einfache Minimale Theorie wird demnach aus S entfernt und innerhalb eines eckigen Klammerpaars mit $(E + F \leftrightarrow H)$ konjunktiv verknüpft zu:

$$[(E + F \leftrightarrow H) \cdot (A + \bar{B}\bar{C} \leftrightarrow E)] \quad (5.81)$$

Aus der übrigbleibenden Menge $\{(D + BC \leftrightarrow F), (\bar{C} + \bar{D} \leftrightarrow G)\}$ wird $(D + BC \leftrightarrow F)$ gewählt und geprüft, ob sie einen Faktor gemeinsam hat mit der Konjunktion aus (5.81), was unter anderem für den Faktor F zutrifft. Die Konjunktion aus (5.81) wird konjunktiv erweitert zu:

$$[(E + F \leftrightarrow H) \cdot (A + \bar{B}\bar{C} \leftrightarrow E) \cdot (D + BC \leftrightarrow F)] \quad (5.82)$$

Der Prozess wird mit der letzten in S enthaltenen einfachen Minimalen Theorie $(\bar{C} + \bar{D} \leftrightarrow G)$ abgeschlossen, die mit der Konjunktion aus (5.82) die Faktoren C und D gemeinsam hat. Es

zeigt sich, dass die Minimale Theorie (5.68) eine komplexe Minimale Theorie darstellt, was durch das den gesamten Ausdruck umrahmende eckige Klammerpaar direkt sichtbar gemacht wird:

$$\left[(E + F \leftrightarrow H) \cdot (A + B\bar{C} \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \right] \quad (5.83)$$

Auch für alle anderen Minimalen Theorien stellt sich heraus, dass sie komplexe Minimale Theorien von $\mathcal{K}_{5,24}$ sind.

Algorithmus 5.7 fasst die Vorgehensweise nochmals mittels Pseudocode zusammen.

Algorithmus 5.7 Suche nach komplexen Minimalen Theorien

Input: Minimale Theorie Ψ von $\mathcal{K}_F$, bestehend aus den einfachen Minimalen Theorien $\beta_1, \beta_2, \dots$

Output: Minimale Theorie Ψ mit durch eckigen Klammern ausgezeichneten komplexen Minimalen Theorien

```

1: K :=  $\emptyset$ 
2: Bilde die Menge S bestehend aus den einfachen Minimalen Theorien  $\beta_1, \beta_2, \dots$ 
3: while S  $\neq \emptyset$  do
4:   Wähle ein  $\beta_i \in \mathbf{S}$  und entferne es aus S
5:   if  $\nexists \beta_j \in \mathbf{S}$ , sodass  $\beta_i$  und  $\beta_j$  einen Faktor gemeinsam haben then
6:     Füge  $\beta_i$  zu K hinzu
7:   else
8:      $\Gamma := \beta_i$ 
9:     while  $\exists \beta_j \in \mathbf{S}$ , sodass  $\Gamma$  und  $\beta_j$  einen Faktor gemeinsam haben do
10:      Bilde  $\Gamma := [\Gamma \cdot \beta_j]$  und entferne  $\beta_j$  aus S
11:    Füge  $\Gamma$  zu K hinzu
12: return Konjunktive Verknüpfung aller Elemente aus K

```

5.4.2. Bestätigung indirekter Zusammenhänge

Ist eine Minimale Theorie Ψ_i der untersuchten Konfigurationstabelle $\mathcal{K}_F$ komplex oder enthält sie komplexe Minimale Theorien, stellt sich abschliessend die Frage, ob diese komplexen Minimalen Theorien Folgen von Literalen aufweisen, die eine Kausalkette nahelegen. Ist dem so, gilt es aufgrund der Intransitivität allfällige indirekte Zusammenhänge noch zu bestätigen oder Switching-Literale als solche auszuzeichnen.

Zuerst wird nach Folgen von Literalen gesucht, die auf kausale Verkettungen hindeuten. Dazu werden in einem ersten Teilschritt alle Literale der komplexen Minimalen Theorien ermittelt,

von denen keine einfachen Minimale Theorien in der komplexen Minimalen Theorie enthalten sind, selbst aber in einfachen Minimalen Theorien von Literalen anderer Faktoren vorkommen. Sofern eine komplexe Minimale Theorie tatsächlich eine Folge von Literalen aufweisen sollte, die auf eine kausale Verkettung hindeutet, sind es diese Literale, die am Ausgangspunkt einer Verkettung stehen. Solche Literale beziehungsweise die entsprechenden Faktoren werden als *Wurzelfaktorlitterale* beziehungsweise Wurzelfaktoren bezeichnet.¹⁴ Ausgehend von diesen Literalen werden allfällige, auf Verkettungen hindeutende Folgen von Literalen Schritt für Schritt aufgebaut. Ist X_i ein Wurzelfaktorliteral einer komplexen Minimalen Theorie, wird jedes Paar (X_i, X_j) gebildet, für das gilt, dass X_i in einer einfachen Minimalen Theorie von X_j enthalten ist. Es kann jeweils deshalb mehrere solche Paare geben, weil es sich bei X_i auch um eine gemeinsame Ursache mehrerer Wirkungen handeln kann. Dieser Schritt wird nun für das jeweilige Literal X_j eines Paares (X_i, X_j) wiederholt und geprüft, ob es einfache Minimale Theorien gibt, die X_j in ihrer RDN-Bedingung enthalten. Für jedes gefundene Literal X_k , das diese Bedingung erfüllt, wird das entsprechende Paar um ein Folglied zu (X_i, X_j, X_k) erweitert. Dieser Erweiterungsschritt wird für jede Folge von Literalen $p_l = (X_1, \dots, X_l)$ solange wiederholt, bis es für das jeweilige Literal X_l am Ende der Folge entweder (a) keine einfache Minimale Theorie eines Literals mehr gibt, die X_l enthält, oder (b) X_l nur noch Teil von einfachen Minimalen Theorien von Literalen ist, die selbst bereits in der Folge p_l enthalten sind. Gilt (a), kann der Prozess für diese Folge $p_l = (X_1, \dots, X_l)$ abgebrochen werden und p_l als Folge von Literalen der Länge l ausgegeben werden, die auf eine kausale Verkettung hindeutet. Gilt (b), deutet die Folge p_l auf eine geschlossene Verkettung hin, die eine kausale Feedbackstruktur beschreibt. Auch in diesem Fall kann der Prozess für die entsprechende Folge $p_l = (X_1, \dots, X_l)$ abgebrochen werden und p_l als Folge von Literalen ausgegeben werden, die auf eine kausale Verkettung hindeutet. Andernfalls würde obiger Prozess in diesem Fall nicht stoppen und eine endlose sich wiederholende Folge generieren. Indem man sich bei diesem Schritt nur für die Folgen von Literalen interessiert, die auf eine Verkettung mit indirekten Relevanzen hindeuten, werden in einem letzten Schritt alle Folgen gelöscht, die aus nur zwei Literalen bestehen.

Es ist wichtig nochmals explizit hervorzuheben, dass es sich bei den obigen Folgen um Folgen von Literalen handelt. Im Gegensatz zur Bildung von komplexen Minimalen Theorien, bei denen einfache Minimale Theorien nur auf gemeinsame Faktoren geprüft werden und noch nicht zwischen positiven und negativen Literalen unterschieden wird, muss letzteres bei auf Verkettungen hindeutenden Folgen auch mitberücksichtigt werden. Ist beispielsweise A in einer einfachen Minimalen Theorie von B und $\bar{B}$ in einer einfachen Minimalen Theorie von C , bilden diese drei Literale keine Folge, die auf eine kausale Verkettung hindeutet. Kausal interpretiert kann aus diesen Regelmässigkeiten keine indirekte kausale Relevanz von A für C hervorgehen, zumal A

¹⁴Eine alternative gebräuchliche Terminologie ist jene des „exogenen“ Literals beziehungsweise des exogenen Faktors.

Schritt 1	Schritt 2	Schritt 3
A	(A, E)	(A, E, H)
B	(B, E)	(B, E, H)
	(B, F)	(B, F, H)
C	(C, F)	(C, F, H)
$\bar{C}$	$(\bar{C}, E)$	$(\bar{C}, E, H)$
	$(\bar{C}, G)$	–
D	(D, F)	(D, F, H)
$\bar{D}$	$(\bar{D}, G)$	–

Tabelle 5.25.: Nach Schritt 3 resultierende, in der komplexen Minimalen Theorie aus (5.68) von $\mathcal{K}_{5,24}$ enthaltene Folgen von Literalen (dritte Spalte), die auf eine kausale Verkettung hindeuten.

eine Ursache für den anwesenden Faktor B ist, dessen Abwesenheit wiederum kausal relevant für die Anwesenheit von C ist. Eine auf eine Verkettung hindeutende Folge würde sich etwa dann bilden lassen, wenn B als positives Literal auch Teil einer einfachen Minimalen Theorie von C ist, sodass sich die Folge (A, B, C) ergibt. Oder es ergibt sich die auf eine Verkettung hindeutende Folge $(A, \bar{B}, C)$, sofern A in einer einfachen Minimalen Theorie von $\bar{B}$ enthalten ist.

Am Beispiel der komplexen Minimalen Theorie aus (5.68) von $\mathcal{K}_{5,24}$ veranschaulicht, lässt sich das obige Vorgehen mittels Tabelle 5.25 zusammenfassen. In einem ersten Schritt werden die Wurzelfaktorlitterale $A, B, C, \bar{C}$ sowie D und $\bar{D}$ ermittelt, die in der komplexen Minimalen Theorie aus (5.68) vorkommen, von denen aber keine einfachen Minimalen Theorien in (5.68) enthalten sind. Für jedes dieser Literale werden in einem zweiten Schritt alle Paare mit jenen Literalen gebildet, in deren einfachen Minimalen Theorien die Wurzelfaktorlitterale enthalten sind. Für B und $\bar{C}$ gibt es dabei zwei solche Paare, da beide Literale jeweils in zwei einfachen Minimalen Theorien von anderen Literalen enthalten sind. In einem dritten Schritt können bei den aus dem zweiten Schritt resultierenden Paaren, an deren zweiter Stelle E oder F steht, noch H hinzugefügt werden, da die einfache Minimale Theorie von H die beiden Literale E und F enthält. Die beiden aus dem zweiten Schritt resultierenden Paare $(\bar{D}, G)$ und $(\bar{C}, G)$ können nicht weiter ergänzt werden, weil G in keiner einfachen Minimalen Theorie eines Literals enthalten ist. Diese aus nur zwei Elementen bestehenden Folgen werden nicht weiter betrachtet, da sie trivialerweise nicht als Verkettungen interpretiert werden können, die indirekte Kausalzusammenhänge beinhalten. Insgesamt ergeben sich damit für die komplexe Minimale Theorie aus (5.68) die sechs Folgen von Literalen aus der dritten Spalte der Tabelle 5.25, die auf kausale Verkettungen hindeuten.

Algorithmus 5.8 fasst das Vorgehen allgemein mittels Pseudocode zusammen. Zwecks einer übersichtlicheren Schreibweise wird dazu die Menge **E** definiert, die jedes Paar (Y, Z) enthält, wobei jeweils gilt, dass Y in einer einfachen Minimalen Theorie von Z enthalten ist.

Algorithmus 5.8 Suche nach auf kausale Verkettungen hindeutenden Folgen von Faktoren

Input: Minimale Theorie Ψ von $\mathcal{K}_F$ mit allfälligen, durch eckige Klammerpaare umrahmten komplexen Minimalen Theorien

Output: Menge **P** von in Ψ enthaltenen, auf kausale Verkettungen hindeutenden Folgen p_k von Literalen der Länge $k > 2$

```

1: if  $\Psi$  enthält keine komplexen Minimalen Theorien then
2:   P =  $\emptyset$ 
3: else
4:   for each Komplexe Minimale Theorie aus  $\Psi$  do
5:     Bilde die Menge E :=  $\{(Y, Z) | Y \text{ ist Teil einer einfachen Minimalen Theorie von } Z\}$ 
6:     Bilde die Menge W1 :=  $\{(X) | \exists Z, \text{ sodass } (X, Z) \in \mathbf{E} \text{ und } \nexists Y, \text{ sodass } (Y, X) \in \mathbf{E}\}$ 
7:     repeat for  $k = 1, 2, \dots$ 
8:       W $k+1$  :=  $\emptyset$ 
9:       P $k$  :=  $\emptyset$ 
10:      for each  $p_k = (X_1, \dots, X_k)$  aus W $k$  do
11:        if  $\exists X_j$ , sodass  $(X_k, X_j) \in \mathbf{E}$  und  $X_j$  ist nicht in  $p_k$  enthalten then
12:          Füge  $p_{k+1} := (X_1, \dots, X_k, X_j)$  zu W $k+1$  hinzu
13:        else if  $\exists X_j$ , sodass  $(X_k, X_j) \in \mathbf{E}$  und  $X_j$  ist in  $p_k$  enthalten then
14:          Füge  $p_k := (X_1, \dots, X_k)$  zu P $k$  hinzu
15:        else
16:          Füge  $p_k$  zu P $k$  hinzu
17:      until W $k+1$  =  $\emptyset$ 
18:    P = P1  $\cup$  P2  $\cup \dots$ 
19:    Lösche alle  $p_k$  aus P mit  $k \leq 2$ 
20: return P

```

Für jede auf eine kausale Verkettung hindeutende Folge muss abschliessend geprüft werden, ob sich die allfällig daraus folgenden indirekten Zusammenhänge im entsprechenden Kausalmodell aufgrund der Intransitivität auch tatsächlich bestätigen lassen. Gibt es eine Folge von Literalen $(\dots, X_i, \dots, X_j, \dots, X_k, \dots)$, die auf ein Kausalmodell ω_i mit indirekter Relevanz von X_i für X_k hinweist, muss aufgrund der Intransitivität bestätigt werden, dass X_i auch in einer einfachen Minimalen Theorie von X_k enthalten ist. Gibt es eine solche einfache Minimale Theorie β_m , bildet diese zusammen mit der Minimalen Theorie Ψ_i von $\mathcal{K}_F$ ein System $\omega_i = [\Psi_i; \beta_m]$ von Minimalen

Theorien. Um Verwechslungen zu vermeiden, wird dabei der Äquivalenzoperator der einfachen Minimalen Theorie β_m , die den indirekten Zusammenhang bestätigt, um ein „i“ ergänzt. Gibt es andererseits keine solche einfache Minimale Theorie, ist X_i ein Switching-Literal und wird als solches in die Menge $\mathbf{F}^s$ aufgenommen.

Bezüglich der Bestätigung einer indirekten Relevanz eines Literals X_i für ein Literal X_k ist es im Grunde unwesentlich, welche einfachen Minimalen Theorien zur Bestätigung herangezogen werden, zumal einzig die Information wesentlich ist, dass X_i in einer einfachen Minimalen Theorie von X_k enthalten ist. Damit das System von Minimalen Theorien übersichtlich bleibt, ist es jedoch sinnvoll zu versuchen möglichst wenige zusätzliche einfache Minimalen Theorien, die die indirekten Zusammenhänge bestätigen, in das System aufzunehmen. Daraus folgt, dass man zur Bestätigung einer indirekten Relevanz von X_i für X_k zuerst prüft, ob bereits eine für einen anderen indirekten Zusammenhang herangezogene einfache Minimale Theorie im System enthalten ist, die auch X_i als Difference-Maker von X_k auszeichnet.

Weiter kann es vorkommen, dass ein Literal X_i bereits im Zusammenhang mit einer anderen Folge von Literalen auf dessen indirekte Relevanz für ein anderes Literal X_j geprüft und das System von Minimalen Theorien entsprechend ergänzt wurde. Dies kann etwa dann der Fall sein, wenn eine Kausalstruktur mehrere Verkettungen aufweist, die identische Teile enthalten. Mit letzterem ist gemeint, dass beispielsweise mehrere Verkettungen bei einem Literal zusammenlaufen und danach (teilweise) dem gleichen Kausalpfad folgen. Aus Effizienzgründen ist es demnach sinnvoll, vor der Prüfung und allfälligen Bestätigung eines indirekten Zusammenhangs zwischen zwei Literalen X_i und X_j zu prüfen, ob das entsprechende Paar (X_i, X_j) bereits innerhalb eines anderen Pfades geprüft wurde. Auf diese Weise werden redundante Schritte vermieden und die Prüfung nur einmal durchgeführt.

Für die Minimale Theorie aus (5.68) von $\mathcal{K}_{5.24}$ über die Faktormenge $\mathbf{M} = \{A, B, C, D, E, F, G, H\}$ kann etwa das System $\omega_1 = [(5.83); (5.49); \mathbf{M}^s = \{C\}]$ gebildet werden. Bei (5.83) und (5.49) handelt es sich dabei um die folgenden Minimalen Theorien:

$$\left[(E + F \leftrightarrow H) \cdot (A + B\bar{C} \leftrightarrow E) \cdot (D + BC \leftrightarrow F) \cdot (\bar{C} + \bar{D} \leftrightarrow G) \right] \quad (5.83)$$

$$A + B + D \overset{i}{\leftrightarrow} H \quad (5.49)$$

C stellt sich als Switching-Literal heraus, da es keine einfache Minimale Theorie von H gibt, die C enthält. Als Kausalgraph dargestellt, ergibt sich damit das mit $\mathcal{K}_{5.24}$ verträgliche Kausalmodell ω_1 , das der datengenerierenden Kausalstruktur aus Abbildung 5.24 entspricht. Zusammen mit den Kausalmodellen ω_2 bis ω_{12} , auf deren konkrete Ausarbeitung aus Platzgründen an dieser Stelle verzichtet wird, die jedoch analog via des in diesem Abschnitt beschriebenen dritten

Hauptschritts für die anderen Minimalen Theorien aus (5.69) bis (5.79) resultieren, liegen alle Kausalmodelle vor, die mit $\mathcal{K}_{5,24}$ verträglich sind.

Dieser letzte Schritt zur Bestätigung indirekter Zusammenhänge beziehungsweise zur Ermittlung von Switching-Literalen wird in Algorithmus 5.9 zusammengefasst. Dabei sind ψ_{X_j} und ϕ_{X_j} RDN-Bedingungen von X_j und bilden über das Bikonditional mit X_j verknüpft einfache Minimale Theorien von X_j . Es gilt zu erwähnen, dass dieser Algorithmus nicht zwangsläufig die gewünschte minimale Anzahl einfacher Minimaler Theorien heranzieht, die für die Bestätigung von indirekten Zusammenhängen notwendig sind.

Algorithmus 5.9 Bestätigung von in der Minimalen Theorie Ψ enthaltenen indirekten Zusammenhängen

Input: Minimale Theorie Ψ von $\mathcal{K}_F$, Menge Δ aller einfachen Minimalen Theorien relativ zu $\mathcal{K}_F$ und Menge $\mathbf{P}$ von auf Verkettungen hindeutenden Folgen $p_1, p_2, \dots$ aus Ψ

Output: Menge $\mathbf{Z}$ von einfachen Minimalen Theorien, die indirekte Zusammenhänge bestätigen sowie Menge $\mathbf{F}^s$ der Switching-Literale

```

1: if  $\mathbf{P} = \emptyset$  then
2:    $\mathbf{F}^s = \emptyset$  und  $\mathbf{Z} = \emptyset$ 
3: else
4:    $\mathbf{P}^* := \emptyset$ 
5:    $\mathbf{F}^s := \emptyset$ 
6:    $\mathbf{Z} := \emptyset$ 
7:   repeat for  $k = 3, 4, \dots$ 
8:     for each  $p_k$  aus  $\mathbf{P}$  do
9:       for each  $X_i$  und  $X_j$  aus  $p_k$ , wobei  $i < i + 1 < j$  sowie  $(X_i, X_j)$  ist keine Teilfolge
         einer Folge aus  $\mathbf{P}^*$  do
10:        if  $X_i$  ist in einer einfachen Minimalen Theorie  $\phi_{X_j} \overset{i}{\leftrightarrow} X_j$  aus  $\mathbf{Z}$  then
11:          Behalte  $\phi_{X_j} \overset{i}{\leftrightarrow} X_j$  unverändert in  $\mathbf{Z}$ 
12:        else if  $X_i$  ist in einer einfachen Minimalen Theorie  $\psi_{X_j} \leftrightarrow X_j$  aus  $\Delta$  then
13:          Füge  $\psi_{X_j} \overset{i}{\leftrightarrow} X_j$  zu  $\mathbf{Z}$  hinzu
14:        else
15:          Füge  $X_i$  zu  $\mathbf{F}^s$  hinzu, sofern  $X_i$  noch nicht in  $\mathbf{F}^s$  enthalten ist
16:        Füge  $p_k$  zu  $\mathbf{P}^*$  hinzu
17:   until  $\mathbf{P}^* = \mathbf{P}$ 
18: return  $\mathbf{Z}$  und  $\mathbf{F}^s$ 

```

5.5. csCNA+ in Kürze

Ausgegangen wird von einer homogenen und maximal diversen Konfigurationstabelle $\mathcal{K}_F$ über $F = \{X_1, \dots, X_n\}$. Gesucht wird jedes Kausalmodell ω_i , das mit $\mathcal{K}_F$ verträglich ist. Die folgenden Schritte fassen das dazu genutzte Aufdeckungsverfahren zusammen:

(I) Bestimmung aller RDN-Bikonditionale: Ermittle die Menge F^* von Faktoren aus F gemäss Algorithmus 5.1 (vgl. S. 182), für die weder (WD1) noch (WD2) gilt. Von den Literalen dieser Faktoren aus F^* werden die RDN-Bikonditionale gesucht. Bilde dazu für jeden Faktor X_j aus F^* die innere Instanzfunktion $f_{X_j} : \mathcal{D} \rightarrow \mathcal{B}$ mit $\mathcal{B} = \{0, 1\}$ und $\mathcal{D} \subseteq \mathcal{B}^{n-1}$:

$$f_{X_j}(X_1, \dots, X_{j-1}, X_{j+1}, \dots, X_n) = \begin{cases} 1, & (X_1, \dots, X_{j-1}, 1, X_{j+1}, \dots, X_n) \in \mathcal{I}(f) \\ 0, & (X_1, \dots, X_{j-1}, 0, X_{j+1}, \dots, X_n) \in \mathcal{I}(f) \end{cases} \quad (5.84)$$

Die Menge $\mathcal{I}(f)$ ist dabei der Instantiierungsbereich der totalen Instanzfunktion $f : \mathcal{B}^n \rightarrow \mathcal{B}$, die das durch $\mathcal{K}_F$ beschriebene Instantiierungsverhalten der Faktoren aus F beschreibt. Das heisst, die Belegungen aus $\mathcal{I}(f)$ entsprechen den Konfigurationen aus $\mathcal{K}_F$ – kurz notiert: $\mathcal{I}(f) = \mathcal{K}_F$.

Finde anschliessend folgendermassen die RDN-Formen von jeder inneren Instanzfunktion f_{X_j} :

- (a) Bestimme die Primimplikanten von f_{X_j} gemäss Algorithmus 5.4 (vgl. S. 211).
- (b) Verknüpfe die aus (a) resultierenden Primimplikanten disjunktiv zu irredundanten Abdeckungen (RDN-Formen) von f_{X_j} gemäss Algorithmus 5.3 (vgl. S. 203).

Jede daraus resultierende RDN-Form ϕ_{X_j} einer inneren Instanzfunktion f_{X_j} wird als RDN-Bikonditional $\phi_{X_j} \leftrightarrow X_j$ in die Menge $\mathbf{A}$ aller RDN-Bikonditionale, die von $\mathcal{K}_F$ impliziert werden, aufgenommen.

(II) Bildung von Minimalen Theorien: Aus den RDN-Bikonditionalen aus $\mathbf{A}$ werden alle Minimalen Theorien von $\mathcal{K}_F$ gemäss Algorithmus 5.6 gebildet (vgl. S. 225). Bilde dazu zuerst alle Konjunktionen von RDN-Bikonditionalen, wobei für jede Konjunktion gilt, dass von jedem Faktor, von dessen Literalen RDN-Bikonditionale aus $\mathcal{K}_F$ folgen, genau ein RDN-Bikonditional darin enthalten ist. Prüfe die resultierenden Konjunktionen auf Äquivalenz zur Konjunktion Φ aller in $\mathbf{A}$ enthaltenen RDN-Bikonditionale und befreie jene, die äquivalent sind zu Φ von Redundanzen.

Die Menge aller RDN-Bikonditionale aller Minimalen Theorien von $\mathcal{K}_F$ entspricht der Menge aller einfachen Minimalen Theorien relativ zu $\mathcal{K}_F$.

(III) Bildung von Kausalmodellen: Ausgehend von den n Minimalen Theorien $\Psi_1, \dots, \Psi_n$ von $\mathcal{K}_F$ aus Schritt (II) werden alle n Kausalmodelle $\omega_1, \dots, \omega_n$ gebildet, die mit $\mathcal{K}_F$ verträglich sind. Jedes dieser Kausalmodelle ω_i besteht aus einer Minimalen Theorie Ψ_i und gegebenenfalls aus einfachen Minimalen Theorien (relativ zu $\mathcal{K}_F$) $\beta_1, \dots, \beta_l$, $l \geq 0$, die indirekte Zusammenhänge bestätigen. Enthält ω_i auch Switching-Literale, wird das System von Minimalen Theorien um die Menge F^s dieser Switching-Literale ergänzt. Etwas konkretisiert wird jedes Kausalmodell $\omega_i := [\Psi_i; \beta_1, \dots, \beta_l; F^s]$ folgendermassen gebildet:

- (a) Ermittle die komplexen Minimalen Theorien aus Ψ_i gemäss Algorithmus 5.7 (vgl. S. 228).
- (b) Suche nach Folgen von Literalen in jeder komplexen Minimalen Theorie aus Ψ_i gemäss Algorithmus 5.8 (vgl. S. 231), die auf kausale Verkettungen hindeuten. Bestätige allfällige, sich daraus ergebende indirekte Zusammenhänge durch die Hinzunahme von entsprechenden einfachen Minimalen Theorien oder bestimme die Menge F^s der Switching-Literale gemäss Algorithmus 5.9 (vgl. S. 233).

Jedes daraus resultierende Kausalmodell $\omega_i := [\Psi_i; \beta_1, \dots, \beta_l; F^s]$ ist mit $\mathcal{K}_F$ verträglich und wird typischerweise mittels kausal wohlgeformten Hypergraphen als Kausalgraph dargestellt.

6. Fazit und Ausblick

Seit der Einführung von QCA durch Charles Ragin im Jahr 1987 wächst die Anzahl Studien stetig, die Kausalanalysen mittels konfiguralen kausalen Modellierens durchführen. Auch die Anwendungsgebiete umfassen mittlerweile viele wissenschaftliche Disziplinen, die von Politikwissenschaften bis hin zur Gesundheitsforschung reichen.¹ Die Grundidee des konfiguralen kausalen Modellierens ist, ausgehend von Konfigurationen auf notwendigen und hinreichenden Bedingungen basierende Regularitäten zu ermitteln, die kausal interpretierbar sind. Die vorliegende Arbeit hat sich eingehend mit diesem konfiguralen kausalen Modellieren für crisp-set Konfigurationstabellen auseinandergesetzt und alle für eine Kausalanalyse wesentlichen Aspekte analysiert, vom kausaltheoretischen Fundament über die empirische Grundlage bis hin zu Methoden zur Aufdeckung von Kausalstrukturen.

6.1. Fazit

In der Literatur findet man mit QCA und CNA zwei vielversprechende Verfahren zur Aufdeckung von Kausalstrukturen beliebiger Komplexität aus Konfigurationstabellen. Obwohl beide dasselbe Ziel verfolgen, unterscheiden sie sich in wesentlichen Punkten voneinander. Besonders hervorzuheben gilt es dabei zum einen die verschiedenen Verursachungsbegriffe, auf denen die beiden Verfahren aufgebaut sind. Zum anderen unterscheiden sie sich hinsichtlich der Methoden, die bei den entsprechenden Aufdeckungsverfahren eingesetzt werden. Unter Berücksichtigung dieser beiden Ansätze wurden in der vorliegenden Arbeit die folgenden drei für das konfigurale kausale Modellieren wesentlichen Themenfelder beziehungsweise Fragen behandelt:

Der Verursachungsbegriff: Was genau wird gesucht?

Die empirische Grundlage: Welche Informationen stehen dafür zur Verfügung?

¹Vgl. COMPASSS für eine Sammlung einer Vielzahl an entsprechenden Studien verschiedener Disziplinen: <http://www.compass.org/bibdata.htm>.

Das Aufdeckungsverfahren: Wie lässt sich basierend auf diesen Informationen das Gesuchte algorithmisch finden?

Die erste, in Kapitel 2 behandelte Frage zielt auf die Ausarbeitung jenes Verursachungsbegriffs ab, der bei auf Konfigurationstabellen basierenden Kausalanalysen zum Einsatz kommt. Es hat sich herausgestellt, dass die Kausaltheorien, auf denen bestehende Verfahren gegenwärtig basieren, ungenügend sind. QCA ist auf der INUS-Theorie aufgebaut und CNA auf deren Nachfolger, den Minimalen Theorien, ohne dabei die strukturellen Redundanzen zu berücksichtigen. Dieser ursprüngliche Begriff der Minimalen Theorie wurde in der vorliegenden Arbeit weiterentwickelt und darauf basierend ein stimmiges begriffliches Fundament erarbeitet, auf dem konfigurationale Methoden aufgebaut sind.

Die zweite der obigen drei Fragen betrifft die Daten, die für das konfigurationale kausale Modellieren zur Verfügung stehen. In Kapitel 3 wurde diese empirische Grundlage beleuchtet und es wurden neue Bedingungen formuliert, mit denen die Verhinderung von kausalen Fehlschlüssen sichergestellt wird. QCA hat sich dieser Fehlschlussproblematik noch gar nicht angenommen und bei CNA kommt eine Homogenitätsannahme zum Einsatz, die im Gegensatz zu jener aus der vorliegenden Arbeit strenger und damit weniger anwendbar als letztere ist.

Beim dritten und letzten der obigen drei Themenfelder geht es schliesslich um das Aufdeckungsverfahren, mit dessen Hilfe systematisch Kausalstrukturen aus crisp-set Konfigurationstabellen abgeleitet werden können. Dazu wurde in Kapitel 4 erstmals im Detail das mathematische Fundament des konfiguralen kausalen Modellierens im crisp-set Fall ausgearbeitet. Schliesslich wurde in Kapitel 5 das Verfahren csCNA+ entwickelt, das im Rahmen eines idealisierten Entdeckungskontextes in der Lage ist, Kausalstrukturen beliebiger Komplexität aufzudecken. Dieses Verfahren berücksichtigt erstens die kausaltheoretischen Weiterentwicklungen aus Kapitel 2. Ferner überwindet es unter Zuhilfenahme des mathematischen Fundaments Nachteile von QCA und CNA durch deren Verschmelzung. Konkret kommt csCNA+ einerseits, im Gegensatz zu csQCA, ohne Datenerweiterungen aus und andererseits verbessert es die hohe Ineffizienz von csCNA. Schliesslich enthält csCNA+ neu Algorithmen zur Bildung von Kausalmodellen aus Minimalen Theorien.

Der Verursachungsbegriff

Es gibt viele verschiedene Kausaltheorien, die unter anderem von probabilistischen Ansätzen über Transferenz- bis hin zu Interventionstheorien reichen. Bis heute hat sich weder eine dieser Kausaltheorien als einzige durchgesetzt, noch wurden die vielversprechendsten modernen Vari-

anten der verschiedenen Theorien zurückgewiesen. Dieser Umstand hat dazu geführt, dass die Position des kausalen Pluralismus populär geworden ist. Gemäss dieser Position widersprechen sich unterschiedliche Kausaltheorien nicht, sondern entwickeln unterschiedliche Verursachungsbegriffe, die unterschiedlichen vor-theoretischen Kausalintuitionen Rechnung tragen. Dementsprechend kann in einem Anwendungskontext eine Kausaltheorie angemessen sein und in einem anderen Anwendungskontext eine andere, wobei die Angemessenheit unter anderem davon abhängt, welche Informationen man für eine Analyse nutzt und was für Aussagen man damit treffen möchte.

Kausalität wird im Rahmen des konfiguralen kausalen Modellierens standardmässig im Sinne von auf Hume zurückgehenden Regularitätstheorien verstanden. Diese Theorien beschreiben Kausalzusammenhänge basierend auf notwendigen und hinreichenden Bedingungen. Moderne Regularitätstheorien haben solche Regularitäten um den wesentlichen Zusatz der Redundanzfreiheit ergänzt. Dieser Zusatz geht auf die Idee zurück, dass die auf hinreichenden und notwendigen Bedingungen basierenden Regelmässigkeiten nur dann kausal interpretiert werden können, wenn sie einen unverzichtbaren Beitrag zur Beschreibung des Instantiierungsverhaltens der betrachteten Faktoren leisten. Mackie (1980) hat diesen Zusatz in seiner bekannten INUS-Theorie eingearbeitet, gemäss der ein Literal X nur dann als kausal relevant für eine Wirkung W interpretiert werden kann, wenn X Teil einer minimal hinreichenden Bedingungen von W ist, die ihrerseits Teil einer notwendigen Bedingung (in disjunktiver Normalform) von W darstellt. Mackie weist seine INUS-Theorie jedoch sogleich wieder zurück, weil er der INUS-Theorie genügende Regularitäten ausmacht, die Kausalität nicht korrekt abbilden. Graßhoff & May (2001) überwinden diese als Manchester-Hooters-Problem bekannte Schwierigkeit, indem sie einen weiteren Typ von Redundanz identifizieren, der durch die zusätzliche Forderung der Minimalität der notwendigen Bedingung ausgemerzt wird. Die vorliegende Arbeit hat gezeigt, dass die sich daraus ergebenden, als RDN-Bikonditionale bezeichneten Regularitäten in einem direkten Zusammenhang mit den eingeführten Difference-Making-Paaren stehen. Konkret gibt es von jedem Literal X , das Teil eines RDN-Bikonditionals eines Literals W ist, ein Difference-Making-Paar bezüglich W . Dieser Zusammenhang formalisiert erstmals zu Teilen jene grundlegende Idee von Regularitätstheorien, dass jede Ursache in mindestens einem Kontext einen entscheidenden Unterschied betreffend die Anwesenheit der Wirkung macht.

Während QCA gegenwärtig die INUS-Theorie als kausaltheoretische Grundlage für Kausalanalysen nutzt, ist CNA auf RDN-Bikonditionalen aufgebaut. Obwohl CNA damit die Einwände überwindet, durch welche die INUS-Theorie nachweislich als unangemessene Theorie der Kausalität zurückgewiesen wurde, hat die vorliegende Arbeit aufgezeigt, dass dem Prinzip der Redundanzfreiheit jedoch auch mit den RDN-Bikonditionalen nicht konsequent genug Rechnung getragen wird. Bei der Betrachtung komplexer Kausalstrukturen mit mehreren Wirkun-

gen hat sich ein weiterer Redundanztyp offenbart, ohne dessen Berücksichtigung auch RDN-Bikonditionale Kausalzusammenhänge nicht korrekt einfangen. Es wurde festgestellt, dass es auch strukturelle Redundanzen gibt, gemäss denen RDN-Bikonditionale als Ganzes relativ zum Instantiierungsverhalten aller Faktoren redundant sein können. Darauf basierend wurde ein neuer Begriff der Minimalen Theorie entwickelt, gemäss dem die Konjunktion aller RDN-Bikonditionalen, die von einer Konfigurationstabelle $\mathcal{K}_F$ über eine Faktormenge F impliziert werden, strukturell minimal sein muss. Jedes einzelne RDN-Bikonditional einer solchen Minimalen Theorie von $\mathcal{K}_F$ entspricht dabei einer einfachen Minimalen Theorie und nur wenn ein Literal X Teil einer solchen einfachen Minimalen Theorie eines anderen Literals W ist, kann X kausal relevant für W sein. Zusammen mit der Vollständigkeits- und der bereits in Baumgartner (2008b) diskutierten Erweiterungsklausel wurde schliesslich ein neuer Begriff der kausalen Relevanz (eines Literals) entwickelt. Von diesem Begriff ausgehend wurden weitere Begriffe für komplexe Ursachen und Alternativursachen einer Wirkung ausgearbeitet. Obwohl dazu der Begriff der kausalen Relevanz nur geringfügig angepasst werden musste, blieb es in der bisherigen Literatur jeweils bei der Definition des Begriffs der kausalen Relevanz eines Literals, der isoliert eine Unterscheidung zwischen komplexen und alternativen Ursachen nicht zulässt.

Die sich daraus ergebende Kausaltheorie auf Typenebene fängt Kausalstrukturen ein, die den vier unter anderem in Graßhoff & May (2001) zu findenden fundamentalen Prinzipien genügen. Diese vier Prinzipien wurden um das Prinzip der kausalen Eindeutigkeit ergänzt. Das Eindeutigkeitsprinzip bringt zum Ausdruck, dass Regularitätstheorien nur auf Welten angewendet werden können, deren kausale Struktur eindeutig ist und dass unsere Welt so eine Welt ist. Damit werden Modellambiguitäten Rechnung getragen, die relativ zu einer Konfigurationstabelle entstehen können und die in den bestehenden Prinzipien unberücksichtigt blieben. Letzteres führte dazu, dass bisherige Regularitätstheorien gezwungen sein konnten, mehrere unterschiedliche Minimale Theorien kausal zu interpretieren, obwohl sich kausal wohlgeformte Modelle nur mit echten Teilen davon bilden lassen. Aus dem Prinzip der kausalen Eindeutigkeit folgt, dass solche Ambiguitäten jeweils auf Mängel beziehungsweise Lücken in den Daten zurückgehen und sich durch die Betrachtung erweiterter Faktormengen auflösen. Mit dem Prinzip geht ebenfalls ein gewisser minimaler Grad an kausaler Komplexität einher, den Kausalstrukturen aufweisen.

Durch den neuen Begriff der einfachen Minimalen Theorie hat sich ein neuer Begriff der komplexen Minimalen Theorie ergeben, der dazu genutzt wird, komplexe Kausalstrukturen mit mehreren Wirkungen, wie beispielsweise Common-Cause-Strukturen oder Kausalketten, einzufangen. Darauf basierend wurden schliesslich auch neue Begriffe für die direkte und indirekte kausale Relevanz präsentiert. In der bestehenden regularitätstheoretischen Literatur wurde eine begrifflich scharfe Unterscheidung zwischen direkter und indirekter Relevanz bisher nicht oder nur ungenügend gemacht. So blieben bei ersten Ansätzen, wie beispielsweise jenem in Baumgartner

(2008b), bei den Definitionen der indirekten Relevanz neben der strukturellen Redundanz auch die Intransitivität unberücksichtigt.

Schliesslich wurde im Detail aufgezeigt, dass das unter anderem in Baumgartner (2008a) diskutierte Kettenproblem ein Problem des kausalen Schliessens ist und kein prinzipielles Problem der Kausaltheorie darstellt. Die mit dem Kettenproblem einhergehenden Ambiguitäten lösen sich nach Erhärtung der Erweiterungsklausel auf.

Die Summe aller eingeführten Begriffe ergab eine neue Regularitätstheorie der Kausalität, die Kausalstrukturen beliebiger Komplexität einfängt. Neben dem Begriff der Kausalkette umfasst diese Theorie erstmalig auch eine Definition der Feedbackstruktur, zu der kausale Zyklen zählen. Regularitäten, die den Bedingungen dieser Theorie genügen, genügen auch den Bedingungen der Vorgängertheorien. Dementsprechend kann die vorliegende Theorie mit Einwänden, denen Regularitätstheorien seit Hume ausgesetzt sind und die von den Vorgängertheorien entkräftet werden konnten, analog umgehen wie ihre Vorgänger (vgl. u. a. Graßhoff & May (2001); Baumgartner (2008b)). Dazu gehören beispielsweise das oben erwähnte Manchester-Hooters-Problem oder das Problem von leeren oder Einzelfall-Regularitäten (vgl. u. a. Armstrong (1983), Kapitel 2).

Die empirische Grundlage

Konfigurationale kausale Methoden decken Kausalstrukturen basierend auf Konfigurationstabellen auf. Je nachdem wie solche Tabellen zustande kommen und welche Eigenschaften sie aufweisen, eignen sie sich besser oder weniger gut für eine Kausalanalyse. Im schlimmsten Fall produzieren Konfigurationen kausale Fehlschlüsse, indem in den Konfigurationen bestehende Regularitäten fälschlich kausal interpretiert werden. Nach einer kurzen Einführung, wie die für das konfigurationale kausale Modellieren zur Verfügung stehenden Informationen erhoben und zu Konfigurationstabellen aufbereitet werden, wurden im zweiten Teil des Kapitels 3 Bedingungen entwickelt und diskutiert, denen Konfigurationstabellen genügen müssen, um kausale Fehlschlüsse zu vermeiden. Ein Schwerpunkt wurde dabei auf den Begriff der Homogenität gelegt. Gerade in der QCA-Literatur wird dieser für das kausale Schliessen sehr wesentlichen Bedingung kaum Beachtung geschenkt, weil man sich der in der vorliegenden Arbeit präsentierten Fehlschluss-Problematik noch gar nicht angenommen hat. Mittels Szenarien wurde aufgezeigt, dass alleine durch geringfügige Variationen von unberücksichtigten Ursachen, sogenannten Störursachen, selbst Daten, die frei von praktischen Mängeln wie Messfehlern sind, kausale Fehlschlüsse produzieren können. Von Baumgartner (2009) ausgehend wurde neben einer Anpassung des Begriffs der potenziellen Störursache – der nun ebenfalls unberücksichtigte

Ko-Literale als potenzielle Störer ausweist – eine neue, auf dem Begriff der Variation beruhende Definition der Homogenität entwickelt. Letztere ist weniger streng als bisherige Ansätze, die bei CNA zu finden sind, schliesst aber nachweislich Fehlschluss provozierende Konfigurationstabellen ebenso aus. Sie ist in dem Sinne weniger streng, dass sie unberücksichtigte Faktoren in ihrem Instantiierungsverhalten viel weniger einschränkt und somit anwendbarer wird. Mit Hinzunahme des Begriffs der maximalen Diversität wurde schliesslich das letzte wesentliche und formal präzisierte Puzzleteil geliefert, um zu definieren, wie ideale Daten beim konfigurationsalen kausalen Modellieren aussehen.

Das Aufdeckungsverfahren csCNA+

Das Aufdeckungsverfahren csCNA+ wurde ausgehend von einer idealen Datenlage entwickelt, um sicherzustellen, dass allfällige Fehlschlüsse ausschliesslich auf ein Fehlverhalten des Verfahrens zurückgeführt werden können. Im Hinblick auf ein solches Verfahren wurde in Kapitel 4 zuerst das in der Literatur zwar immer erwähnte, aber nie im Detail ausgearbeitete mathematische Fundament des konfigurationsalen kausalen Modellierens im crisp-set Fall entwickelt. Es wurde plausibel gemacht, dass das konfigurationsale kausale Modellieren eine konkrete Anwendung von diesem als Konfigurationsalgebra bezeichneten Modell einer zweielementigen Booleschen Algebra ist. Die bestehende Literatur beruft sich zwar jeweils auf eine solche Boolesche Grundlage, eine entsprechende und umfassende mathematische Fundierung blieb sie jedoch bis heute schuldig.

Mit der Einführung der Konfigurationsalgebra und der damit einhergehenden Instanzfunktion konnte der Begriff der Minimalen Theorie explizit mit gängigen Begriffen aus der Booleschen Algebra ausgedrückt und so das konfigurationsale kausale Modellieren direkt für eine Vielzahl an etablierten Methoden anderer Anwendungen, wie beispielsweise der Schaltalgebra, zugänglich gemacht werden. Die Konfigurationsalgebra hat sich ebenfalls als starkes Instrument für die Kontrastierung und fruchtbare Verbindung jener Methoden herausgestellt, die bei den bereits existierenden Aufdeckungsverfahren csQCA und csCNA zum Einsatz kommen. In Kapitel 5 wurden beide Verfahren im Detail diskutiert und ihre grundlegenden Prinzipien ausgearbeitet. Mithilfe der Konfigurationsalgebra konnten die Verfahren, wie etwa anhand der kompakten Darstellung der Redundanzregel ersichtlich, direkt gegenübergestellt und so ihre Unterschiede präzise lokalisiert werden. Bei csQCA hat sich dabei im Fall von Kausalanalysen von Kausalstrukturen mit mehreren Wirkungen das unter anderem bereits in Baumgartner (2013a) erwähnte Problem der Datenerweiterungen gezeigt. CNA hat sich hingegen als wenig recheneffizient erwiesen. Die beiden Nachteile berücksichtigend, wurde schliesslich csCNA+ entwickelt. Dieses Verfahren orientiert sich stark an der crisp-set Variante von CNA, die ergänzt, recheneffizienter

gestaltet und weiterentwickelt wurde. Konkret berücksichtigt csCNA+ die strukturellen Redundanzen, von denen Regularitäten zusätzlich zu befreien sind, um kausal interpretiert werden zu können. Weiter wurden einzelne Teile aus Gründen der Recheneffizienz aus dem von QCA genutzten Quine-McCluskey-Verfahren übernommen. Schliesslich wurden dahingehend Ergänzungen vorgenommen, dass auch eine systematische Aufdeckung von Kausalmodellen mit darin enthaltenen indirekten Kausalzusammenhängen möglich ist. Dieser letzte Schritt der kausalen Interpretation, der von den abgeleiteten Regularitäten ausgehend die mit den Daten verträglichen Kausalmodelle liefert, wurde in der bestehenden Literatur bis anhin vernachlässigt, indem die Analyse typischerweise mit der Ermittlung der Minimalen Theorien endet. Es hat sich jedoch gezeigt, dass, etwa aufgrund von möglichen Switching-Literalen, der Schritt von den Minimalen Theorien zu Kausalmodellen nicht trivial ist. In csCNA+ sind erstmals entsprechende Algorithmen enthalten, die dem Rechnung tragen.

6.2. Ausblick

Die vorliegende Arbeit beschränkt sich auf die Kausalanalyse von crisp-set Konfigurationstabellen, die innerhalb eines idealen Entdeckungskontextes betrachtet werden. Dementsprechend ist auch das entwickelte Aufdeckungsverfahren csCNA+ vorerst nur vor einem derart idealen Hintergrund anwendbar. Ein nächster Schritt, der sich besonders aufdrängt, ist es, diese idealisierenden Annahmen abzubauen. Es stellt sich also die Frage, inwiefern csCNA+ sich für die Analyse von nicht maximal diversen und/oder inhomogenen crisp-set Konfigurationstabellen eignet und wie das Verfahren allenfalls angepasst werden müsste. Als weitere Aufweichung stellt sich die Frage, wie mit mehrwertigen oder fuzzy-set Konfigurationstabellen umzugehen ist. Während sich die einen dieser Fragen verhältnismässig einfach beantworten und durch geringfügige Adaptionen von csCNA+ meistern lassen, verlangen andere tiefgreifendere Anpassungen.

Hinsichtlich einer Aufweichung der Diversität von crisp-set Konfigurationstabellen verändern sich die Prozessschritte von csCNA+ kaum. Liegt eine homogene crisp-set Konfigurationstabelle vor, bei der nicht ausgeschlossen ist, dass empirisch mögliche Konfigurationen nicht darin enthalten sind, kann csCNA+ sequenziell zuerst auf diese ursprüngliche Konfigurationstabelle und anschliessend jeweils auf jede weitere Konfigurationstabelle angewendet werden, die aus der ursprünglichen Konfigurationstabelle sowie zusätzlichen Konfigurationen besteht. Der wesentliche Unterschied zwischen einer Analyse von maximal diversen Konfigurationstabellen und solchen, die nicht maximal divers sind, ist somit der, dass csCNA+ mehrmals angewendet wird und dadurch allenfalls auch die Anzahl verträglicher Kausalmodelle steigt.

Bei einer Aufweichung der Homogenität von crisp-set Konfigurationstabellen über eine Faktormenge $\mathbf{F}$ kann davon ausgegangen werden, dass dies bereits nach grösseren Weiterentwicklungen von csCNA+ verlangt. Der erste Prozessschritt der Wirkungsdeklaration (vgl. Algorithmus 5.1) ist vor diesem Hintergrund nicht mehr sinnvoll anwendbar, weil die Inhomogenität oftmals dazu führen könnte, dass es für jeden Faktor Konfigurationen gibt, bei denen ausschliesslich dieser variiert. Damit würden bei der Kausalanalyse einer entsprechenden Konfigurationstabelle mittels csCNA+ mit dem ersten Prozessschritt des Ausschlusses von Faktoren via (WD2) keine Faktoren resultieren, für deren Literale RDN-Bikonditionalen gesucht und folglich auch keine Kausalmodelle abgeleitet werden. Überwinden lässt sich dieses Problem mit dem Einbezug von Modellfit-Parametern, wie etwa das Konsistenzmass aus QCA oder Signifikanzniveaus. Konkret würde man Algorithmus 5.1 aus csCNA+ entfernen und ohne Ausnahme für Literale jedes Faktors nach RDN-Bikonditionalen suchen. Dabei würde man jedoch beispielsweise für die Bildung der RDN-Bikonditionale nur jene Konjunktionen $X_1 \cdots X_{j-1} X_{j+1} \cdots X_n$ als Minterme der inneren Instanzfunktionen f_{X_j} (bzw. $\overline{f_{X_j}}$) von X_j betrachten, von denen es in den Konfigurationen signifikant mehr Fälle mit konstant bleibenden Faktoren $X_1, \dots, X_{j-1}, X_{j+1}, \dots, X_n$ gibt, bei denen X_j den Wert 1 (bzw. 0) annimmt, gegenüber den Fällen, bei denen X_j den Wert 0 (bzw. 1) annimmt.

Bei einer Aufweichung der angenommenen Zweiwertigkeit auf mehrwertige oder fuzzy-set Konfigurationstabellen sind tiefgreifendere Anpassungen notwendig. Das kausaltheorietische Fundament bleibt grundsätzlich zwar dasselbe, mit dem kleinen Unterschied, dass nicht mehr von Zusammenhängen zwischen (positiven und negativen) Literalen die Rede ist, die die Werte 1 oder 0 annehmen, sondern allgemeiner von Zusammenhängen zwischen Faktoren, die einen bestimmten Wert annehmen. Das hier entwickelte mathematische Fundament der Konfigurationsalgebra jedoch, die ein Modell der zweielementigen Booleschen Algebra ist und daher spezifisch für das konfigurationale kausale Modellieren im crisp-set Fall entwickelt wurde, muss durch eine mehrwertige oder Fuzzy-Logik ersetzt werden. Zumal csCNA+ die Konfigurationsalgebra als Fundament nutzt, hat dies zur Folge, dass eine entsprechende Adaption von csCNA+ beispielsweise hin zu einer fuzzy-set Variante fsCNA+ anders aussieht. Zwar könnte eine fuzzy-set Konfigurationstabelle durch die Hinzunahme eines weiteren Vorbereitungsschrittes dichotomisiert werden, um die resultierende Tabelle wiederum csCNA+ als Input mitzugeben. Ein solches Vorgehen wird im Wesentlichen von QCA umgesetzt. Dies entspricht jedoch nicht einem idealen Lösungsansatz, zumal damit die grundlegende Kritik der Verzerrung der Daten bei crisp-set Analysen, die überhaupt erst dazu motiviert hat, fuzzy-set Analysen durchzuführen, nicht vollständig aus dem Weg geschaffen ist. Vielmehr müsste man zuerst erneut das mathematische Fundament von fsCNA+ ausarbeiten. Man müsste unter anderem die dazugehörigen Operationen definieren und darauf basierend ein neues Verfahren zur Minimierung von Fuzzyformeln in Minimalformen entwickeln, die RDN-Bedingungen darstellen. Als Grundlage solcher Ansät-

ze können unter anderem wiederum Methoden und Erkenntnisse dienen, die vor allem auf die Digitaltechnik und das damit eng verknüpfte Gebiet der Logik-Optimierung zurückgehen. Im Rahmen der Computertechnik befasst man sich bereits mehrere Jahrzehnte intensiv mit dieser Thematik. Unter anderem hat man in diesem Bereich bereits kurz nach der Entwicklung der Fuzzy-Mengen durch Lotfi Zadeh im Jahre 1965 damit begonnen, sich mit der Frage auseinanderzusetzen, wie Normalformen entsprechender Fuzzy-Funktionen minimiert werden können (vgl. u. a. Siy & Chen (1972), Kandel (1973), Bhat (1981)).

Allgemein lohnt es sich, künftig vermehrt einen Blick auf bestimmte Gebiete der Informationstechnologie zu werfen, von denen es zu einigen eine Analogie im Gebiet des konfiguralen kausalen Modellierens zu geben scheint. Gerade auch die jüngsten Entwicklungen in den Bereichen der künstlichen Intelligenz und der neuronalen Netze fördern Methoden und Instrumente zutage, die für das konfigurale kausale Modellieren von grossem Interesse sein könnten. Denkt man diesbezüglich etwa an Algorithmen der Bilderkennung, die beispielsweise trotz Lichteffekten oder unterschiedlichen Positionen Gesichter auf Bildern identifizieren, darf davon ausgegangen werden, dass mit ähnlichen Algorithmen die Erkennung von Kausalstrukturen aus einem verhältnismässig sehr rudimentären crisp-set oder fuzzy-set Muster ebenfalls möglich ist. Diese zugegebenermassen sehr plakative Querverbindung zeigt, dass sich der Blick über den Tellerrand lohnt. Letzteres ist seit den Anfängen charakteristisch für die Literatur zum konfiguralen kausalen Modellieren, als Ragin in den Sozialwissenschaften eingesetzte Forschungsmethoden mit Erkenntnissen aus der Philosophie sowie der Digitaltechnik verbunden hat. Man darf gespannt bleiben, ob entsprechende Adaptionen aus anderen Disziplinen auch künftig derart fruchtbar sind und wie diese aussehen könnten.

A. Anhang

A.1. Notation

$A, B, X_1, X_2, \dots$	Faktorlitterale bzw. Faktoren
$\mathbf{F}, \mathbf{M}, \mathbf{F}_1, \mathbf{F}_2, \dots$	Mengen von Faktoren, Faktorlitteralen oder Booleschen Ausdrücken
$\mathcal{X}, \mathcal{X}_i$	Unberücksichtigte Ko-Litterale
$\mathcal{Y}, \mathcal{Y}_i$	Latente Alternativursachen
$\mathbf{F}^w, \mathbf{G}^w, \dots$	Menge der Faktoren potenzieller Wirkungen aus der Faktormenge $\mathbf{F}$ bzw. $\mathbf{G}$ etc.
$\mathbf{F}^s, \mathbf{G}^s, \dots$	Menge der Switching-Litterale von Faktoren aus der Menge $\mathbf{F}$ bzw. $\mathbf{G}$ etc.
$\mathcal{K}_{\mathbf{F}}, \mathcal{K}_{x,y}$	Konfigurationstabelle über eine Faktormenge $\mathbf{F}$ bzw. Konfigurationstabelle aus Abbildung/Tabelle $x.y$
Ω	Datengenerierende Kausalstruktur
$\omega, \omega_1, \omega_2, \dots$	Kausalmodelle
$f, g, g_1, g_2, \dots$	Instanzfunktionen
$f_W, f_{X_i}, f_{X_j}, \dots$	Innere Instanzfunktionen von f mit abhängigem Faktor W bzw. X_i bzw. X_j etc.
$\epsilon, \epsilon_i, \sigma$	Konjunktionsterme
δ, ρ	Disjunktive Normalformen
$\phi_W, \phi_{X_i}, \phi_{X_j}, \dots$	RDN-Bedingungen von W bzw. X_i bzw. X_j etc.
$\beta_1, \beta_2, \beta_i, \dots$	Bikonditionale (RDN-Bikonditionale, einfache Minimale Theorien)
Φ, Ψ, Γ	Konjunktionen von (RDN-)Bikonditionalen ($\Phi = \beta_1 \cdot \beta_2 \cdot \dots$)
$\mathcal{B}$	Menge $\{0, 1\}$
$\mathcal{B}^n, \mathcal{D}$	n -faches kartesisches Produkt von $\mathcal{B}$ bzw. Teilmenge des n -fachen kartesischen Produktes von $\mathcal{B}$
$\mathcal{I}(f), \mathcal{NI}(f), \mathcal{ND}(f)$	Instantiierungs-, Non-Instantiierungs- und nicht definierter Bereich der Instanzfunktion f
$\mathbf{m}, \mathbf{p}_i, \mathbf{r}, \mathbf{s}$	Konjunktionsterme von mittels Booleschen Ausdrücken dargestellten Instanzfunktionen (Minterme, Primimplikanten, etc.)

A.2. Beweise

Konfigurationsalgebra und die Booleschen Gesetze

Es wird gezeigt, dass die Struktur $(0, 1, +, \cdot, \bar{})$, bestehend aus der Menge $\mathcal{B} = \{0, 1\}$ und (a) der konjunktiven Verknüpfung, (b) der disjunktiven Verknüpfung sowie (c) dem Komplement aus Tabelle A.1, eine Boolesche Algebra ist.

A	B	$A \cdot B$	A	B	$A + B$	A	$\bar{A}$
1	1	1	1	1	1	1	0
1	0	0	1	0	1	0	1
0	1	0	0	1	1		
0	0	0	0	0	0		
(a)			(b)			(c)	

Tabelle A.1.: Die drei Verknüpfungen (a) der Konjunktion, (b) der Disjunktion und (c) des Komplements.

Der Nachweis erfolgt durch die folgenden Wahrheitstabellen, die aufzeigen, dass für die obige Struktur die Booleschen Gesetze gelten.

Assoziativgesetze:

A	B	C	$(A \cdot B) \cdot C$	$A \cdot (B \cdot C)$	$(A + B) + C$	$A + (B + C)$
1	1	1	1	1	1	1
1	1	0	0	0	1	1
1	0	1	0	0	1	1
1	0	0	0	0	1	1
0	1	1	0	0	1	1
0	1	0	0	0	1	1
0	0	1	0	0	1	1
0	0	0	0	0	0	0

Konklusion: $\forall A, B, C \in \mathcal{B} : (A \cdot B) \cdot C = A \cdot (B \cdot C)$ bzw. $(A + B) + C = A + (B + C)$.

Kommutativgesetze:

A	B	$A \cdot B$	$B \cdot A$	$A + B$	$B + A$
1	1	1	1	1	1
1	0	0	0	1	1
0	1	0	0	1	1
0	0	0	0	0	0

Konklusion: $\forall A, B \in \mathcal{B} : A \cdot B = B \cdot A$ bzw. $A + B = B + A$.

Absorptionsgesetze:

A	B	$A \cdot (A + B)$	$A + (A \cdot B)$
1	1	1	1
1	0	1	1
0	1	0	0
0	0	0	0

Konklusion: $\forall A, B \in \mathcal{B} : A = A \cdot (A + B)$ bzw. $A = A + (A \cdot B)$.

Distributivgesetze:

A	B	C	$(A + B) \cdot C$	$(A \cdot C) + (B \cdot C)$	$(A \cdot B) + C$	$(A + C) \cdot (B + C)$
1	1	1	1	1	1	1
1	1	0	0	0	1	1
1	0	1	1	1	1	1
1	0	0	0	0	0	0
0	1	1	1	1	1	1
0	1	0	0	0	0	0
0	0	1	1	1	1	1
0	0	0	0	0	0	0

Konklusion: $\forall A, B, C \in \mathcal{B} : (A + B) \cdot C = (A \cdot C) + (B \cdot C)$ bzw. $(A \cdot B) + C = (A + C) \cdot (B + C)$.

Neutrale Elemente:

A	$A \cdot 1$	$A \cdot 0$	$A + 0$	$A + 1$
1	1	0	1	1
0	0	0	0	1

Konklusion: $\forall A \in \mathcal{B} : A \cdot 1 = A, A \cdot 0 = 0, A + 0 = A$ und $A + 1 = 1$.

Komplement:

A	$A \cdot \bar{A}$	$A + \bar{A}$
1	0	1
0	0	1

Konklusion: $\forall A \in \mathcal{B} : A \cdot \bar{A} = 0$ und $A + \bar{A} = 1$.

Herleitung der weiteren RDN-Formen von g_E (vgl. S. 186)

$$\begin{aligned}
 g_E(\cdot) &= AB + A\bar{B}\bar{D} + A\bar{B}\bar{C} + \bar{A}\bar{B} \\
 &= AB + A\bar{B}\bar{D} + AB + A\bar{B}\bar{C} + \bar{A}\bar{B} \\
 &= A(B + \bar{B}\bar{D}) + A(B + \bar{B}\bar{C}) + \bar{A}\bar{B} \\
 &= A(BD + \bar{B}\bar{D} + B\bar{D} + \bar{B}\bar{D}) + A(BC + \bar{B}\bar{C} + B\bar{C} + \bar{B}\bar{C}) + \bar{A}\bar{B} \\
 &= A(B + \bar{D}) + A(B + \bar{C}) + \bar{A}\bar{B} \\
 &= AB + A\bar{D} + AB + A\bar{C} + \bar{A}\bar{B} \\
 &= AB + A\bar{D} + A\bar{C} + \bar{A}\bar{B}
 \end{aligned}$$

$$\begin{aligned}
 g_E(\cdot) &= AB + A\bar{B}\bar{D} + A\bar{B}\bar{C} + \bar{A}\bar{B} \\
 &= AB + A\bar{B}\bar{D} + \bar{A}\bar{B} + A\bar{B}\bar{C} + \bar{A}\bar{B} \\
 &= AB + \bar{B}(A\bar{D} + \bar{A}) + \bar{B}(A\bar{C} + \bar{A}) \\
 &= AB + \bar{B}(A\bar{D} + \bar{A}D + \bar{A}\bar{D} + \bar{A}\bar{D}) + \bar{B}(A\bar{C} + \bar{A}C + \bar{A}\bar{C} + \bar{A}\bar{C}) \\
 &= AB + \bar{B}(\bar{D} + \bar{A}) + \bar{B}(\bar{C} + \bar{A}) \\
 &= AB + \bar{B}\bar{D} + \bar{A}\bar{B} + \bar{B}\bar{C} + \bar{A}\bar{B} \\
 &= AB + \bar{B}\bar{D} + \bar{B}\bar{C} + \bar{A}\bar{B}
 \end{aligned}$$

Nachweis von (RES) und (RED) (vgl. S. 189 und 205)

Resolutionsregel (RES):

$$\begin{aligned} & (X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W) \cdot (X_1 \cdot \dots \cdot X_{i-1} \cdot \bar{X}_i \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W) \\ & = X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W. \end{aligned}$$

Nachweis von (RES):

$$\begin{aligned} & (X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W) \cdot (X_1 \cdot \dots \cdot X_{i-1} \cdot \bar{X}_i \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W) \\ & = (\overline{X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n} + W) \cdot (\overline{X_1 \cdot \dots \cdot X_{i-1} \cdot \bar{X}_i \cdot X_{i+1} \cdot \dots \cdot X_n} + W) \\ & \stackrel{(*)}{=} (\bar{X}_1 + \dots + \bar{X}_{i-1} + \bar{X}_i + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \cdot (\bar{X}_1 + \dots + \bar{X}_{i-1} + X_i + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \\ & = \bar{X}_1 \cdot (\bar{X}_1 + \dots + \bar{X}_{i-1} + X_i + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \\ & \quad + \dots \\ & \quad + \bar{X}_{i-1} \cdot (\bar{X}_1 + \dots + \bar{X}_{i-1} + X_i + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \\ & \quad + \bar{X}_i \cdot (\bar{X}_1 + \dots + \bar{X}_{i-1} + X_i + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \\ & \quad + \bar{X}_{i+1} \cdot (\bar{X}_1 + \dots + \bar{X}_{i-1} + X_i + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \\ & \quad + \dots \\ & \quad + \bar{X}_n \cdot (\bar{X}_1 + \dots + \bar{X}_{i-1} + X_i + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \\ & \quad + W \cdot (\bar{X}_1 + \dots + \bar{X}_{i-1} + X_i + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \\ & \stackrel{(**)}{=} \bar{X}_1 + \dots + \bar{X}_{i-1} + \bar{X}_{i+1} + \dots + \bar{X}_n + W \\ & \stackrel{(*)}{=} \overline{X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n} + W \\ & = X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W \end{aligned}$$

Redundanzregel (RED):

$$\begin{aligned} & (X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W) \cdot (\overline{X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \cdot \bar{W}}) \\ & = X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W \end{aligned}$$

Nachweis von (RED):

$$\begin{aligned} & (X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W) \cdot (\overline{X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \cdot \bar{W}}) \\ & = (\overline{X_1 \cdot \dots \cdot X_{i-1} \cdot X_i \cdot X_{i+1} \cdot \dots \cdot X_n} + W) \cdot (\bar{X}_1 + \dots + \bar{X}_{i-1} + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \\ & = (\bar{X}_1 + \dots + \bar{X}_{i-1} + \bar{X}_i + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \cdot (\bar{X}_1 + \dots + \bar{X}_{i-1} + \bar{X}_{i+1} + \dots + \bar{X}_n + W) \end{aligned}$$

$$\begin{aligned}
 &= \overline{X_1} \cdot (\overline{X_1} + \dots + \overline{X_{i-1}} + \overline{X_{i+1}} + \dots + \overline{X_n} + W) \\
 &\quad + \dots \\
 &\quad + \overline{X_{i-1}} \cdot (\overline{X_1} + \dots + \overline{X_{i-1}} + \overline{X_{i+1}} + \dots + \overline{X_n} + W) \\
 &\quad + \overline{X_i} \cdot (\overline{X_1} + \dots + \overline{X_{i-1}} + \overline{X_{i+1}} + \dots + \overline{X_n} + W) \\
 &\quad + \overline{X_{i+1}} \cdot (\overline{X_1} + \dots + \overline{X_{i-1}} + \overline{X_{i+1}} + \dots + \overline{X_n} + W) \\
 &\quad + \dots \\
 &\quad + \overline{X_n} \cdot (\overline{X_1} + \dots + \overline{X_{i-1}} + \overline{X_{i+1}} + \dots + \overline{X_n} + W) \\
 &\quad + W \cdot (\overline{X_1} + \dots + \overline{X_{i-1}} + \overline{X_{i+1}} + \dots + \overline{X_n} + W) \\
 &\stackrel{(**)}{=} \overline{X_1} + \dots + \overline{X_{i-1}} + \overline{X_{i+1}} + \dots + \overline{X_n} + W \\
 &\stackrel{(*)}{=} \overline{X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n} + W \\
 &= X_1 \cdot \dots \cdot X_{i-1} \cdot X_{i+1} \cdot \dots \cdot X_n \rightarrow W
 \end{aligned}$$

Bemerkung: Bei Schritt (*) kommen die De Morganschen Gesetze zum Einsatz, wonach für zwei Boolesche Terme S und T gilt: $\overline{S \cdot T} = \overline{S} + \overline{T}$ beziehungsweise $\overline{S + T} = \overline{S} \cdot \overline{T}$. Die Vereinfachung des gesamten Booleschen Ausdrucks mittels Schritt (**) erfolgt im Wesentlichen durch mehrmaliges Verwenden des Absorptionsgesetzes $S + S \cdot T = S$. Dabei verschwindet X_i (bzw. $\overline{X_i}$) deshalb aus dem gesamten Ausdruck, weil X_i Teil einer Konjunktion $X_i \cdot \overline{X_i} (= 0)$ ist oder Teil der konjunktiven Verknüpfung $X_i \cdot \overline{X_j}$ (bzw. $\overline{X_i} \cdot \overline{X_j}$). Eine solche Konjunktion $X_i \cdot \overline{X_j}$ kann jeweils via Absorption aufgrund des vorhandenen Terms $\overline{X_j} \cdot \overline{X_j} (= \overline{X_j})$ als Ganzes wegreduziert werden.

A.3. R-Replikationsskript

```

# Notwendiges R-Package:
library(cna)

#-----
# Konfigurationstabelle K2.9 (vgl. S. 40)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(A*B+C <-> D)*(c+a*E <-> F)")
tab <- tab[order(tab$A,tab$B,tab$C,decreasing=T),]
tab
# Berechnung und Ausgabe aller RDN-Bikonditionale:
RDN <- cna(tab)
asf(RDN)

```

```
#-----
# Konfigurationstabelle K2.11 (vgl. S. 45)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(a+b+c <-> D)*(A+B <-> E)*(A+C <-> E)")
tab <- tab[order(tab$A,tab$B,tab$C,tab$D,tab$E,decreasing=T),]
tab
# Berechnung aller RDN-Bikonditionale:
mt <- cna(tab)
asf(mt)
# Bildung der Konjunktion Phi aller RDN-Bikonditionale:
Phi <- "(d + B*c + b*e <-> A)*(d + b*C + c*e <-> A)*(d + b*e + c*e <-> A)*
(a + b + c <-> D)*(A + B <-> E)*(A + C <-> E)"
# Ausgabe aller strukturell redundanzfreien Konjunktionen Psi aus Phi:
minimalizeCsf(Phi)

#-----
# Konfigurationstabelle K2.6 (vgl. S. 31 bzw. S. 47)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(A+E <-> B)*(A+D <-> C)",full.tt("A*B*C*D*E"))
tab <- tab[order(tab$A,tab$B,tab$C,tab$D,tab$E,decreasing=T),]
tab
# Berechnung und Ausgabe der RDN-Bikonditionale:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
asf(mt)
# Ausgabe aller Konjunktionen von RDN-Bikonditionalen von Faktorliterals
# paarweise verschiedener Faktoren:
csf(mt)
# Berechnung und Ausgabe der Minimalen Theorien:
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K2.12 (vgl. S. 49)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(d+B*c+b*e <-> A)*(a+b+c <-> D)*(A+C*f <-> E)",
                    full.tt("A*B*C*D*E*f"))
tab <- tab[order(tab$A,tab$B,tab$C,tab$D,tab$E,tab$f,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)
```

```
#-----
# Konfigurationstabelle K2.14 (vgl. S. 52)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(A*B+C <-> D)*(f+a*D <-> C)")
tab <- tab[order(tab$A,tab$B,tab$C,tab$D,tab$F,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K2.19 (vgl. S. 70)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(A+B <-> C)*(C+D <-> E)")
tab <- tab[order(tab$A,tab$B,tab$C,tab$D,tab$E,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K2.21 (vgl. S. 73)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(B*S+U <-> O)*(B*s+A <-> K)*(K+O <-> E)",
                    full.tt("B*S*A*U*K*O*E"))
tab <- tab[order(tab$B,tab$S,tab$A,tab$U,tab$K,tab$O,tab$E,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K2.26 (vgl. S. 88)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(A+B+F <-> C)*(A+B+D <-> E)",full.tt("A*B*C*D*E*F"))
tab <- tab[order(tab$A,tab$B,tab$C,tab$D,tab$E,tab$F,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
```

```

minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K3.4 (vgl. S. 102)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(A+B <-> C)*(C+D <-> E)")
tab <- tab[order(tab$A,tab$B,tab$C,tab$D,tab$E,decreasing=T),]
tab
# Berechnung und Ausgabe der RDN-Bikonditionale:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
asf(mt)
# Ausgabe aller Konjunktionen von RDN-Bikonditionalen von Faktorliterals
# paarweise verschiedener Faktoren:
csf(mt)
# Berechnung und Ausgabe der Minimalen Theorien:
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K3.6 (vgl. S. 113)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(D+E+F+G <-> H)",full.tt("D+E*F*G*H"))
tab <- tab[order(tab$D,tab$E,tab$F,tab$G,tab$H,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K3.9 (vgl. S. 115)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(D+E+F+G <-> H)",full.tt("D+E*F*G*H"))
tab <- tab[order(tab$D,tab$E,tab$F,tab$G,tab$H,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K3.15a (vgl. S. 122)
#-----
# Erzeugung der Konfigurationstabelle:

```

```
tab <- selectCases("(F+G <-> H)",full.tt("D*E*F*G*H"))
tab <- tab[order(tab$D,tab$E,tab$F,tab$G,tab$H,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K3.15b (vgl. S. 122)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(E+F+G <-> H)",full.tt("D*E*F*G*H"))
tab <- tab[order(tab$D,tab$E,tab$F,tab$G,tab$H,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K3.15c (vgl. S. 122)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(D+F+G <-> H)",full.tt("D*E*F*G*H"))
tab <- tab[order(tab$D,tab$E,tab$F,tab$G,tab$H,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K3.15d (vgl. S. 122)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(D+E+F+G <-> H)",full.tt("D*E*F*G*H"))
tab <- tab[order(tab$D,tab$E,tab$F,tab$G,tab$H,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K3.16a (vgl. S. 123)
#-----
```

```
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(I+J <-> L)*(J+K <-> M)",full.tt("I*J*K*L*M"))
tab <- tab[order(tab$I,tab$J,tab$K,tab$L,tab$M,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K3.16b (vgl. S. 123)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(I+J <-> L)*M",full.tt("I*J*K*L*M"))
tab <- tab[order(tab$I,tab$J,tab$K,tab$L,tab$M,decreasing=T),]
tab
# Berechnung und Ausgabe der Minimalen Theorien:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K5.1 (vgl. S. 177)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(A*B+A*d+b*c+a*b <-> E)",full.tt("A*B*C*D*E"))
tab <- tab[order(tab$A,tab$B,tab$C,tab$D,tab$E,decreasing=T),]
tab
# Berechnung und Ausgabe der RDN-Bikonditionale:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
asf(mt)
# Ausgabe aller Konjunktionen von RDN-Bikonditionalen von Faktorliterals
# paarweise verschiedener Faktoren:
csf(mt)
# Berechnung und Ausgabe der Minimalen Theorien:
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K5.2 (vgl. S. 177)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(R+S <-> V)*(S+T <-> U)",full.tt("R*S*T*U*V"))
tab <- tab[order(tab$R,tab$S,tab$T,tab$U,tab$V,decreasing=T),]
tab
# Berechnung und Ausgabe der RDN-Bikonditionale:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
```

```
asf(mt)
# Ausgabe aller Konjunktionen von RDN-Bikonditionalen von Faktorliteralen
# paarweise verschiedener Faktoren:
csf(mt)
# Berechnung und Ausgabe der Minimalen Theorien:
csfAll <- csf(mt, "all")
minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)

#-----
# Konfigurationstabelle K5.24 (vgl. S. 220)
#-----
# Erzeugung der Konfigurationstabelle:
tab <- selectCases("(E+F<->H)*(A+B*c<->E)*(D+B*C<->F)*(c+d<->G)",
                    full.tt("A*B*C*D*E*F*G*H"))
tab <- tab[order(tab$A,tab$B,tab$C,tab$D,tab$E,tab$F,tab$G,tab$H,
                  decreasing=T),]

tab
# Berechnung und Ausgabe der RDN-Bikonditionale:
mt <- cna(tab,maxstep=c(4, 5, 15), details = TRUE)
asf(mt)
# Ausgabe aller Konjunktionen von RDN-Bikonditionalen von Faktorliteralen
# paarweise verschiedener Faktoren:
csf(mt)
# Berechnung und Ausgabe der Minimalen Theorien:
csfAll <- csf(mt, "all")
Psi <- minimalizeCsf(subset(csfAll, exhaustiveness == 1)$condition)
print(Psi, subset = 1:15)
```

Sachregister

- Abdeckung, 152
 - irredundante, 153, 194
- Abdeckungsmass, 135
- Abdeckungstabelle, 213
- allumfassende Kausalstruktur, 63
- Ambiguität, 48
- analytisches Moment, 100–104
- Anti-Nezessitarismus, 14
- Bedingung
 - hinreichende, 25
 - INUS-, 28
 - minimal hinreichende, 29, 33
 - minimal notwendige, 34
 - notwendige, 25
 - RDN-, 34
- Belegung, 35
- Bereich
 - Definitions-, 154
 - Instantiierungs-, 154
 - ND-, 154
 - Noninstantiierungs-, 154
- Bikonditional, 25
 - RDN-, 35
- Boolesche Algebra, 140
- Boolesche Funktion, 145
 - partiell definierte, 145
- Boolescher Ausdruck, 145–146
- brute fact, 14
- ceteris-paribus, 29
- Coincidence Analysis (CNA), 6, 9–10
- Common-Cause-Struktur, 18, 69–70
- crisp-set CNA (csCNA), 204–215
- crisp-set QCA (csQCA), 100, 187–203
- De Morgansches Gesetz, 163
- Difference-Maker, 28, 47
- Difference-Making-Paar, 36–38
- Differenzmethode, 59
- Digraph, 19
- disjunktive Minimalform (DMF), 151
- disjunktive Normalform (DNF), 34, 149
 - kanonische (KDNF), 149
 - redundanzfreie (RDNF), 168–171
- Dominanzprüfung, 196–200
- Dualität, 141
- Erweiterung (Instanzfunktion), 155–156
- Erweiterungsklausel, 51–53
- Faktor, 15–16
- Faktorliteral, 16, *siehe* Literal
- Faktormenge, 21
- Faktorwert, 16
- Feedbackstruktur, 19, 91
- Forschungsdesign, 96
- Fuzzy-Menge, 100
- fuzzy-set QCA (fsQCA), 100
- Güte, 101, 109

- Hauptsatz der Booleschen Algebra, 150
homogen, *siehe* Homogenität
Homogenität, 128
 strikte, 120
 triviale, 124
Hume'scher Aktualismus, 14
Hypergraph, 19, 21
Hyperkante, 20

Implikant, 151
Instanz (eines Faktors), 15
Instanzfunktion, 144
 innere, 165–168
 partielle, 153, 155–156
 totale, 153, 154
Intransitivität, 72

Kalibrierung, 97
kausal relevant, *siehe* Kausale Relevanz
kausale Komplexität, 2, 17–19
kausale Relevanz, 54
 direkte, 90
 indirekte, 90
kausaler Fehlschluss, 110–111
kausaler Zyklus, 18
kausales Feld, 29
Kausalgraph, 21
Kausalkette, 18, 79
Kausalmodell, 102–103
 verträgliches, 102
Kausalrelation
 auf Token-Ebene, 14–15
 auf Typen-Ebene, 14–15
Kernprimimplikant, 152
Kettenproblem, 71, 80–86
Ko-Faktorliteral, *siehe* Ko-Literal
Koinzidenz, *siehe* Konfiguration
Konfiguration, 5, 24, 98–99
Konfigurationsalgebra, 143
Konfigurationstabelle, 26
 crisp-set, 99
 komplette, 59
 maximal diverse, 137
Konjunktionsterm, 33, 149
 0-stelliger, 105
Konsistenzmass, 133–135

Literal, 16, 146
 Ko-, 23
 negatives, 16, 146
 positives, 16, 146
 Switching-, 75

Minimale Theorie, 46
 einfache, 36, 47, 66
 komplexe, 66–68
Minterm, 149
modal unabhängig, *siehe* modale Unabhängigkeit
modale Unabhängigkeit, 22
MT-Form (MTF), 164, 171–172
Multikollinearität, 4

Pfad, 21
potenzielle Störursache, 119
potenzielle Wirkung, 118
Primimplikant, 151
Primimplikantentabelle, 194
Prinzip
 der kausalen Eindeutigkeit, 60
 der Redundanzfreiheit, 28
 Determinismus-, 58
 Difference-Making-, 62
 Kausalitäts-, 59
 Permanenz-, 62

Qualitative Comparative Analysis (QCA), 6–9

- redundant, 28
- Redundanzregel, 205
- Regressionsanalyse, 2–5
- Relation
 - antisymmetrische, 65
 - asymmetrische, 65
 - intransitive, 72
 - symmetrische, 65
 - transitive, 72
- Repräsentant (einer Instanzfunktion), 148
- Resolutionsregel, 188–189
- Rohdaten, 95–97
- strukturelle
 - Minimalität, 43
 - Redundanz, 44
- Ursache
 - alternative, 17, 56
 - gemeinsame, 18, 69
 - komplexe, 17, 56
 - latente, 23
 - Schein-, 61
- Variable, 143
- Variation, 117
 - ausschlaggebende, 125–128
- Verfahren
 - nach Baumgartner, 204–215, 217
 - nach Petrick, 200–203, 217
 - nach Quine und McCluskey (QMC), 188–200, 217
- Verhaltensmuster (von Faktoren), 158
- Verhaltensregel (von Faktoren), 159
- Verknüpfungsbasis, 147
- Vollständigkeitsklausel, 50
- Wahrheitstabelle, *siehe* Wertetabelle
- Wahrheitswertetabelle, *siehe* Wertetabelle
- Wechselwirkung, 18, 93
- Wertetabelle, 35
- Wohlgeformtheit, 62

Literaturverzeichnis

- Adams, R. M. (1974). Theories of Actuality. *Noûs*, 8(3), 211–231.
- Antonakis, J., Bendahan, S., Jacquart, P., & Lalive, R. (2010). On making causal claims: A review and recommendations. *The Leadership Quarterly*, 21, 1086 – 1120.
- Armstrong, D. M. (1983). *What is a Law of Nature?* Cambridge: Cambridge University Press.
- Ausiello, G., Franciosa, P. G., & Frigioni, D. (2001). Directed Hypergraphs: Problems, Algorithmic Results, and a Novel Decremental Approach. In A. Restivo, S. R. D. Rocca, & L. Roversi (Eds.), *Theoretical Computer Science: 7th Italian Conference, ICTCS 2001 Torino, Italy, October 4-6, 2001 Proceedings* (pp. 312–328). Berlin, Heidelberg: Springer-Verlag.
- Basurto, X. & Speer, J. (2012). Structuring the Calibration of Qualitative Data as Sets for Qualitative Comparative Analysis (QCA). *Field Methods*, 24(2), 155–174.
- Baumgartner, M. (2006). *Complex Causal Structures. Extensions of a Regularity Theory of Causation*. PhD thesis, University of Bern.
- Baumgartner, M. (2008a). The Causal Chain Problem. *Erkenntnis*, 69, 201–226.
- Baumgartner, M. (2008b). Regularity Theories Reassessed. *Philosophia*, 36, 327–354.
- Baumgartner, M. (2009). Uncovering deterministic causal structures: a Boolean approach. *Synthese*, 170, 71–96.
- Baumgartner, M. (2013a). Detecting Causal Chains in Small-*n* Data. *Field Methods*, 25(1), 3–24.
- Baumgartner, M. (2013b). A Regularity Theoretic Approach to Actual Causation. *Erkenntnis*, 78, 85–109.

- Baumgartner, M. (2015). Parsimony and Causality. *Quality & Quantity*, 49(2), 839–856.
- Baumgartner, M. & Ambühl, M. (2019). *cna: An R Package for Configurational Causal Inference and Modeling*. Norway and Switzerland: University of Bergen, ConsultAG. (Software Version 2.2.0).
- Baumgartner, M. & Eppe, R. (2014). A Coincidence Analysis of a Causal Chain. The Swiss Minaret Vote. *Sociological Methods & Research*, 43(2), 280–312.
- Baumgartner, M. & Graßhoff, G. (2004). *Kausalität und kausales Schliessen. Eine Einführung mit interaktiven Übungen*. Bern: Bern Studies in the History and Philosophy of Science. Educational materials.
- Becker, B. & Molitor, P. (2008). *Technische Informatik. Eine einführende Darstellung*. Oldenbourg Verlag.
- Berg-Schlosser, D. & Meur, G. D. (2009). *Configurational Comparative Methods: Qualitative Comparative Analysis (QCA) and Related Techniques*, chapter Comparative Research Design: Case and Variable Selection, (pp. 19–32). SAGE Publications.
- Bhat, K. V. (1981). Minimization of disjunctive normal forms of fuzzy logic functions. *Journal of the Franklin Institute*, 311(3), 171 – 185.
- Böhme, G. (1981). *Algebra*. Springer-Verlag.
- Büning, H. K. & Lettmann, T. (1994). *Aussagenlogik. Deduktion und Algorithmen*. Stuttgart: B.G. Teubner.
- Boole, G. (1847). *Mathematical Analysis of Logic*. Cambridge: Macmillan, Barclay, & Macmillan.
- Boros, E., Hammer, P. L., Ibaraki, T., Kogan, A., Mayoraz, E., & Muchnik, I. (2000). An Implementation of Logical Analysis of Data. *IEEE Transactions on Knowledge and Data Engineering*, 12(2), 292–306.
- Bortz, J. & Döring, N. (2009). *Forschungsmethoden und Evaluation* (4 ed.). Heidelberg: Springer Medizin Verlag.

- Bortz, J. & Schuster, C. (2010). *Statistik für Human- und Sozialwissenschaftler* (7 ed.). Berlin, Heidelberg: Springer-Verlag.
- Brambor, T., Clark, W. R., & Golder, M. (2006). Understanding Interaction Models: Improving Empirical Analyses. *Political Analysis*, 14, 63–82.
- Braumoeller, B. F. (2004). Hypothesis Testing and Multiplicative Interaction Terms. *International Organization*, 58(4), 807–820.
- Broad, C. D. (1925). *The Mind and its Place in Nature*. New York: Harcourt, Brace & Company, Inc.
- Bronstein, I. N., Semendjaev, K. A., Musiol, G., & Mühlig, H. (2001). *Taschenbuch der Mathematik* (5 ed.). Thun and Frankfurt am Main: Verlag Harri Deutsch.
- Caren, N. & Panofsky, A. (2005). TQCA: A Technique for Adding Temporality to Qualitative Comparative Analysis. *Sociological Methods & Research*, 34(2), 147–172.
- Clark, W. R., Gilligan, M. J., & Golder, M. (2006). A Simple Multivariate Test for Asymmetric Hypotheses. *Political Analysis*, 14, 311–331.
- Crama, Y. & Hammer, P. L. (2011). *Boolean Functions: Theory, Algorithms, and Applications*. Cambridge University Press.
- Crama, Y., Hammer, P. L., & Iwakura, T. (1988). Cause-effect relationships and partially defined Boolean functions. *Annals of Operations Research*, 16(1), 299–325.
- Denis-Papin, M., Faure, R., Kaufmann, A., & Malgrange, Y. (1974). *Theorie und Praxis der Booleschen Algebra*. Braunschweig: Vieweg.
- Dierker, P. F. & Voxman, W. L. (1986). *Discrete Mathematics*. Harcourt Brace Jovanovich.
- Diestel, R. (2010). *Graph Theory* (4 ed.). Heidelberg: Springer-Verlag.
- Dowe, P. (1995). What's right and what's wrong with transference theories. *Erkenntnis*, 42(3), 363 – 374.
- Eberhardt, F. (2014). Direct Causes and the Trouble with Soft Interventions. *Erkenntnis*, 79, 755–777.

- Eschermann, B. (1993). *Funktionaler Entwurf digitaler Schaltungen*. Springer-Verlag.
- Fabricius, E. D. (1992). *Modern Digital Design and Switching Theory*. CRC Press LCC.
- Fodor, J. (1997). Special Sciences: Still Autonomous After All These Years. *Noûs*, 31(11), 149–163.
- Friedrich, R. J. (1982). In Defense of Multiplicative Terms in Multiple Regression Equations. *American Journal of Political Science*, 26(4), 797–833.
- Garey, M. R. & Johnson, D. S. (1979). *Computers and Intractability. A Guide to the Theory of NP-Completeness*. New York: W. H. Freeman and Company.
- Garnier, R. & Taylor, J. (2002). *Discrete Mathematics for New Technology* (2 ed.). London: Institute of Physics Publishing Ltd.
- Givone, D. D. (1970). *Introduction to switching circuit theory*. New York: McGraw-Hill.
- Goodman, N. (1979). *Fact, Fiction, and Forecast*. Cambridge: Harvard University Press.
- Gotthardt, K. (2005). *Grundlagen der Informationstechnik* (2 ed.). Münster: Lit Verlag.
- Graßhoff, G. & May, M. (2001). Causal Regularities. In W. Spoon, M. Ledwig, & M. Esfeld (Eds.), *Current Issues in Causation* (pp. 85–114). Mentis Verlag.
- Gross, J. L. & Yellen, J. (Eds.). (2004). *Handbook of Graph Theory*. CRC Press.
- Gumm, H.-P. & Poguntke, W. (1981). *Boolesche Algebra*. Mannheim: B.I.-Wissenschaftsverlag.
- Hausman, D. M. (2005). Causal Relata: Tokens, Types, or Variables? *Erkenntnis*, 63, 33–54.
- Hitchcock, C. (2001). The Intransitivity of Causation Revealed in Equations and Graphs. *The Journal of Philosophy*, 98(6), 273–299.
- Hitchcock, C. (2007). Prevention, Preemption, and the Principle of Sufficient Reason. *Philosophical Review*, 116(4).
- Hoffmann, D. W. (2016). *Grundlagen der Technischen Informatik* (5 ed.). München: Carl Hanser Verlag.

- Hofmann, U. & Baumgartner, M. (2011). Determinism and the Method of Difference. *THEORIA*, 71, 155–176.
- Hume, D. (1739). *A Treatise of Human Nature*. Oxford: Clarendon Press 1896.
- Hume, D. (1748). *An Enquiry concerning Human Understanding*. Chicago and La Salle, Illinois: Open Court 1988.
- Ihringer, T. (1994). *Diskrete Mathematik* (2 ed.). Vieweg+Teubner Verlag.
- Kalisch, M., Mächler, M., Colombo, D., Maathuis, M. H., & Bühlmann, P. (2012). Journal of Statistical Software. *Journal of Statistical Software*, 47(11), 1–26.
- Kandel, A. (1973). On Minimization of Fuzzy Functions. *IEEE Transactions on Computers*, C-22(9), 826 – 832.
- Karnaugh, M. (1953). The Map Method for Synthesis of Combinational Logic Circuits. *Transactions of the AIEE*, 72(9), 593–599.
- Karp, R. M. (1972). Reducibility Among Combinatorial Problems. In I. R. E. Miller & J. W. Thatcher (Eds.), *Complexity of Computer Computations* (pp. 85–103). New York: Plenum.
- Keller, J. & Paul, W. J. (1997). *Hardware Design. Formaler Entwurf digitaler Schaltungen*. Leipzig: B. G. Teubner Verlagsgesellschaft.
- Kenny, D. A. (1979). *Correlation and causality*. New York: Wiley-Interscience.
- Krook, M. L. (2010). Women’s Representation in Parliament: A Qualitative Comparative Analysis. *Political Studies*, 58, 886–908.
- Kumar, A. A. (2014). *Fundamentals of Digital Circuits*. PHI Learning Pvt. Ltd.
- Kunz, W. & Stoffel, D. (1997). *Reasoning in Boolean Networks. Logic Synthesis and Verification Using Testing Techniques*. Dordrecht: Springer-Science & Business Media, B. V.
- Kvart, I. (2001). The Counterfactual Analysis of Cause. *Synthese*, 127, 389–427.
- Lemmon, E. J. (1998). *Beginning Logic* (2 ed.). Boca Raton: Chapman and Hall / CRC.
- Lewis, D. (1973). Causation. *The Journal of Philosophy*, 70, 556–567.

- Lewis, D. (1983). New work for a theory of universals. *Australasian Journal of Philosophy*, 61(4), 343–377.
- Lewis, D. K. (1986). *Counterfactuals*. Oxford: Baseil Blackwell.
- Mackie, J. L. (1980). *The Cement of the Universe. A Study of Causation* (2 ed.). Oxford: Clarendon Press.
- Mahoney, J. & Goertz, G. (2006). A Tale of Two Cultures: Contrasting Quantitative and Qualitative Research. *Political Analysis*, 14, 227–249.
- McCluskey, E. J. (1956). Minimization of Boolean Functions. *The Bell system technical journal*, 35(6), 1417–1444.
- Mendelson, E. (1970). *Schaum's Outline of Boolean Algebra and Switching Circuits*. New York: McGraw-Hill.
- Menzel, C. (1990). Actualism, Ontological Commitment, and Possible Worlds Semantics. *Synthese*, 85(3), 3.
- Mill, J. S. (1843). *A System of Logic*, volume 1. London: John W. Parker, West Strand.
- Morgenstern, B. (1997). *Elektronik: Digitale Schaltungen und Systeme*, volume 2. Braunschweig/Wiesbaden: Vieweg.
- Nelson, V. P., Nagle, H. T., Carroll, B. D., & Irwin, J. D. (1995). *Digital Circuit Analysis and Design*. Upper Saddle River, New Jersey: Prentice-Hall.
- Pereboom, D. (Ed.). (2009). *Free Will* (2 ed.). Indianapolis: Hackett Publishing Company, Inc.
- Psillos, S. (2009). Causal pluralism. In R. Vanderbeeken & B. D'Hooghe (Eds.), *Worldviews, Science and Us: Studies of Analytical Metaphysics* (pp. 131–151). Singapore: World Scientific Publishing Co. Pte. Ltd.
- Quine, W. V. O. (1952). The Problem of Simplifying Truth Functions. *The American Mathematical Monthly*, 59(8), 521–531.
- Quine, W. V. O. (1959). On Cores and Prime Implicants of Truth Functions. *The American Mathematical Monthly*, 66(9), 755–760.

- R Core Team (2016). *R: A Language and Environment for Statistical Computing*. Vienna, Austria: R Foundation for Statistical Computing. <https://www.R-project.org/>.
- Ragin, C. (1987). *The Comparative Method*. University of California Press.
- Ragin, C. (2008). *Redesigning Social Inquiry. Fuzzy Sets and Beyond*. Chicago and London: University of Chicago Press.
- Ragin, C. & Davey, S. (2009). *fs/QCA: Fuzzy-Set/Qualitative Comparative Analysis*. Tucson: Department of Sociology, University of Arizona. (Software Version 2.5).
- Ragin, C. C. (2000). *Fuzzy-Set Social Science*. Chicago and London: University Chicago Press.
- Richter, R., Sander, P., & Stucky, W. (1997). *Der Rechner als System: Organisation, Daten, Programme*. Stuttgart: B. G. Teubner.
- Rihoux, B. & Ragin, C. (2009). *Configurational Comparative Methods: Qualitative Comparative Analysis (QCA) and Related Techniques*, chapter Introduction, (pp. xvii–xxv). SAGE Publications.
- Roth, C. H. & Kinney, L. L. (2010). *Fundamentals of Logic Design* (6 ed.). Stamford: Cengage Learning.
- Scarbata, G. (2001). *Synthese und Analyse digitaler Schaltungen*. München: Oldenbourg Verlag.
- Schneider, C. Q. & Wagemann, C. (2007). *Qualitative Comparative Analysis (QCA) und Fuzzy Sets. Ein Lehrbuch für Anwender und jene, die es werden wollen*. Opladen and Farmington Hills: Verlag Barbara Burdich.
- Schneider, C. Q. & Wagemann, C. (2012). *Set-Theoretic Methods for the Social Sciences. A Guide to Qualitative Comparative Analysis*. New York: Cambridge University Press.
- Shadish, W. R., Cook, T. D., & Campbell, D. T. (2001). *Experimental and Quasi-experimental Designs for Generalized Causal Inference*. Boston: Houghton Mifflin Company.
- Simon, H. A. (1954). Spurious correlation: a causal interpretation. *Journal of the American Statistical Association*, 49, 467–479.
- Siy, P. & Chen, C. S. (1972). Minimization of Fuzzy Functions. *IEEE Transactions on Computers*, C-21(1), 100 – 102.

- Spirtes, P., Glymour, C., & Scheines, R. (2000). *Causation, Prediction, and Search* (2 ed.). London: The MIT Press.
- Staab, F. (2007). *Logik und Algebra*. R. Oldenbourg Verlag.
- Suppes, P. (1970). *A Probabilistic Theory of Causality*. Amsterdam: North Holland.
- Thiem, A., Baumgartner, M., & Bol, D. (2016). Still Lost in Translation! A Correction of Three Misunderstandings Between Configurational Comparativists and Regressional Analysts. *Comparative Political Studies*, 49(6), 742–774.
- Thiem, A. & Dusa, A. (2013). Boolean Minimization in Social Science Research: A Review of Current Software for Qualitative Comparative Analysis (QCA). *SAGE*, 31(4), 505–521.
- Urban, D. & Mayerl, J. (2008). *Regressionsanalyse: Theorie, Technik und Anwendung* (3 ed.). Wiesbaden: VS Verlag für Sozialwissenschaften.
- Veitch, E. W. (1952). A chart method for simplifying truth functions. *Proc. Assoc. for Computing Machinery, Pittsburgh Mai 1952*, 127–133.
- Wheeler, J. A. & Zurek, W. H. (Eds.). (1983). *Quantum Theory and Measurement*. Princeton: Princeton University Press.
- Whitesitt, J. E. (1964). *Boolesche Algebra und ihre Anwendungen*. Braunschweig: Vieweg.
- Whitesitt, J. E. & Stumpf, B. (1973). *Einführung in die Boolesche Algebra*. Braunschweig: Vieweg.
- Woodward, J. (2003). *Making Things Happen: A Theory of Causal Explanation*. Oxford University Press.
- Zadeh, L. A. (1965). Fuzzy sets. *Information and Control*, 8, 338–353.