

7 Forschungsmethodisches Vorgehen

Für die Bearbeitung der Fragestellungen wird auf die Daten des Schweizer Kinder- und Jugendsurvey COCON (*Competence and Context*) zurückgegriffen. Das gross angelegte Projekt untersucht die sozialen Bedingungen, Lebenserfahrungen und die psychosoziale Entwicklung von Kindern und Jugendlichen in der Schweiz aus einer Lebenslaufperspektive. Die benötigten Variablen für die intersektionale Analyse sowie jene für Sozialkapital sind in den Daten von COCON enthalten. Des Weiteren kann auch auf Angaben der Lehrpersonen und Hauptbetreuungspersonen zurückgegriffen werden, die insbesondere für das Erstellen eines Sozialkapitalportfolios wertvoll sind.

Neben der Beschreibung der Datengrundlage soll im Folgenden die Auswahl der Sekundäranalyse als Untersuchungsmethode begründet und die Variablenauswahl anhand der theoretischen Konstrukte und methodischen Grundlagen dargestellt werden. Ausführungen zum Datenanalyseverfahren bezüglich latenter Konstrukte schliessen dieses Kapitel ab.

7.1 Datengrundlage

Es besteht eine vertragliche Verpflichtung, bei der Benutzung der Daten von COCON für Sekundäranalysen die Angaben zum Projekt der in der Dokumentation vorgegebenen Beschreibung entsprechend wiederzugeben:

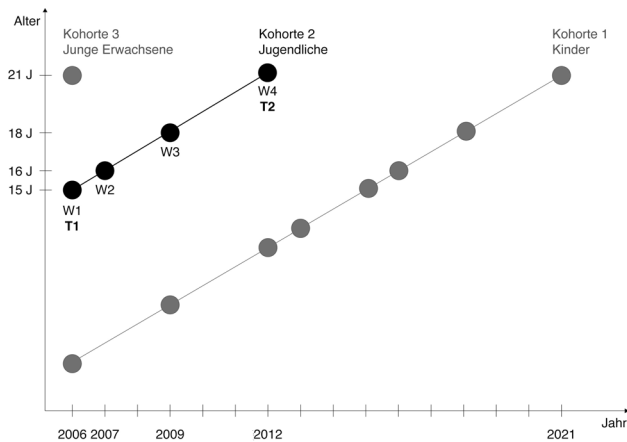
Der Schweizer Kinder- und Jugendsurvey COCON (Competence and Context; www.cocon.uzh.ch) ist ein interdisziplinär angelegtes Forschungsprojekt, das die sozialen Bedingungen des Aufwachsens von Heranwachsenden aus einer Lebenslaufperspektive untersucht. Unter der Leitung von Prof. Dr. M. Buchmann wird diese Längsschnittstudie seit 2006 in der deutsch- und französischsprachigen Schweiz durchgeführt. COCON wurde bisher durch den Schweizerischen Nationalfonds (SNF), das Jacobs Center for Productive

Youth Development sowie die Universität Zürich finanziert. (M. Buchmann et al., 2016)

In der vorliegenden Studie werden Daten der Kohorte 2 der Jugendlichen bei Welle 1 mit 15 Jahren und bei Welle 4 mit 21 Jahren verwendet (Abbildung 6).

Abbildung 6

COCON-Kohorten im Zeitverlauf (vgl. M. Buchmann et al., 2016)



Die drei zentralen Aspekte, die in COCON untersucht werden, sind die Sozialisationskontexte, die institutionalisierten (Status-)Übergänge im Lebenslauf und die individuelle Kompetenzentwicklung. Untersuchungsgegenstand in den Sozialisationskontexten ist die Familie, die Schule mit den Gleichaltrigen (Peers), Freunde und Freundinnen, aber auch die Freizeitgestaltung und die Mediennutzung bezogen auf die interaktive Ausgestaltung und die vorherrschenden Wert- und Handlungsorientierungen:

- Die institutionalisierten (Status-)Übergänge als wichtige Weichenstellungen spielen mit ihren Opportunitäten und Restriktionen eine wichtige Rolle im Entwicklungsprozess von Kindern und Jugendlichen.
- Die Bewältigung von komplexen Übergängen gilt im Sozialisationsprozess als wesentliche Voraussetzung für die gesellschaftliche Integration.
- Der individuellen Kompetenzentwicklung wird im Fokus der Identifikation derjenigen spezifischen Kompetenzen betrachtet, die einen besonde-

ren Stellenwert für die Beeinflussung der Bewältigung der beiden vorge-
nannten Aspekte haben. Neben dem effektiven und effizienten Handeln,
sind auch die Persönlichkeit, die Werthaltungen und die Herausbildung
von sozialen und produktiven Kompetenzen Untersuchungsgegenstand
(M. Buchmann et al., 2016, S. 5).

Tabelle 7

Übersicht der COCON-Befragungswellen der Kohorte 2 (M. Buchmann et al., 2016)

Kohorte 2 - Welle 1	
Befragungszeitraum	März bis Juni 2006
Befragte	Jugendliche, deren Hauptbezugspersonen und Lehrpersonen
Befragungsinstrumente	Jugendliche: persönliches Interview mit Laptop zuhause (CAPI) Hauptbezugspersonen: schriftlicher Fragebogen Lehrpersonen: schriftlicher Fragebogen
Stichprobengröße	N = 1258
Kohorte 2 - Welle 2	
Befragungszeitraum	Mai bis August 2007
Befragte	Jugendliche
Befragungsinstrumente	Jugendliche: telefonisches Interview (CATI)
Stichprobengröße	N = 1162
Kohorte 2 - Welle 2 Zwischenbefragung	
Befragungszeitraum	Mai bis August 2008
Befragte	Jugendliche (nur diejenigen, die zum Zeitpunkt der 2. Welle 2007 noch nicht in einer Lehre oder einer weiterführenden Schule waren)
Befragungsinstrumente	Jugendliche: telefonisches Interview (CATI)
Stichprobengröße	N = 404
Kohorte 2 - Welle 3	
Befragungszeitraum	Mai bis Oktober 2009
Befragte	Jugendliche
Befragungsinstrumente	Jugendliche: persönliches Interview mit Laptop zuhause (CAPI)
Stichprobengröße	N = 952
Kohorte 2 - Welle 4	
Befragungszeitraum	April bis Juli 2012
Befragte	Junge Erwachsene
Befragungsinstrumente	Junge Erwachsene: persönliches Interview mit Laptop zuhause (CAPI)
Stichprobengröße	N = 816

Wie in Tabelle 7 ersichtlich ist, waren die Jugendlichen der Kohorte 2 zu Be-
ginn der Studie im Jahr 2006 15 Jahre alt. Sie standen damit grösstenteils kurz
vor dem Ende der obligatorischen Schule und kurz vor dem Übertritt in eine
Berufsausbildung oder weiterführende Schule. Das Ziel war es, diese Perso-
nen bis zum Alter von 21 Jahren, im Jahr 2012, zu begleiten, um die prägende
Phase des Übertritts von der Schule in den Beruf beziehungsweise in nach-
obligatorische Bildungsangebote zu dokumentieren können. Die Stichprobe

wurde zufällig aus dem vorgesehenen Jahrgang aus zuvor ausgewählten Gemeinden gezogen (M. Buchmann et al., 2016). Die Daten aus COCON sind der Forschungsgemeinschaft über die Datenbank *FORSbase* des Schweizer Kompetenzzentrums für Sozialwissenschaften (FORS) zugänglich.

In der ersten Befragungswelle wurden Informationen sowohl im Querschnitt wie auch im Längsschnitt erfragt. Retrospektiv wurden auf den Monat genau der bisherige Ausbildungsverlauf und die Wohn- und Familiensituation mittels elektronischer Fragebögen erfasst. Der Grossteil der Fragen war geschlossen oder halb-offen, offene Fragen dienten zur Auflockerung, da die Fragebögen doch sehr umfangreich und dadurch aufwändig in der Bearbeitung waren. Die Fragebögen wurden mit Antwortvorlagen ergänzt, um die Beantwortung zu erleichtern und zu beschleunigen. (M. Buchmann et al., 2016).

Kohorte 2 wird, neben den bereits genannten Eigenschaften, ausgewählt, weil für alle Variablen der intersektionalen Kategorien und des Sozialkapitals eine ausreichende Datengrundlage zur Beantwortung der Fragestellung vorhanden ist.

7.2 Stichprobenbeschreibung

Die Stichprobe zu Untersuchungsbeginn umfasste 1258 Schülerinnen und Schüler. Wie in Tabelle 8 ersichtlich, haben 81.5 % der Jugendlichen die gleiche Muttersprache wie die Unterrichtssprache am jeweiligen Schulort. Das Bundesamt für Statistik (BfS) lieferte vor 2018 lediglich Daten zum Ausländer- und Ausländerinnenanteil (ausländische Nationalität, entweder im Ausland oder in der Schweiz geboren) in der obligatorischen Schule. Im Schuljahr 2010/11 betrug dieser 24.3 % (BfS, 2019a). In der vorliegenden Untersuchung wird darauf verzichtet die nationalstaatliche Herkunft allein als Kriterium für einen Migrationshintergrund zu verwenden, da dieser, wie bereits erwähnt, bezogen auf den interaktionistischen Ansatz nur bedingt Aussagekraft hat. Mit 18.3 % der Jugendlichen, die zu Hause primär eine andere Sprache sprechen als am Unterrichtsort, liegt der Anteil etwas tiefer als beim herkömmlichen Ausländerinnen- und Ausländeranteil. Dies kann darauf zurückgeführt werden, dass beispielsweise Jugendliche, die aus Deutschland migriert sind, zu Hause und in der Schule die gleiche Sprache sprechen.

Tabelle 8
Stichprobenbeschreibung T1 (N = 1258)

Kategorien		Absolut	Relativ
Migration (n = 1255)	Gleiche Muttersprache wie Unterrichtssprache	1025	81.5 %
	Andere Muttersprache als Unterrichtssprache	230	18.3 %
Behinderung (n = 1257)	Regelschule	916	91.6 %
	Besonderer Lehrplan	104	8.4 %
SES (n = 1251)	Höher	551	44.0 %
	Tiefer	700	56.0 %
Geschlecht (n = 1258)	Weiblich	683	54.3 %
	Männlich	575	45.7 %

Anmerkung. Das durchschnittliche Alter betrug $M = 15.3$ Jahre (min = 14.5, max = 16.2, $SD = .2$).

Die Anzahl Jugendlicher, die in ihrer obligatorischen Schullaufbahn mit einem besonderen Lehrplan unterrichtet wurden, beträgt in der Stichprobe 8.4 %. In der schweizerischen Gesamtpopulation betrug 2010/11 der Anteil Schülerinnen und Schüler mit besonderem Lehrplan 4.1 %. Der Unterschied der beiden Anteile resultiert aus den Erhebungsmethoden des BfS. Vor 2018 ist nämlich der Lehrplanstatus von Schülerinnen und Schüler mit besonderen Bedürfnissen nach einer dreistufigen Skala erhoben worden:

Die Schülerin/der Schüler wird

1. durchgehend nach Regellehrplan unterrichtet.
2. teilweise nach individuellen, nicht dem Regellehrplan entsprechenden Zielsetzungen unterrichtet. *Kriterium:* Der Unterricht ist in ein bis zwei Fächern nicht auf das Erreichen der Mindestanforderungen des Regellehrplans ausgerichtet.
3. mehrheitlich nach individuellen, nicht dem Regellehrplan entsprechenden Zielsetzungen unterrichtet. *Kriterium:* Der Unterricht ist in drei und mehr Fächern nicht auf das Erreichen der Mindestanforderungen des Regellehrplans ausgerichtet.

Sonderklassen und Sonderschulen werden über das Merkmal »Schulart« erfasst. (BfS, 2010, S. 24)

Diese Herangehensweise hatte zu Folge, dass in den Erhebungen des BfS Schülerinnen und Schüler, die eine Sonderklasse oder eine Sonderschule besuchten, nicht mit dem Kriterium besonderer Lehrplan erfasst wurden. In

der vorliegen Untersuchung werden sowohl die Schülerinnen und Schüler der obengenannten Skala wie auch jene in Sonderklassen und -schulen zusammengefasst. Dieses Konzept bildet die Strukturkategorie Behinderung umfassender ab.

Der sozioökonomische Status wird für die Berechnungen am Median dichotomisiert, dies weil alle Strukturkategorievariablen für die statistischen Berechnungen die gleiche Ausprägung haben müssen, welche durch die nicht zu umgehende Dichotomie der Variable Geschlecht bestimmt ist. Daraus ergibt sich für 44.0 % der Schülerinnen und Schüler ein höherer und für 56.0 % ein tieferer sozioökonomischer Status. Werden die z-standardisierten Werte näher betrachtet, zeigt sich ein weniger einheitliches Bild. Zwar liegt der Mittelwert mit $M = -.01$ sehr nahe am Median, aber die Standardabweichung beträgt $SD = .905$. Die Streuung wird durch die Dichotomisierung weniger gut abgebildet ($M = .56$, $SD = .497$), muss aber aus forschungsmethodischen Gründen so vorgenommen werden.

54.3 % der Studienteilnehmenden sind weiblich. In der Schweizer Gesamtpopulation betrug 2010/11 der Anteil Schülerinnen 48.6 % (BfS, 2019a).

Tabelle 9
Stichprobenbeschreibung T2 (N = 816)

Kategorien		Absolut	Relativ
Migration (n = 816)	Gleiche Muttersprache wie Unterrichtssprache	708	86.8 %
	Andere Muttersprache als Unterrichtssprache	108	13.2 %
Behinderung (n = 816)	Regelschule	795	97.4 %
	Besonderer Lehrplan	21	2.6 %
SES (n = 814)	Höher	391	48.0 %
	Tiefer	423	52.0 %
Geschlecht (n = 816)	Weiblich	448	54.9 %
	Männlich	368	45.1 %

Anmerkung. Das durchschnittliche Alter betrug $M = 21.4$ Jahre (min = 20.4, max = 22.3, $SD = .2$).

Zum zweiten Messzeitpunkt (Tabelle 9) hatte sich die Stichprobe um etwa ein Drittel auf 816 Personen reduziert. Die Verhältnisse in den Kategorien Migration, sozioökonomischer Status und Geschlecht sind in etwas ähnlich zum ersten Messzeitpunkt. In der Kategorie Behinderung ist der relative Anteil an Probandinnen und Probanden, die mit einem besonderen Lehrplan unterrichtet worden waren, deutlich gesunken.

7.3 Durchführung einer Sekundäranalyse

Sekundäranalysen haben im Bereich der quantitativen Methoden eine lange Tradition, in sozialwissenschaftlichen Studien gelten sie als wichtige Herangehensweise der Theoriebildung. Historisch gesehen geht die Sekundäranalyse der Primäranalyse voraus. Durkheim (1897/2009) hat sein Werk *Le suicide*, welches noch heute als methodisches Vorbild der quantitativen Datenanalyse gesehen werden kann, auf der Basis von Sekundärdaten verfasst. Interessanterweise hat Lazarsfeld in der gleichen Zeit, in der er die Fundamente für die Analyse latenter Klassen legte, den Begriff Sekundäranalyse (*secondary analysis*) geprägt, indem er mit seinen Studierenden Sekundäranalysen auf der Grundlage von *The American Soldier* (Stouffer, 1949) durchgeführt und dokumentiert hat (Kendall & Lazarsfeld, 1950). Ab den späten 1950er-Jahren wurden erhobene Daten in Datenarchiven gesammelt und den Forschenden zur Verfügung gestellt. In Europa war das Zentralarchiv für empirische Sozialforschung an der Universität zu Köln, heute Datenarchiv für Sozialwissenschaften des GESIS-Leibniz-Instituts für Sozialwissenschaften (GESIS-DAS), die erste solche Institution (Mochmann, 2014).

Nach Mochmann (2014) sind die Vorteile der Sekundäranalyse vielfältig: Es kann auf qualitativ hochwertige und umfangreiche Datensätze zugegriffen werden, der Fokus liegt stärker auf den theoretischen Zielen und Grundlagen der Studie, es können statistische Analysen auf einem hohen wissenschaftlichen Niveau durchgeführt werden und es gibt sowohl Zeit- wie Kostensparnisse. Aus ethischer und sozialer Perspektive nehmen Sekundäranalysen Rücksicht auf Befragte, was besonders vulnerable Populationen vor Überbefragung schützt. Aus der Sicht der sozialwissenschaftlichen Forschungsgemeinschaft werden ausserdem verfestigte Strukturen und Privilegien aufgebrochen, um unterschiedliche Ideen sowie neues Wissen einzubringen.

Mit der Anwendung von Sekundäranalysen wird unweigerlich ein Feld betreten, das aber gemäss dem Vorbild des »independent researchers« lediglich ein Weg in der soziologischen Forschung sein sollte:

I suggest secondary analysis as only one possible aspect (not a complete style) of a sociologist's research career. In doing this I have tried to locate its very appropriate use in the social structure of sociology, in the sense that it can be used to solve some typical problems faced by different types of independent researchers – again, usually only one aspect of a sociologist's career. Insofar as secondary analysis allows some people to mobilize their meager

resources to tempt a sociological contribution, it can help save time, money, careers, degrees, research interest, vitality, and talent, self-images and myriads of data from untimely, unnecessary, and unfortunate loss. (Glaser, 1963, S. 14)

In gewissem Sinne warnt Glaser davor, Sekundäranalysen als einfacheren wissenschaftlichen Weg zu beschreiten. Burzan (2005) meint dazu, dass deswegen in jedem Fall unbedingt vorab geklärt werden muss, ob die vorhandenen Daten nicht nur zur eigenen Fragestellung, sondern zur gesamten Anlage der Studie passen.

Je nachdem sind anschliessend verschiedene Formen der Sekundäranalyse möglich. In der *Supraanalyse* oder *transzendierende Analyse* werden die bereits erhobenen Daten unter neuen theoretischen und empirischen Forschungsperspektiven ausgewertet. Die *ergänzende Analyse* führt Fragestellungen und Forschungsperspektiven weiter, die in der ursprünglichen Bearbeitung der Primäranalyse entstanden sind. In der *Reanalyse* werden die Daten aus der Primäranalyse unter derselben Fragestellung ausgewertet, gegebenenfalls mittels anderer Methoden (Heaton, 2008).

Die Auswahl der Variablen erfolgt nach dem gleichen Prinzip der Passung zwischen vorhandenen Daten und Fragestellung, gegebenenfalls auch nach Hypothesen. Indikatoren und Variablen, die in die Berechnungen aufgenommen werden sollen, müssen zu den Dimensionen passen. Auf einer allgemeinen Ebene muss auch die Qualität der Daten bezogen auf das Erhebungsinstrument beurteilt werden. Bei quantitativen Analysen bedeutet dies, dass die Gütekriterien Reliabilität, Validität und Objektivität der verwendeten Skalen überprüft werden. Je nach Vorgehensweise können die konkreten Hypothesen entlang der vorhandenen Daten formuliert werden. Bei dieser Abstimmung ist es wichtig, die Kongruenz zum Erkenntnisinteresse und dem theoretischen Hintergrund kontinuierlich zu überprüfen. Idealerweise liegen aber Daten vor, die bezogen auf die ursprüngliche Fragestellung und Hypothesen verwendet werden können.

Für die vorliegende Arbeit ist die Voraussetzung der Passung zwischen vorhandenen Daten und den Erkenntnisinteressen, dem theoretischen Hintergrund und den zu untersuchenden Dimensionen erfüllt. In der Projektdokumentation zu COCON wird beschrieben, dass der Einfluss der Familie, der Schule, von Gleichaltrigen und der weiteren sozialen Umgebung auf die Entwicklung von Kindern und Jugendlichen untersucht wird. Des Weiteren liegt ein wichtiger theoretischer und empirischer Fokus auf der Wechselwirkung

zwischen elementaren Handlungskompetenzen und Beziehungsressourcen (M. Buchmann et al., 2016). Die vorliegende Untersuchung folgt dem Prinzip einer Supraanalyse.

7.4 Operationalisierung der Variablen

In der vorliegenden Untersuchung liegen alle Skalen bereits vor. Die Herausforderung dieser Sekundäranalyse liegt darin, die Indikatoren auf der Basis vorausgehender Studien und theoretischer Grundlagen so auszuwählen, dass sie als Klassifikationsmerkmale operationalisiert werden können und bestenfalls in der Datenanalyse konfirmatorisch bestätigt werden können.

Sowohl für das Konstrukt Sozialkapital wie auch für die intersektionale Analyse von Behinderung, Migration, sozioökonomischem Status und Geschlecht liegen in der COCON-Studie valide Daten vor. In Bezugnahme auf die bereits dargestellten theoretischen und empirischen Grundlagen werden im Folgenden ausschliesslich die technischen Aspekte der Operationalisierungen erläutert.

Behinderung

In der ursprünglichen Befragung von COCON wurden keine Formen von Behinderungen oder Beeinträchtigungen erhoben. Aus den Daten können von zwei Variablen Indikatoren abgeleitet werden. Einerseits wurde der sprachfreie Intelligenztest *Culture Fair Test* (CFT) in einer Kurzversion angewendet, und andererseits wurde nach Ausbildungsepisoden gefragt, in welchen Unterricht nach besonderem Lehrplan mittels *Besuch der Integrationsklasse* oder *Besuch einer Sonderklasse/Kleinklasse* stattgefunden hatte. Auf die Verwendung des Intelligenztests wurde gemäss Dokumentation von COCON aus zwei Gründen verzichtet: Einerseits handelt es sich bei den kognitiven Fähigkeiten um ein individuelles, psychologisches Merkmal, welches andererseits durch die Fokussierung des CFTs auf das logische Denken wahrscheinlich ungenügend abgebildet wird (C. Jacobs & Petermann, 2007).

So wird für die vorliegende Untersuchung die Kategorie Behinderung ausschliesslich über das institutionelle Merkmal des besonderen Lehrplans beziehungsweise der besonderen Förderung gebildet.

Migration

Für die Erfassung der Variable Migration konnte im Datensatz auf die Indikatoren Staatsangehörigkeit, Geburtsort, Lebensdauer in der Schweiz sowie Muttersprache und Familiensprache zurückgegriffen werden. Die Interviews waren in der ursprünglichen Befragung von COCON entweder auf Französisch oder Deutsch beziehungsweise Schweizerdeutsch geführt worden.

Aus diesen Informationen wird die dichotome Variable Migration als Diskrepanz zwischen Interviewsprache und hauptsächlicher Familiensprache gebildet.

Sozioökonomischer Status (SES)

Der SES wird aus den im ursprünglichen Datensatz von COCON erfassten höchsten Bildungsabschluss sowie dem höchsten Beruf der Eltern errechnet.

Die Angaben zum ausgeübten Beruf liegen in vier Codierungsvarianten (ISCO-88, ISEI-08, SIOPS und BfS 5- und 8-stufig) für jeweils Vater und Mutter vor. Mit dem von Ganzeboom et al. (1992) entwickelten *International Socio-Economic Index of occupational status* (ISEI) kann der sozioökonomische Status im internationalen Vergleich errechnet werden. Dies, indem eine Konversionssyntax für die *International Standard Classification of Occupations* (ISCO) oder für die *Standard Index of Occupational Prestige Scala* (SIOPS), auch Treiman-Index genannt (Treiman, 1977), angewendet wird. Beim ISEI wird davon ausgegangen, dass jede berufliche Tätigkeit einen bestimmten Bildungsgrad erfordert und entsprechend entlohnt wird, was wiederum in Chancen zur Teilhabe an Macht umgesetzt wird.

Geschlecht

Es kann nicht genau nachvollzogen werden, ob die Jugendlichen in der ursprünglichen Befragung von COCON nach ihrem Geschlecht gefragt wurden oder ob eine Einschätzung durch die interviewende Person vorgenommen wurde. Es wird eher von letzterem ausgegangen, da der gesamte Datensatz an beiden Messzeitpunkten für die Variable Geschlecht keine fehlenden Werte aufweist.

Die Daten geben eine dichotome Variable mit den Ausprägungen *männlich* und *weiblich* wieder.

Sozialkapital

In der ursprünglichen Befragung von COCON wurden verschiedene Dimensionen von Sozialkapital erhoben. Neben Fragen zu Freizeitaktivitäten mit Kolleginnen und Kollegen, zur Anzahl Freundinnen und Freunden sowie zur Anzahl Kolleginnen und Kollegen wurden auch organisierte Freizeitaktivitäten erfragt. Klassische Sozialkapitalvariablen wie allgemeine Werte (Deutsches Jugendinstitut und Infas, 2003), politische Handlungsorientierung und soziale Verantwortung (Grob & Maag Merki, 2001), allgemeines Vertrauen (Deutsches Jugendinstitut und Infas, 2003) wurden mit Fragen zur emotionalen Nähe zu Bezugspersonen (Gerhard et al., 1997) und Aushandlungsprozessen mit Gleichaltrigen (Parker & Asher, 1993) ergänzt. Weiter wurden auch die Zusammenarbeit zwischen Schule und Eltern, Familienstrukturen und -konstellationen oder das Vertrauen in einen sozialen Aufstieg erfasst.

In Tabelle 10 wird die Auswahl der Indikatoren, die das Konstrukt Sozialkapital für die vorliegende Untersuchung beim ersten Messzeitpunkt (T1) abbilden, dargestellt. Die grosse Anzahl Indikatoren im Datensatz wurde faktoranalytisch auf ein Set von 25 in sechs Dimensionen reduziert. Angaben zum schrittweisen methodischen Vorgehen finden sich im Ergebnisteil.

Tabelle 11 zeigt die Dimensionen und Indikatoren von Sozialkapital am zweiten Messzeitpunkt (T2). Da die Jugendlichen in der Zwischenzeit die obligatorische Schule verlassen hatten, entfällt die Dimension des schulbezogenen Sozialkapitals. Leider liegen im Datensatz keine Indikatoren vor, welche zu berufsbezogenem Sozialkapital hätten verarbeitet werden können. Die wiederum relativ grosse Anzahl Indikatoren wurde auf 21 reduziert.

7.5 Begründung der Auswahl der statistischen Verfahren

Das Fehlen einer herausragenden Methode für intersektionale Analysen (Bowleg, 2008; Cho et al., 2013; Nash, 2008) verlangt danach, dass der methodische Rahmen aus den jeweiligen Ressourcen, welche innerhalb einer Disziplin vorhanden sind, angewendet wird, beziehungsweise welche für die Bearbeitung des Erkenntnisinteressens zielführend sind. Aus den historisch gewachsenen Bezugnahmen kommen in intersektionalen Analysen meistens sozialkonstruktivistische und interpretative Methoden zur Anwendung (Rodriguez et al., 2016). Im deutschen Sprachraum wird dabei oft auf den mehrerebenenanalytischen Vorschlag von Winker und Degele (2009) zurückgegriffen. Quantitative Zugänge, hingegen, haben sich noch nicht durchgesetzt und werden

Tabelle 10
Dimensionen, Items und Ursprungsskalen von Sozialkapital T1

Dimensionen und Items	Skala
<i>Schulbezogenes Sozialkapital</i>	
Wenn ich in der Schule eine schwierige Aufgabe machen muss, fange ich damit sobald als möglich an.	EJR
Ich strenge mich in der Schule/beim Schaffen sehr an.	TREE
Auch bei einer mühsamen Arbeit gebe ich nicht auf, bis ich fertig bin.	EJR
Ich lerne so fleissig wie möglich.	TREE
Auch wenn ich bei einer Aufgabe auf Schwierigkeiten stosse, bleibe ich hartnäckig dran.	EJR
<i>Elternhausbezogenes Sozialkapital: Strenge Kontrolle</i>	
Meine Mutter/mein Vater duldet häufig keinen Widerspruch.	EJ
Meine Mutter/mein Vater erwartet, dass ich mich immer allem füge, was sie/er mir vorschreibt.	EJ
Meine Mutter/mein Vater und ich geraten häufig aneinander, aber trotzdem haben wir uns sehr gern.	GeAmb
<i>Empathie</i>	
Wenn ich sehe, dass es einem anderen Jugendlichen schlecht geht, habe ich Mitleid mit ihm oder ihr.	EmpMes
Ich habe Mitgefühl mit anderen Jugendlichen, wo traurig sind oder Schwierigkeiten haben.	EmpMes
Wenn ich sehe, dass jemand geplatzt wird, habe ich Mitgefühl mit ihm oder ihr.	EmpMes

nur vereinzelt eingesetzt (Dubrow & Ilinca, 2019; McCall, 2005). Zunehmend sprechen sich Forschende für einen Methodenpluralismus in der intersektionalen Forschung aus, der es erlauben würde, partizipative, längsschnittliche oder mehrebenenanalytische Untersuchungen durchzuführen (Creswell, 2014; Heiskanen et al., 2015; Rodriguez et al., 2016; Woodhams & Lupton, 2014). Gerade letztere können für sozialwissenschaftliche Fragestellungen gewinnbringend sein, da sie die Beziehungen zwischen den Individuen sowie strukturelle Mechanismen darstellen können.

Weldon (2006) argumentiert ebenfalls für eine sozialstrukturelle Herangehensweise und lehnt gleichzeitig den statistischen Interaktionsbegriff ab, da es der Methode nicht gelinge, die der Intersektionalität inhärente Idee der Wechselwirkungen und Interdependenzen abzubilden. In diesem Sinne wird davor gewarnt, addierende und multiplizierende Methoden als intersektionale Modelle zu gestalten.

It is not often recognized that structural analysis is *required* by the idea of intersectionality: It is the intersection of social *structures*, not identities, to which the concept refers. We cannot conceptualize »interstices« unless we have a concept of the structures that intersect to create these points of interaction. (Weldon, 2006, S. 239)

Tabelle 10 (Fortsetzung)

Dimensionen, Items und Ursprungsskalen von Sozialkapital T1

Dimensionen und Items	Skala
<i>Elternhausbezogenes Sozialkapital: Emotionale Nähe</i>	
Wie häufig zeigt Ihnen Ihre Mutter/Ihr Vater, dass sie/er Sie wirklich gern hat?	FndT
Wie häufig gibt Ihnen Ihre Mutter/Ihr Vater das Gefühl, dass sie/er Ihnen wirklich vertraut?	FndT
Wenn Sie etwas machen, wo Ihre Mutter/Ihr Vater gut findet: Wie häufig zeigt sie/er Ihnen dann, dass sie/er sich darüber freut?	FndT
Wie häufig redet Ihre Mutter/Ihr Vater mit Ihnen über das, was Sie machen und erlebt haben?	FndT
<i>Allgemeine Werte</i>	
Alle Menschen gleichberechtigt behandeln	DJI
Ungleichheiten zwischen Menschen abbauen	DJI
Sich selbstverwirklichen	DJI
Im Umgang mit anderen fair sein	DJI
<i>Peerbezogenes Sozialkapital</i>	
Mein/e Freundin und ich helfen einander.	FEEA
Mein/e Freundin und ich vertrauen uns gegenseitig bei Ratschlägen.	FEEA
Mit meinem Freund/meiner Freundin rede ich darüber, was mir im Leben wichtig ist.	FEEA
Mit meinem Freund/meiner Freundin bespreche ich Probleme, wo ich mit meinen Eltern habe.	FEEA
Reden Sie über Probleme?	DJI
Diskutieren Sie zusammen?	DJI

Anmerkungen. Die Angaben zu den Skalen können der Projektdokumentation der COCON-Studie entnommen werden (M. Buchmann, 2006). Abkürzungen: EJR: Eidgenössische Jugend und Rekrutenbefragung (2000/2001), TREE: Transitionen von der Erstausbildung ins Erwerbsleben (2001), Eij: Entwicklung im Jugendalter (1986), Ge-Amb: Generationenambivalenzen operationalisieren: Konzeptuelle, methodische und forschungspraktische Grundlagen (2002), EmpMes: Empathy and its measurement (2003), FndT: Familienentwicklung nach der Trennung (1997), DJI: DJI-Jugendsurvey (2003), FEEA: Fragebogen zur Erfassung von Entwicklungsnormen in der Adoleszenz (2003). Kodierungen: 1 bis 6, stimmt gar nicht bis stimmt völlig (schulbezogenes Sozialkapital, peerbezogenes Sozialkapital, elternhausbezogenes Sozialkapital und Empathie), 1 bis 10, gar nicht wichtig bis extrem wichtig (allgemeine Werte).

Darauf aufbauend plädieren Dubrow und Ilinca (2019) für die Verwendung von *Faktorenanalysen*. Denn der Faktorenanalyse kann es gelingen, Konstrukte abzubilden, die zu komplex sind, um direkt gemessen zu werden, und intersektionale Gruppen sind komplexe Gruppen. Latente Konstrukte sind dafür geeignet, die Faktorenanalyse allein erlaubt aber nur die Bezugnahme auf ein gemeinsames Konstrukt. Anders gesagt eignet sich die Faktorenanalyse dann, wenn lediglich die Verhältnisse intersektionaler Kategorien, wie beispielsweise Geschlecht, ethnische Herkunft und sozioökonomischer Status ohne weitere sozialstrukturellen Aspekte untersucht werden sollen.

Tabelle 11
Dimensionen, Items und Ursprungsskalen von Sozialkapital T2

Dimensionen und Items	Skala
<i>Freizeitaktivitäten mit Kolleginnen und Kollegen</i>	
Wie oft reden Sie zusammen über Probleme?	DJI
Wie oft diskutieren sie zusammen?	DJI
Wie oft diskutieren sie zusammen?	DJI
Wie oft reden Sie über Sachen wo sie selber beschäftigen?	DJI
Wie oft diskutieren Sie zusammen über Ihr späteres Leben?	DJI
Wie oft reden Sie über vertrauliche Sachen?	DJI
<i>Allgemeine Werte</i>	
Im Umgang mit anderen fair sein	DJI
Viel leisten	DJI
Sich selbst verwirklichen	DJI
Pflichtbewusst sein	DJI
Ein aufregendes und spannendes Leben führen	DJI
Sich anstrengen	DJI
Eigene Fähigkeiten entwickeln	DJI
<i>Gesellschaftliche Verantwortung</i>	
Der Reichtum sollte weltweit gerechter verteilt werden.	EJR
Es ist eine wichtige Aufgabe vom Staat, hilfsbedürftige Menschen zu unterstützen.	EJR
Ich finde es wichtig, dass der Staat Arbeitslose unterstützt.	EJR
Wenn man gegen etwas ist, dann muss man protestieren.	EJR
Wer gut verdient, sollte einen grösseren Beitrag für die Allgemeinheit leisten.	EJR
<i>Empathie</i>	
Wenn ich sehe, dass jemand geplatzt wird, habe ich Mitgefühl mit ihm oder ihr.	EmpMes
Ich habe Mitgefühl mit anderen Menschen, wo traurig sind oder Schwierigkeiten haben.	EmpMes
Wenn ich sehe, dass es einem anderen Menschen schlecht geht, habe ich Mitleid mit ihm oder ihr.	EmpMes

Anmerkungen. Die Angaben zu den Ursprungsskalen können der Projektdokumentation der COCON-Studie entnommen werden (M. Buchmann, 2012). Abkürzungen: DJI: DJI-Jugendsurvey (2003), EJR: Eidgenössische Jugend und Rekrutenbefragung (2000/2001), EmpMes: Empathy and its measurement (2003). Kodierungen: 1 bis 6, nie bis täglich (Freizeitaktivitäten mit Kolleginnen und Kollegen), 1 bis 6, stimmt gar nicht bis stimmt völlig (allgemeine Werte, gesellschaftliche Verantwortung, Empathie).

Solga et al. (2009) postulieren für die intersektionale Analyse von sozialen Ungleichheiten unter Bezugnahme der Nutzung von Handlungsressourcen und Lebensstil- beziehungsweise Milieukonzepten eine *Clusteranalyse*. Während soziale Klassen von aussen vorgegebene Gruppenzugehörigkeiten definieren, werden mittels des empirisch-induktiven Verfahrens Klassifikationen abgebildet. Personen werden typisiert und mit anderen, die ähnliche Merkmalskombinationen aufweisen, zu einem Cluster zusammengefasst. Die Beschreibung und Benennung der einzelnen Cluster erfolgt erst, nachdem sich die Clusterprofile im Modell eindeutig abgrenzen.

An dieser Stelle gilt es auch zu diskutieren, ob ein personen- oder ein variablenbezogener Ansatz beziehungsweise eine Kombination von beidem für die Analyse verwendet werden. Grundsätzlich sollen statistische Modelle dem Charakter der zu untersuchenden Strukturen und Prozesse entsprechen. Für die Verwendung des personenbezogenen Ansatzes votiert beispielsweise Magnusson (1999). Er begründet dies folgendermassen:

If empirical research is to help us to understand the mechanisms at work in the synchronization process at the level of individuals, it needs a theoretical framework for studies on specific elements of the processes, a framework that regards the individual as an integrated organism functioning as an undivided totality in an integrated person-environment system. (Magnusson, 1999, S. 208)

Im Gegensatz dazu ist der variablenbezogene Ansatz an Beziehungen zwischen Variablen und von Variablen zu spezifischen Kriterien, an der Stabilität von Variablen, deren Beziehungen zu Umweltfaktoren und Entwicklungsbedingungen interessiert. Grundsätzlich dient der variablenbezogene Ansatz dazu, Gruppen auf der Aggregatebene zu beschreiben, jedoch wird häufig fälschlicherweise davon ausgegangen, dass solche Berechnungen und auf diesem Wege gewonnene Ergebnisse angemessen sind, um Verhältnisse und Zusammenhänge innerhalb eines Individuums zu beschreiben (Magnusson, 1999, S. 215). Weiter können personenbezogene Ansätze auf einem Kontinuum von induktiv zu deduktiv beziehungsweise von explorativ bis konfirmatorisch verstanden werden. Wurde induktive Forschung als *Bottom-up*, datengestützte und explorative Forschung konzipiert, spiegeln deduktive Ansätze das *Top-down* Vorgehen mit a-priori Hypothesen in Kontext eines klar definierten theoretischen Modells wider (Woo et al., 2017). In der aktuellen Forschung werden, bezogen auf die methodischen Ansätze der vorliegenden Arbeit, zunehmend auch Anwendungen von qualitativ hochwertigen induktiven Vorgehen postuliert, da rein deduktive Ansätze den Fortschritt in der Theoriebildung begrenzen können (Jebb et al., 2017).

Der entscheidende Unterschied zwischen dem personenbezogenen und dem variablenbezogenen Ansatz bezüglich der Behandlung von Daten besteht, in Abhängigkeit von unterschiedlichen theoretischen Modellen, in der Interpretation des Wertes, der die Position eines Individuums auf einer latenten Dimension anzeigt. Im variablenbezogenen Ansatz erlangt jeder einzelne Datenpunkt seine Bedeutung in Abhängigkeit von seiner Position relativ zur Position anderer Individuen auf der gleichen Dimension. Im personenbezo-

genen Ansatz erlangt jeder einzelne Datenpunkt die Bedeutung seiner Position innerhalb eines Musters von Daten für das gleiche Individuum. Individuelle Unterschiede werden demnach in Mustern von Daten für die relevanten Dimensionen ermittelt. Es kann damit von einer holistischeren Herangehensweise gesprochen werden.

Eine Möglichkeit, die den personenbezogenen Ansatz in latenten Konstrukten modelliert, ist die Analyse latenter Klassen, *Latent Class Analysis* (LCA). Es handelt sich um ein Verfahren der Klassifizierung von Objekten (z.B. Personen) mithilfe einer Anzahl von vorab an ihnen erhobenen beziehungsweise gemessenen Merkmalen (manifeste Variablen). Die Aufteilung der Objekte in homogene Klassen wird aufgrund einer latenten Variablen vorgenommen, die aus den direkt beobachtbaren Variablen, den Indikatoren, abgeleitet wird. Personen aus einer mittels der LCA identifizierten Klasse zeigen beispielsweise ein ähnliches (homogenes) Antwortverhalten bezüglich der Variablen, die in der Analyse berücksichtigt werden. Personen aus verschiedenen Klassen zeigen hinsichtlich der entsprechenden erhobenen Merkmale unterschiedliche (heterogene) Antwortmuster beziehungsweise -profile. Ein Vorteil der LCA besteht unter anderem darin, dass Variablen mit unterschiedlichem Skalenniveau simultan analysiert werden können. Insbesondere die häufig in den Sozialwissenschaften vorkommenden nominalskalierten Variablen können angemessen in die Analyse aufgenommen werden. Die LCA liefert zudem Statistiken, mittels derer, unter Angaben einer bestimmten Irrtumswahrscheinlichkeit, entschieden werden kann, wie viele homogene Subgruppen in der untersuchten Stichprobe von Objekten vorliegen und wie gross diese sind. Damit können mittels der LCA neben explorativen Analysen auch hypothesenprüfende (konfirmatorische) Analysen durchgeführt werden. Für jede Person wird ausserdem angegeben, mit welcher Wahrscheinlichkeit sie den jeweiligen Gruppen beziehungsweise Klassen angehört (Bacher & Vermunt, 2010, 554f.).

In der Ausgangslage ist die LCA ein exploratives Verfahren, es wird jeweils lediglich vermutet, dass hinsichtlich bestimmter Merkmale oder Kategorien latente Klassen vorliegen könnten. Damit diese ex-post Vorgehensweise nicht zu unbefriedigenden Ergebnissen oder dem voreiligen Schluss führt, dass die LCA ungeeignet sei, ist möglichen Ursachen auf den Grund zu gehen. So können beispielsweise die ausgewählten Variablen für die Klassenbildung unbrauchbar sein. In diesem Fall empfiehlt sich ein deduktives, konfirmatorisch orientiertes Vorgehen, welches bestenfalls bereits vor der Datenerhebung zur Anwendung kommt (Bacher & Vermunt, 2010).

Zusammengefasst kann gesagt werden, dass die Vorteile der LCA darin liegen, dass kein Ähnlichkeitsmass festgelegt werden muss, es keinen Algorithmus für die Clusterung braucht und die Distanz zwischen den Clustern nicht definiert werden muss. Es wird eine probabilistische Zuordnung vorgenommen, welche auch höherdimensionale Zusammenhänge abbildet. Die einer LCA zugrunde liegenden Daten müssen auch nicht normalverteilt sein, was als spezifischer Vorteil der Methode für sozialwissenschaftliche Untersuchungen gesehen werden kann: »Since the social world seems to have been created with less multivariate normality than many researchers are willing to assume, it is likely that latent class analysis will continue to enjoy an increasingly prominent role in social research« (McCutcheon, 1987, S. 79).

Für die Darstellung der Zusammenhänge im Konstrukt Sozialkapital sowie für die Reduktion der ursprünglichen Indikatoren aus der COCON-Studie, wird auf die exploratorische Faktorenanalyse, *Exploratory Factor Analysis* (EFA), zurückgegriffen. Moosbrugger und Schermelleh-Engel (2012) begründen die Anwendung einer EFA wie folgt:

Die EFA ist ein Verfahren, das immer dann zur Anwendung kommt, wenn der Untersucher die Anzahl der einem Datensatz zugrunde liegenden Faktoren analysieren möchte, jedoch keine konkreten Hypothesen über die Zuordnung der beobachteten Variablen zu den Faktoren hat. (Moosbrugger & Schermelleh-Engel, 2012, 236f.)

In der vorliegenden Untersuchung dient die EFA einerseits dazu eine Auswahl der Indikatoren vornehmen und andererseits um in der Interpretation der latenten Klassen auf die Faktoren zurückgreifen zu können. Die finalen Modelle werden dann wiederum mit den extrahierten Indikatoren, welche am höchsten auf die jeweiligen Faktoren laden, gerechnet. Die EFA ermöglicht ausserdem, dass ein Portfolioansatz beibehalten werden kann, in welchem sowohl aus verschiedenen Informationsquellen, wie auch auf verschiedenen Ebenen die wahrgenommenen und tatsächlichen Beziehungsressourcen als Sozialkapital abgebildet werden (Rose et al., 2013).

In Tabelle 12 wird gezeigt, wie latente Variablenmodelle eingeordnet werden können. Dies unabhängig davon, ob die latente Variable kategorisch oder kontinuierlich ist und ob die Indikatorvariablen als kategorisch oder kontinuierlich behandelt werden. Manchmal sind die Unterscheidungen zwischen den verschiedenen Modellen etwas arbiträr, auch weil mit modernen Statistikprogramme gemischte Modelle gerechnet werden können. Die Darstellung

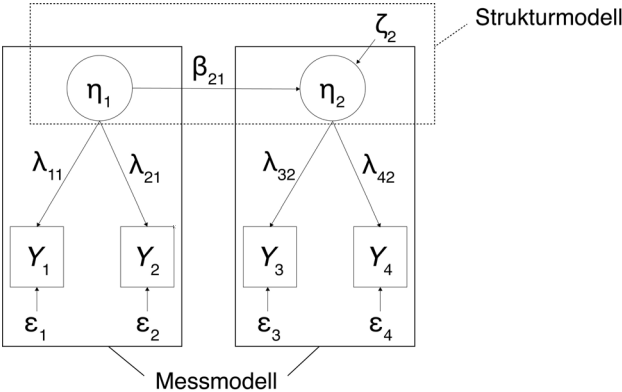
dient demnach lediglich einer Veranschaulichung, um zu zeigen, wie sich LCA und EFA zu anderen latenten Variablenmodellen einordnen lassen.

Tabelle 12
Modellierungen zur Analyse latenter Variablen (vgl. L. M. Collins & Lanza, 2013, S. 7)

	Kontinuierliche latente Variablen	Kategoriale latente Variablen
Kontinuierliche manifeste Variablen	Faktorenanalyse	Latente Profilanalyse
Kategoriale manifeste Variablen	Probabilistische Testtheorie	Latente Klassenanalyse

In der Modellierung mit latenten Variablen, werden diese nicht direkt gemessen, sondern von mindestens zwei manifesten Variablen auch Indikatoren oder Items genannt. In der LCA entsprechen die latenten Klassen und in der EFA die Faktoren den latenten Variablen. In Abbildung 7 ist die Modellierung mit latenten Variablen demnach der messende Anteil, das Messmodell, eines sogenannten Strukturgleichungsmodells, *Structural Equation Model* (SEM).

Abbildung 7
Beispiel für ein einfaches lineares Strukturgleichungsmodell (nach Geiser, 2011, S. 42)



Das abgebildete Strukturgleichungsmodell hat zwei latente Variablen (η_1 und η_2), die jeweils durch zwei Indikatoren (Y_1 und Y_2 , sowie Y_3 und Y_4) gemessen werden. Die Parameter λ_{11} , λ_{21} , λ_{32} und λ_{42} sind Faktorladungen. Die Variablen ε_1 bis ε_4 , bezeichnen Fehlervariablen für die Indikatoren (Geiser, 2011).

Strukturgleichungsmodelle werden in den Sozialwissenschaften verbreitet eingesetzt, da Messfehler in der Verwendung latenter Variablen explizit berücksichtigt werden können und sich dadurch Zusammenhänge korrekter schätzen lassen als in Basisanalysen, die auf der direkten Betrachtung fehlerbehafteter Variablen beruhen. Als zweiter wichtiger Vorteil erlauben es Strukturgleichungsmodelle auch komplexe Theorien einer empirischen Prüfung zwischen den Zusammenhangsstrukturen der Variablen zu testen und verschiedenen Modelle miteinander zu vergleichen. Die Flexibilität von Strukturgleichungsmodellen zeigt sich in deren vielfältigen Anwendung zur Bearbeitung komplexer Fragestellungen. So können Strukturgleichungsmodelle bei Veränderungsmessungen, bei situationsbedingten Einflüssen auf psychologische Messungen, zur Messinvarianztestung, zur Analyse latenter Veränderungen und latenter Wachstumskurven (L. M. Collins & Lanza, 2013) oder in der Auswertung multimethodal erhobener Daten (Nussbeck & Eid, 2015) angewendet werden.

Entgegen der auch unter Forschenden weit verbreiteten Annahme handelt es sich bei der LCA nicht um ein neues Verfahren, sondern um eines, das wie viele andere auch, in neuerer Zeit mehr Beachtung findet und weiterentwickelt wird. Clogg (1981) hat zu Beginn der 1980er-Jahre in seinem wegweisenden Beitrag *New Developments in Latent Structure Analysis* geschrieben:

This paper surveys new developments in latent structure analysis, a statistical method somewhat related to factor analysis but different from it in important ways. This method has a rich tradition on social research beginning with the studies in item analysis occasioned by World War II. Lazarsfeld's creative genius was largely responsible for the conceptual underpinning of the method. (Clogg, 1981, S. 215)

7.6 Analyse latenter Klassen

Die Analyse latenter Klassen, *Latent Class Analysis* (LCA) oder auch *Latent Structure Analysis* (Gibson, 1959; Goodman, 1974b, 1974a; Lazarsfeld & Henry, 1968),

ist ein multivariates statistisches Verfahren, welches es erlaubt Individuen in homogene Untergruppen, den latenten Klassen, einzuteilen.

Bei der Analyse wird von verschiedenen Unbekannten ausgegangen, welche bestimmt werden sollen (vgl. Bacher & Vermunt, 2010; Geiser, 2011; Gollwitzer, 2012):

- Wie viele latente Klassen lassen sich in einer Stichprobe auf der Basis von beobachteten Antwortmustern finden?
- Wie viele Individuen gehören den einzelnen Klassen an?
- Welche Individuen gehören welcher Klasse an?

Die LCA ist als Modell der probabilistischen Testtheorie, beziehungsweise *Item Response Theory* (IRT) oder *Latent Trait Theory*, zu verstehen. Demnach ist die Zugehörigkeit eines Individuums zu einer Klasse nie deterministisch, sondern probabilistisch. Eine Person, beziehungsweise ihr Antwortmuster, wird mit einer gewissen Wahrscheinlichkeit einer bestimmten latenten Klasse zugeordnet. Ebenfalls geschätzt wird die relative Klassengrösse. Die Anzahl der Klassen hingegen sollte im besten Fall theoriegeleitet festgelegt werden. Diese Vorgabe kann anschliessend durch die Anwendung verschiedener Modelllösungen hinsichtlich ihrer Gütekriterien verglichen werden (Gollwitzer, 2012).

Neben diesen Verwandtschaften zeigt die LCA Ähnlichkeiten zur Clusteranalyse. Im Unterschied zur LCA, bei welcher die Anzahl Klassen theoretisch abgeleitet werden, ist die Zahl der Cluster in der Regel unbekannt und soll empirisch ermittelt werden. Dazu werden bestimmte Verteilungsparameter in den Klassifikationsmerkmalen angewendet, welche in der LCA durch Verteilungsannahmen bezüglich der Klassifikationsmerkmale ersetzt werden. Die modellbasierten, beziehungsweise theoretisch abgeleiteten Verteilungsannahmen haben den Vorteil, dass die Zahl der latenten Klassen formal besser abgesichert ist als die Clusterzahl. Auf der anderen Seite bietet die Clusteranalyse die Möglichkeit durch ein hierarchisches Verfahren auch kleinere Stichproben zu untersuchen (Bacher & Vermunt, 2010).

Ein LCA-Modell für eine dichotome Variable lässt sich gemäss Geiser (2011, S. 236) für einen Indikator wie folgt notieren und erklären:

$$p(\mathbf{X}_{vi} = \mathbf{1}) = \sum_{c=1}^C \pi_c \pi_{ic}.$$

(Formel 1)

Dabei entspricht $\mathbf{p}(\mathbf{X}_{vi} = \mathbf{1})$ der Wahrscheinlichkeit, dass ein Individuum $\mathbf{v}$ auf einem Indikator $\mathbf{i}$ einen Wert 1 hat. Der Klassengrößenparameter $\pi_{\mathbf{c}}$ steht für die unbedingte Wahrscheinlichkeit, der latenten Klasse $\mathbf{c}$ anzugehören. Das Modell nimmt an, dass jedes Individuum einer und nur einer Klasse angehört. Somit gilt

$$\sum_{\mathbf{c}=1}^{\mathbf{C}} \pi_{\mathbf{c}} = \mathbf{1}.$$

(Formel 2)

Der Parameter $\pi_{\mathbf{ic}}$ gibt die Wahrscheinlichkeit für den Wert 1 bei einem Indikator $\mathbf{i}$ in Klasse $\mathbf{c}$ an.

Ziel der LCA ist eine möglichst zuverlässige und genaue Schätzung der unbekannten Parameter. Es ist zu beachten, dass die Anzahl der Klassen, die benötigt wird, um die beobachteten Antwortmuster angemessen abbilden zu können, kein zu schätzender Parameter ist.

Die LCA kann sowohl explorativ wie auch konfirmatorisch modelliert werden. Die explorative Anwendung kommt dann zum Zuge, wenn vermutet wird, dass gemessene Daten Teil einer komplexeren Struktur sind, diese aber gegebenenfalls nicht theoretisch belegt werden können oder die zugrunde liegende Theorie auf ihre Angemessenheit befragt werden soll. In diesem Sinne sind für eine explorative LCA auch keine expliziten Hypothesen nötig (McCutcheon, 1987).

Parameterschätzung und Überprüfung der Modellgüte

Bestenfalls wird die Anzahl latenter Klassen aus der Theorie abgeleitet. Da dies aber nicht immer möglich ist, kann auf eine Anzahl wichtiger Kriterien zur Beurteilung der Güte von LCA-Modellen zurückgegriffen werden. In der explorativen Vorgehensweise wird die Modellgüte, *Modelfit*, bezüglich der Anzahl Klassen mithilfe von statistischen Indizes, *indices*, verglichen. Das Modell mit der besten Datenpassung und adäquaten Parameterschätzungen wird zur Interpretation ausgewählt. Es werden dabei die absolute und die relative Modellgüte unterschieden.

Das Ziel der LCA ist eine möglichst genaue Schätzung der unbekannten Modellparameter *Klassengrößen* und *klassenspezifischen Antwortmöglichkeiten* (vgl. L. M. Collins & Lanza, 2013; Geiser, 2011; Gollwitzer, 2012). Diese werden iterativ geschätzt, was bedeutet, dass die Parameter schrittweise dem

Ziel bestimmte Optimierungskriterien zu erfüllen, angepasst werden. In iterativen Ansätzen werden aufeinanderfolgende Sätze von Parameterschätzungen mit einem prinzipiellen Suchalgorithmus getestet. Parameter in der LCA werden oft mittels einer Version eines iterativen Verfahrens geschätzt, das als *Erwartungsmaximierungs* (EM)-Algorithmus bezeichnet wird (Dempster et al., 1977), manchmal in Kombination mit einem anderen iterativen Verfahren, dem sogenannten *Newton-Raphson*-Algorithmus (Agresti, 1990/2002).

Die Startwerte für die Parameterschätzungen sind dabei relativ weit vom Optimum entfernt und werden in weiteren Schritten, den *Iterationen*, so verändert, dass sie sich dem Optimum möglichst annähern beziehungsweise dieses nicht mehr verbessert werden kann. Bei der LCA wird das Verfahren zur Bestimmung des Optimums *Maximum Likelihood* (ML)-Verfahren genannt. Mit ihm wird die Wahrscheinlichkeit, *Likelihood*, als Produkt der unbedingten Antwortmusterwahrscheinlichkeiten über alle beobachteten Antwortmuster als grösstmöglicher Wert festgelegt. Wann immer ein iterativer Schätzalgorithmus verwendet wird, um ML-Parameterschätzungen zu erhalten, ist es notwendig, Kriterien für die Einstellung des Verfahrens festzulegen, andernfalls könnte der Algorithmus auf unbestimmte Zeit fortgesetzt werden. In der Regel werden zwei verschiedene Kriterien angegeben. Eine davon ist die maximale Anzahl der Iterationen, die das Verfahren durchführen darf. Das andere, wichtigere Kriterium ist eine Stoppregel, welche darauf basiert festzulegen, wann die Suche nahe genug an einer Reihe von Parameterschätzungen liegt, welche die *Likelihood*-Funktion maximiert oder zumindest fast maximiert (L. M. Collins & Lanza, 2013).

Das iterative Verfahren scheint einleuchtend, da es den Unterschied zwischen einem schlechten und einem optimalen Modell über die Höhe der Antwortmusterwahrscheinlichkeit festlegt. Mathematisch wirft dies aber einige Probleme auf. So beispielsweise, wenn mehrere Kombinationen von Modellparametern zum gleichen *Maximum Likelihood* führen, was besonders bei kleineren Stichproben auftreten kann. Aus diesem Grund wird statt des absoluten ein *standardisierter Likelihood* verwendet, in welchem die geschätzten unbedingten zu den empirisch beobachteten relativen Antwortmusterwahrscheinlichkeiten in Beziehung gesetzt werden, dem *Likelihood-Ratio-G²*-Test (Gollwitzer, 2012; Lanza et al., 2012).

Die zweite zu beachtende Komponente zur Bestimmung des absoluten *Modelfits* ist der *Pearson-²*-Test. Signifikante Werte, sowohl beim *Likelihood-Ratio-G²*- oder beim *Pearson-²*-Test weisen darauf hin, »dass es eine statistisch bedeutsame Abweichung zwischen beobachteten und modellimplizier-

ten Patternhäufigkeiten gibt. Das würde bedeuten, dass das Modell die Daten nicht perfekt abbildet« (Geiser, 2011, S. 260). Den p -Werten für die beiden Teststatistiken sollte aber nur getraut werden, wenn sie für dasselbe Modell nicht zu unterschiedlich ausfallen. Einen Hinweis darauf, ob das Modell trotz signifikanter Werte beibehalten werden kann, bieten dabei die Freiheitsgrade df (*degree of freedom*). Liegen diese nicht viel höher als die Werte der Teststatistiken, kann davon ausgegangen werden, dass die Fehlpassung, *Misfit*, nicht sehr gross ist.

Beide Tests folgen nur bei grossen Stichproben und relativ geringen Itemzahlen der theoretischen χ^2 -Verteilung (L. M. Collins et al., 1993). Für die Ermittlung der notwendigen Stichprobengrösse wird vorgeschlagen, dass diese mindestens so viele Personen umfassen sollte, wie es maximale Antwortmöglichkeiten pro Item gibt oder die Itemanzahl multipliziert mit den maximalen Antwortmöglichkeiten multipliziert mit fünf (Formann, 1984; Gollwitzer, 2012).

Durch die eben beschriebenen Probleme der asymptotischen, sprich der sich nicht annähernden Verteilung, bei der *Likelihood-Ratio- G^2* - und *Pearson- χ^2* -Statistik wird der Einfachheit halber oft der relative *Modelfit* zur Beurteilung der Klassenlösungen mittels dem *Bootstrap-Likelihood-Ratio- χ^2* (BLRT)-Differenztest und dem *Vuong-Lo-Mendell-Rubin* (VLMR)-Test herangezogen. Ergänzt durch die informationstheoretischen Masse AIC (*Akaike's Information Criterion*), BIC (*Bayesian Information Criterion*) und aBIC (*sample-size adjusted Bayesian Information Criterion*) lassen sich die geschätzten Modelle ebenfalls vergleichen (Geiser, 2011).

In einem *Bootstrap*-Verfahren werden bei zu kleiner Stichprobengrösse simulierte Daten selbst erzeugt. Efron (1979) verwendet für deren Beschreibung die Allegorie der Münchhausenschen Fähigkeit, sich am eigenen Zopf beziehungsweise an den eigenen Stiefelschlaufen (engl. *bootstraps*), aus dem Sumpf zu ziehen. So wird anhand einer parametrischen *Bootstrap*-Prozedur (McLachlan & Peel, 2000) der ungefähre korrekte p -Wert für den *Bootstrap-Likelihood-Ratio- χ^2* -Differenzwert aufgrund von *Monte-Carlo-Bootstrap*-Stichproben geschätzt (Geiser, 2011). In *Monte-Carlo* (MC)-Simulationen wird eine sehr grosse Anzahl von Zufallsdatensätzen generiert, welche durch das *Gesetz der grossen Zahlen* (Henze, 2013) eine Annäherung an einen bestimmten Wert erlauben.

Eine MC-Simulation kann alternativ durch die Untersuchung informationstheoretischer Masse beziehungsweise Informationskriterien ersetzt werden. Der Vorteil liegt darin, dass die Informationskriterien Modelle mit zu

vielen Parametern bestrafen. Anders gesagt passt ein Modell mit vielen Klassen logischerweise besser auf die Daten als ein sparsameres. Nach dem Parsimonitätsprinzip sind sparsamere Modelle, also solche mit weniger Klassen, zu belohnen, entsprechend wird ein (zu) komplexes Modell mithilfe eines Bestrafungsfaktors abgewertet (Gollwitzer, 2012). Das Parsimonitätsprinzip ist auf die philosophischen Arbeiten von Wilhelm von Ockham zurückzuführen. Seine Lehre besagt, dass in Aussagen unnötige Vervielfachungen vermieden werden sollen. Diese pragmatische Zweckmässigkeitsregel ist für die wissenschaftliche Beschreibung von Phänomenen vorgesehen und will nicht die Annahme vermitteln, dass die Welt möglichst sparsam aufgebaut sein sollte. Wird die Sparsamkeitsregel verletzt, heisst das nach Ockham nicht, dass die Aussage nicht wahr ist, sondern lediglich, dass sie für die epistemologische Bearbeitung einer Fragestellung nicht nötig, beziehungsweise nicht zielführend, ist (Beckmann, 2010).

Der *Vuong-Lo-Mendell-Rubin* (VLMR)-Test basiert auf einem ähnlichen Prinzip wie der *Bootstrap-Likelihood-Ratio* (BLRT)-Differenztest. Mit ihm werden ebenfalls Klassenlösungen gegeneinander getestet, wobei ein signifikanter Wert darauf hinweist, dass ein Modell mit mehr Klassen gegenüber jenem mit weniger zu bevorzugen ist. Nylund et al. (2007) konnten in einer Simulationsstudie zeigen, dass der BLRT-Differenztest die Anzahl Klassen genauer ausgibt als der VLMR-Test. Besonders bei grossen Stichproben wurden die Anzahl Klassen überschätzt.

Weiter werden drei ähnliche Informationskriterien unterschieden: das *Akaike's Information Criterion* (AIC), das *Bayesian Information Criterion* (BIC) sowie das *sample-size adjusted Bayesian Information Criterion* (aBIC). Bei allen drei ist darauf zu achten, dass je niedriger der Wert ist, desto besser passt das Modell auf die Daten (Gollwitzer, 2012). Diese informationstheoretischen Masse liefern deskriptive Indizes für die Beurteilung der Modellgüte. Alle drei berücksichtigen die Güte der Anpassung des Modells an die Daten sowie die Modellsparsamkeit. Nylund et al. (2007) empfehlen das BIC gefolgt vom aBIC als beste Indikatoren. Das AIC hat in ihrer Simulationsstudie ungenügende Werte geliefert. Auch Bacher und Vermunt (2010) argumentieren, dass das AIC dazu tendiert, Modellpassung und Klassenzahl zu überschätzen.

Geiser (2011) gibt weiter an, dass die mittleren Klassenzuordnungswahrscheinlichkeiten möglichst grösser als .80 sein sollten. Nach Nagin (2005) ist eine Wahrscheinlichkeit von .70 ausreichend.

Das Globalmass für die Zuverlässigkeit der Klassifikation wird als Entropie, *entropy*, bezeichnet (Celeux & Soromenho, 1996). Das in dieser Untersu-

chung verwendete statistische Analyseprogramm *Mplus* gibt lediglich eine relative Entropie aus, welche die Stichprobengrösse in die Berechnung einbezieht. Obwohl es keinen klaren Grenzwert für die Entropie gibt, bezeichnen J. Wang und Wang (2012) einen Wert von .80 als hoch, .60 als mittel und .40 als niedrig.

Durch eine *Residuenanalyse* können Ursachen für eine schlechte Modellanpassung analysiert werden. Sie gibt Auskunft darüber, welche beobachteten Antworten beziehungsweise welche Individuen mit abweichenden Antwortmustern (*Outliers*), zum *Misfit* beitragen.

Abschliessend muss jedes Modell beziehungsweise seine Klassen eindeutig identifizierbar sein. Geiser (2011, S. 270) schlägt dafür vor, dass »alle Items innerhalb jeder Klasse entweder hohe oder niedrige (nicht aber mittlere) Antwortwahrscheinlichkeiten aufweisen« sollen. Die interpretierbare Lösung ist in jedem Fall gegenüber der nichtinterpretierbaren vorzuziehen.

Identifikation bedeutet auch die inhaltliche Interpretierbarkeit der Klassen beziehungsweise der Klassenbezeichnungen. Grundsätzlich ist das sparsamste Modell zu verwenden. Ein Modell mit einer Klasse mehr und ähnlichen Gütekriterien kann aber eine differenziertere Betrachtungsweise erlauben. Modelle, welche lediglich die Klassenbezeichnungen von beispielsweise *tief*, *mittel* und *hoch* erlauben, sind zu verwerfen (L. M. Collins & Lanza, 2013; Lanza et al., 2012). Die Klassenbezeichnungen müssen in Bezugnahme auf die Fragestellungen und somit das Erkenntnisinteresse schlüssig sein.

7.7 Analyse latenter Klassen mit Kovariaten

Das LCA-Modell kann unter Einbezug von *Kovariaten*, auch *Kovariablen* oder *Zusatzvariablen* genannt, erweitert werden. Die Kovariaten können als mögliche Prädiktoren wie auch als Ergebnisse der latenten Klassenzugehörigkeit fungieren (Lanza et al., 2012).

Die Idee dazu wurde aus dem ursprünglich von Clogg (1981) beschriebenen *Latent Class Model* (LCM) mit externen Variablen übernommen beziehungsweise weiterentwickelt. Diese Methode liefert Informationen über die Zugehörigkeit zu einer bestimmten Gruppe, indem die strukturellen und prädiktiven Assoziationen zwischen der latenten Klassenvariablen und den Kovariaten untersucht werden (Lanza et al., 2013; Masyn, 2013; Nylund-Gibson et al., 2019).

In der vorliegenden Untersuchung sollen die möglichen prädiktiven Aspekte der Kovariaten Migration, Behinderung, sozioökonomischer Status und Geschlecht auf die Klassenzugehörigkeit Sozialkapital untersucht werden. Es wird daher aus den verschiedenen Methoden der LCA mit Kovariaten jene der *Latent Class Regression* (LCR) vorgestellt und von den anderen abgegrenzt.

Abbildung 8

Modell einer LCA mit Klasse c , Indikatoren u_1 bis u_6 und Kovariaten $x_1 \dots x_g$ (J. Wang & Wang, 2012, S. 309)

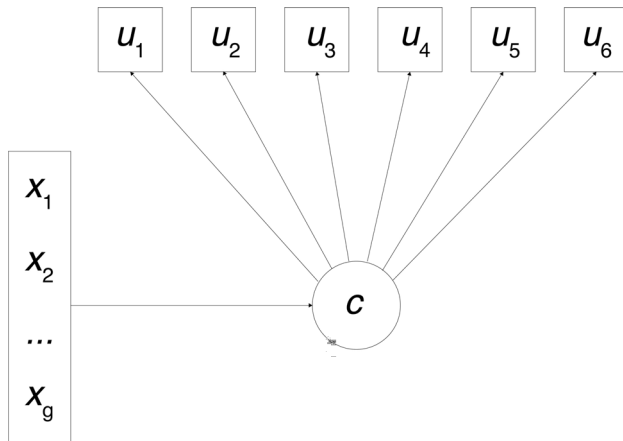


Abbildung 8 zeigt ein Modell, in welchem Kovariaten x_1 bis x_g als Prädiktoren der latenten Klasse funktionieren. Die Beziehung zwischen der latenten Klassenzugehörigkeit und den Kovariaten wird durch ein multinomiales Logit-Regressionsmodell modelliert und in folgender Formel ausgedrückt (Nylund-Gibson & Choi, 2018, S. 450):

Es wird dabei schrittweise nach klassenübergreifenden Effekten gesucht, indem klassenspezifische Mittelwert- und Varianzschätzungen für jede Kovariate geschätzt und dann paarweise Vergleiche durchgeführt werden, um festzustellen, wo zwischen den Klassen die Kovariaten signifikant unterschiedlich sind (Nylund-Gibson & Choi, 2018).

Klassifizierung und Vorhersage der Klassenzugehörigkeit werden gleichzeitig berechnet. Die Kovariaten müssen dabei immer Teil des jeweiligen Mo-

dells sein, da diese sonst falsch spezifiziert werden, was zu verzerrten Parameterschätzungen führt (B. O. Muthén, 2004).

$$\Pr(\mathbf{c}_i = \mathbf{k} | \mathbf{x}_i) = \frac{\exp(\mathbf{y}_{0k} + \mathbf{y}_{1k}\mathbf{x}_i)}{\sum_{j=1}^K \exp(\mathbf{y}_{0k} + \mathbf{y}_{1k}\mathbf{x}_i)}$$

(Formel 3)

Der Vorteil dieser Vorgehensweise liegen darin, dass auf Grund fehlender Werte bei den Indikatorvariablen oder bei den Kovariaten keine Fälle ausgeschlossen werden müssen. Weiter werden Unsicherheiten bei der latenten Klassenzugehörigkeit in den Berechnungen des Zusammenhangs der latenten Klassen und der Kovariate berücksichtigt. Ein Nachteil dieses Verfahrens liegt darin, dass die Kovariate durch die Einbezugnahme direkt Einfluss auf die Mittelwerte der Indikatoren sowie auf die Klassengrößen haben kann. Auch können die Regressionskoeffizienten nicht definiert werden, wenn Kovariaten keine Varianz in den latenten Klassen aufweisen.

Es gibt in diesem Sinne keine analytische Möglichkeit, zwischen latenten Klassenindikatoren und Kovariaten in der *Likelihood*-Funktion eines Mischmodells zu unterscheiden, das gleichzeitig sowohl das Messmodell für eine latente Klassenvariable als auch eine strukturelle prädiktive Beziehung beinhaltet, die die Mitgliedschaft in einer latenten Klasse mit einer Kovariate verbindet.

Eine Übersicht sowie Bewertungen zu den gegenwärtig bekannten Verfahren liefern Nylund-Gibson et al. (2019). Empfohlen ist das *One-step (distal-as-indicator)*-Verfahren, bei dessen Anwendung aber in Betracht gezogen werden muss, dass es die Kovariate direkt als Indikator der latenten Klassenvariable behandelt (Asparouhov & Muthén, 2014; Nylund-Gibson et al., 2014; Vermunt, 2010). Bei den Verfahren in zwei Schritten, handelt es sich entweder um historische Vorläufer der gegenwärtigen Herangehensweisen oder um solche, die noch nicht weit genug entwickelt sind, um zweckdienlich in statistische Programme implementiert zu werden. Des Weiteren gibt es drei Verfahren in drei Schritten, welche häufig zur Anwendung kommen:

- Maximum Likelihood (ML) three-step approach (Vermunt, 2010)
- Bolck-Croon-Hagenaars (BCH) method (Bolck et al., 2004)
- Lanza-Tan-Bray (LTB) approach (Lanza et al., 2013)

Alle drei 3-Schritt-Verfahren werden mit entsprechenden Vorbehalten bezüglich der automatisierten oder manuellen Durchführung, empfohlen. Die Entscheidung darüber, in welcher Form das Verfahren durchgeführt wird hängt fast ausschliesslich davon ab, wie Kovariaten implementiert werden (Nylund-Gibson et al., 2019, S. 7). Grundsätzlich wird im 3-Schritt-Verfahren der Klassenbildungsschritt von den subsequenten Modellen getrennt (Nylund-Gibson & Masyn, 2016). Im ersten Schritt wird die LCA nur mit den massgebenden Indikatoren modelliert, um im zweiten Schritt gemäss den bereits erwähnten Gütekriterien analytisch das am besten passende Modell auszuwählen (Nylund et al., 2007). Im dritten Schritt werden die Kovariaten in das Modell einbezogen. In allen drei genannten Ansätzen werden die Messparameter der latenten Klassen unter Berücksichtigung von Klassifizierungsfehlern fixiert, anschliessend werden die Kovariaten aufgenommen und deren Zusammenhang mit der latenten Klassenvariablen geschätzt (Nylund-Gibson & Choi, 2018).

Für die vorliegende Analyse ist es in erster Linie von Interesse, wie die Kovariaten die latente Klassenzugehörigkeit beeinflussen. J. Wang und Wang (2012, S. 375) empfehlen dafür folgende Vorgehensweise: »If our interest is only to see how covariate x affects the latent class membership (...), this can be done simply by specifying regressions of latent classes c_1 and c_2 on x in the following Mplus program.«

Formel 4 stellt das Modell mit Kovariaten mathematisch dar:

$$P(Y = y | X = x) = \sum_{c=1}^C \gamma_c(x) \prod_{j=1}^J \prod_{r_j=1}^{R_j} \rho_{j,r_j|c}^{I(y_j=r_j)}.$$

(Formel 4)

Es entspricht dem bereits vorgestellten Modell, ergänzt um die multinomiale logistische Regression mit welcher die Zusammenhänge zur Kovariate x geschätzt werden (Agresti, 1990/2002; L. M. Collins & Lanza, 2013, S. 153).

Damit das 1-Schritt-Modell erfolgreich konvergiert, müssen relativ anspruchsvolle Bedingungen erfüllt sein. In praktischen Anwendungen konnte gezeigt werden, dass dies der Fall ist, wenn die Daten aus idealen Bedingungen stammen, so zum Beispiel eine grosse Stichprobe, mittlere bis hohe Klassentrennungsgüte und kein Vorhandensein eines geringen Klassenanteils (Kamata et al., 2018).

Die Bezeichnung 1-Schritt-Modell bezieht sich auf die Modellierung mit den Kovariaten; damit das Modell spezifiziert werden kann, sind nach

Nylund-Gibson und Masyn (2016) aber mehrere Schritte notwendig. Zuerst werden Modelle ohne Kovariaten gerechnet. Mittels der genannten Gütekriterien wird über die am besten passende Anzahl Klassen entschieden. Anschliessend werden die Kovariaten in das Modell integriert. Die Auswahl und Reihenfolge der Implementation der Kovariaten sollte theoriebezogen sein. Obwohl aus dem ersten Schritt bereits bekannt ist, welche Anzahl Klassen am besten auf die Daten passt, wird empfohlen, Modelle mit noch je mindestens einer Klasse mehr beziehungsweise einer Klasse weniger zu rechnen. So geht man sicher, dass auch für das Modell mit Kovariaten die am besten passende Anzahl Klassen gefunden wurde.

Auch sollte beim Einbezug von Kovariaten in latente Klassenanalysen auf *Messinvarianz* getestet werden. Dies beinhaltet die Anwendung von Modell-Restriktionen so, dass jede Item-Antwort-Wahrscheinlichkeit für jede Ausprägung der Kovariaten (z.B. männlich vs. weiblich, behindert vs. nicht behindert usw.) die gleiche ist. Der Messinvarianz-Test kann aber sehr empfindlich sein, wenn viele Parameter beteiligt sind. Daher ist es möglich lediglich auf die Informationskriterien der relativen und absoluten Modelgüte zurückzugreifen. Bei Erweiterungen der latenten Klassenanalysen für mehrere Messzeitpunkte (*Latent Transition Analysis*) oder beim Einfügen von Gruppierungsvariablen sollte aber unbedingt auf Messinvarianz getestet werden (Lanza et al., 2012).

Um die Stärke des Zusammenhangs zwischen der Klassenzuteilung und den Kovariaten zu quantifizieren werden *Odds Ratios* (OR) ausgegeben. Sie beschreiben bei dichotomen Indikatoren die Eintretenswahrscheinlichkeit beim Vorhandensein beziehungsweise bei der Abwesenheit eines Merkmals.

$$\text{oddsratio} = \frac{\frac{P(\text{event1} | \mathbf{x}=1)}{P(\text{event2} | \mathbf{x}=1)}}{\frac{P(\text{event1} | \mathbf{x}=0)}{P(\text{event2} | \mathbf{x}=0)}}$$

(Formel 5)

Formel 5 bildet dieses Verhältnis mathematisch ab, $\mathbf{x}$ entspricht der Kovariante in ihren möglichen Ausprägungen. Als einfaches Beispiel kann Geschlecht genannt werden, wobei 0 weiblich und 1 männlich entsprechen könnte (vgl. L. M. Collins & Lanza, 2013, S. 156).

Mit der OR können deskriptive Aussagen bezüglich der Effektstärke und damit des Zusammenhangs oder der Unabhängigkeit zweier Variablen gemacht werden. Bei einer OR grösser als 1, kann davon ausgegangen werden, dass das Vorhandensein eines Merkmals (Kovariate) die Wahrscheinlichkeit

des Vorhandenseins des anderen Merkmals (Klassenzugehörigkeit) erhöht. Im umgekehrten Fall, also wenn die OR kleiner als 1 ist, senkt ein Vorhandensein des einen Merkmals (Kovariate) die Wahrscheinlichkeit des Vorhandenseins des anderen Merkmals (Klassenzugehörigkeit). Bei einer OR von 1 sind die Variablen unabhängig. Die Ausprägungen der OR verlaufen auf einem Kontinuum von 0 bis unendlich.

ORs in einer LCA werden in Bezug auf eine Referenzklasse interpretiert, das bedeutet der Effekt einer Kovariate bildet die Veränderung der Eintretenswahrscheinlichkeit in Bezug auf die Referenzklasse ab. Von ORs kann nicht auf Kausalität geschlossen werden, sprich es muss immer davon ausgegangen werden, dass weitere Einflussfaktoren existieren, welche das Ergebnis konfundieren, daher wird *Effekt* in statistischem und nicht kausalen Sinn verwendet (L. M. Collins & Lanza, 2013, 155ff.).

7.8 Exploratorische Faktoranalyse

Mit Faktorenanalyse wird eine Gruppe von multivariaten Analyseverfahren bezeichnet, welche unter anderem zum Ziel haben Daten zu reduzieren, indem eine Vielzahl von Variablen auf eine oder mehrere gemeinsame Dimensionen zurückgeführt wird und um gegebenenfalls die Validität von Konstrukten zu überprüfen. Die exploratorische Faktorenanalyse, *Exploratory Factor Analysis* (EFA), kommt dann zur Anwendung, wenn keine konkreten Zuordnungshypothesen vorliegen und in erster Linie die Anzahl der Faktoren ermittelt werden soll (Moosbrugger & Schermelleh-Engel, 2012, S. 326).

Bei der Anwendung der EFA müssen vorgängig verschiedene Durchführungsentscheidungen, die Auswirkungen auf die Ergebnisse haben, getroffen werden. In Abhängigkeit zum Erkenntnisinteresse werden die Art der Extraktionsmethode, die Wahl des Abbruchkriteriums und die Methode der Faktorenrotation festgelegt.

Als Extraktionsmethoden stehen sich die Hauptkomponentenanalyse, *Principal Component Analysis* (PCA), und die Hauptachsenanalyse, *Principal Axes Factor Analysis* (PFA), gegenüber. Bei der PCA wird implizit angenommen, dass die gesamte Varianz der Indikatoren durch eine gemeinsame Hauptkomponente erklärt werden kann. In der Realität sind die beobachteten Variablen aber selten messfehlerfrei. Da aber bei der EFA eine Variablenbeziehungsweise Dimensionsreduktion angestrebt wird, können empirisch weniger Hauptkomponenten festgelegt werden, als theoretisch extrahiert

wurden, und es wird vereinfacht eine möglichst hohe Varianzaufklärung angestrebt. Somit ermöglicht die PCA eine hohe Varianzaufklärung zwischen den beobachteten Indikatoren und bestimmt sogenannte Hauptkomponenten, die Faktoren. Die PFA, hingegen, geht davon aus, dass die beobachteten Variablen sowohl wahre Varianz wie auch Messfehlervarianz aufweisen. Sie kommt dann zur Anwendung, wenn eine einfache Datenreduktion vorgenommen werden soll. Ziel der PFA ist es, latente Konstrukte, also Faktoren, zu identifizieren, die die Korrelationen zwischen den Indikatoren aufzeigen, ohne die gesamte Varianz zu erklären. Erklärt wird also nur die beobachtbare Varianz (Moosbrugger & Schermelleh-Engel, 2012, 327f.).

Um die Anzahl der Faktoren zu bestimmen, werden Abbruchkriterien anhand des *Kaiser-Kriteriums*, des *Scree-Tests* und der *Parallelanalyse* festgelegt. Beim Kaiser-Kriterium werden nur jene Faktoren berücksichtigt, die einen Eigenwert höher als 1 aufweisen. Das bedeutet, dass der Faktor mehr Varianz aufklärt als der einzelne standardisierte Indikator. In der Realität führt dies aber oft zu einer Überschätzung der Anzahl Faktoren, daher wird als weiteres Abbruchkriterium der Scree-Test zu Hilfe genommen. Der Eigenwertverlauf wird als Grafik dargestellt, in der die Faktoren ordinal der Grösse nach geordnet sind. In der Regel bildet der Eigenwertverlauf an einer bestimmten Stelle einen Knick, alle inhaltlich relevanten Faktoren befinden sich vor diesem Knick. Abschliessend kann eine Parallelanalyse mit mindestens hundert Datensätzen mit zufälligen Zahlen Aufschluss über die Anzahl relevanter Faktoren geben (Moosbrugger & Schermelleh-Engel, 2012, 330f.).

Eine Faktorenrotation erlaubt es, die Faktorenextraktion, welche lediglich sukzessive maximale Eigenwerte sucht, im Faktorenraum zu drehen. Ziel ist es, eine sogenannte Einfachstruktur zu erreichen, bei welcher jede Variable nur auf einen einzigen Faktor eine hohe Ladung aufweist (Primärladung). Das gängigste Verfahren ist die *Varimax-Rotation* (Varianzmaximierung). Mittels orthogonaler Rotation, welche die Unkorreliertheit der Faktoren beibehält, können die Faktoren unabhängig voneinander interpretiert werden. Bei der Anwendung werden die Faktoren in fortlaufenden Schritten, den Iterationen, so lange im Raum gedreht, bis die Varianz der quadrierten Ladungen pro Faktor maximal ist. Dies führt im Idealfall zu hohen Primärladungen. Bei den obliquen, den schiefwinkligen, Rotationsverfahren wird die Unkorreliertheit der Variablen aufgegeben. Die verbreitete *Oblimin-Rotation* strebt die simultane Optimierung eines orthogonalen und eines obliquen Rotationskriteriums an. Faktorenlösungen, die damit erstellt werden, sind sowohl nahe an der Realität wie auch einfach interpretierbar. Moosbrugger und Schermelleh-

Engel (2012) geben abschliessend zu den Rotationsverfahren folgende Auswahlempfehlung:

Erfolgt eine Faktorenanalyse primär mit dem Ziel der Datenreduktion und ohne theoretisch fundierte Annahmen über die Dimensionalität der untersuchten Variablen, ist immer ein orthogonales Rotationsverfahren empfehlenswert. Liegen dagegen theoretische Anhaltspunkte vor, die auf korrelierte Faktoren hinweisen, so ist der Einsatz eines obliquen Rotationsverfahrens zweckmässig. (Moosbrugger & Schermelleh-Engel, 2012, S. 332)

In der vorliegenden Untersuchung wird die EFA dazu eingesetzt, die vorhandenen Indikatoren von Sozialkapital auf zugrunde liegende Faktoren zu untersuchen. Diese bilden bestenfalls Dimensionen von Sozialkapital ab, welche die Anschaulichkeit in der Interpretation der Ergebnisse erhöhen.

7.9 Statistische Analyseprogramme

Die Daten für die vorliegende Untersuchung werden als SPSS-Dateien geliefert. Für die Aufbereitung wird das entsprechende Statistikprogramm *SPSS* in der Version 25 (IBM Corp., 2017) verwendet. Die EFA wird ebenfalls mit *SPSS* durchgeführt. Für die weiterführenden Analysen der LCA und LCR wird das Statistikprogramm *Mplus* in der Version 8 (B. O. Muthén & Muthén, 2017) verwendet. *Mplus* ist leistungsstark und flexibel. Das Modellierungssystem basiert auf dem verbindenden Thema der latenten Variablen und seiner einzigartigen Verwendung von sowohl kontinuierlichen als auch kategorialen latenten Variablen. Das Statistikprogramm wird von einem ausführlichen Handbuch, dem *Mplus User Guide*, begleitet (L. K. Muthén & Muthén, 2017), das die einzelnen Modelle erklärt und die anzuwendende Syntax beschreibt.