

4. Auswertung: Q-Methode, Ergebnisse und Interpretation

Mit dem im dritten Kapitel ausgearbeiteten Forschungsdesign konnten Probanden verschiedener Bevölkerungsgruppen Sprachspiele so spielen, dass für einen forschenden Blick Daten zurückgeblieben waren, die sich auf Strukturmuster hin auswerten ließen. Aus welchen Erhebungs- und Formatierungsprozessen diese Daten hervorgingen, wird im ersten Abschnitt thematisiert (4.1). Wie man mit der Q-Methode aus derartigen Daten Erkenntnisse gewinnen kann, sei im Anschluss erläutert (4.2). Die Ergebnisse werden in den darauffolgenden Abschnitten ausgestellt, um an ihnen dann einen interpretierenden Vergleich vorzuführen (4.3).

4.1 Die Daten

Sind alle Sprachspiele gespielt, stehen uns die Ergebnisse als ein Haufen von je persönlich ausgefüllten 20×9 -Tabellen vor Augen. Für jede der 20 zu bewertenden Modellaussagen haben die Befragten einen Wert aus dem Bereich von -4 bis 4 angegeben,¹ der anzeigen soll, inwiefern man sich auf die jeweilige Aussage in der beschriebenen Situation bezieht. Tabelle 4.1 veranschaulicht, wie man sich einen ausgefüllten Fragebogen vorstellen kann; freilich sahen die Daten nicht tatsächlich so aus. In einem geeigneten Tabellenformat (.csv) wurden sie aus der Server-Datenbank exportiert, um sie anschließend mit der Statistik-Software R weiterzuverarbeiten. Außerdem erschienen die Aussagen während der Befragung je einzeln und in zufälliger Reihenfolge, um mögliche Tendenzen auszuschließen. In Tabelle 4.1 finden sich die im Vollzug zufällig präsentierten Fragen bereits gemäß

1 Dass die Wahl hier auf eine neunwertige Likert-Skala fiel, hatte schlicht den Grund, *a priori* möglichst viel Varianz offenzuhalten. Unter anderem deshalb habe ich auch auf einen sogenannten *Q-Sort*, d. h. auf eine Begrenzung der jeweiligen Bewertungsslots, verzichtet, vor allem aber auch, weil sich derartige Begrenzungen hier inhaltlich kaum rechtfertigen ließen. Denn warum sollte man sich, zumindest prinzipiell, nicht auf alle Aussagen mit unterschiedloser Zustimmung oder Ablehnung beziehen können?

ihren Modellfamilien gruppiert. Der Übersicht wegen stellen wir die Ergebnisse durchgängig in dieser Anordnung dar; die tatsächlichen, randomisierten, Bewertungsreihenfolgen bleiben außen vor.

Tabelle 4.1: Symbolische Darstellung eines ausgefüllten Online-Fragebogens

Bitte geben Sie für jede der Aussagen an, wie Sie sich in Ihrer Rede auf sie beziehen werden.									
		-3	-2	-1					
KL1	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	✗	2	3	4 (»stark zustimmend«)
KL2	-4 (»stark ablehnend«)	-3	✗	-1	o (»überhaupt nicht«)	1	2	3	4 (»stark zustimmend«)
KL3	-4 (»stark ablehnend«)	-3	-2	✗	o (»überhaupt nicht«)	1	2	3	4 (»stark zustimmend«)
KL4	-4 (»stark ablehnend«)	✗	-2	-1	o (»überhaupt nicht«)	1	2	3	4 (»stark zustimmend«)
KL5	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	1	✗	3	4 (»stark zustimmend«)
KY1	-4 (»stark ablehnend«)	-3	-2	-1	✗ (»überhaupt nicht«)	1	2	3	4 (»stark zustimmend«)
KY2	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	1	2	✗	4 (»stark zustimmend«)
KY3	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	✗	2	3	4 (»stark zustimmend«)
KY4	-4 (»stark ablehnend«)	-3	✗	-1	o (»überhaupt nicht«)	1	2	3	4 (»stark zustimmend«)
KY5	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	1	2	✗	4 (»stark zustimmend«)
MG1	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	1	2	✗	4 (»stark zustimmend«)
MG2	✗	-3	-2	-1	o (»überhaupt nicht«)	1	2	3	4 (»stark zustimmend«)
MG3	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	1	✗	3	4 (»stark zustimmend«)
MG4	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	1	2	✗	4 (»stark zustimmend«)
MG5	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	✗	2	3	4 (»stark zustimmend«)
AG1	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	1	2	✗	4 (»stark zustimmend«)
AG2	✗	-3	-2	-1	o (»überhaupt nicht«)	1	2	3	4 (»stark zustimmend«)
AG3	-4 (»stark ablehnend«)	✗	-2	-1	o (»überhaupt nicht«)	1	2	3	4 (»stark zustimmend«)
AG4	-4 (»stark ablehnend«)	-3	-2	✗	o (»überhaupt nicht«)	1	2	3	4 (»stark zustimmend«)
AG5	-4 (»stark ablehnend«)	-3	-2	-1	o (»überhaupt nicht«)	1	2	3	✗

Die Subjekttypen der ökonomischen Modellwelten haben die Rekrutierung der Probanden angeleitet. Gesucht wurden je 15 Arbeitnehmer, Arbeitgeber und Arbeitslose, die sich über ihre Auswahl im Bereich der persönlichen Angaben selbstkategorisiert haben. Die Teilnehmer stammten aus Panels eines Dienstleisters, welcher sie an den von mir gehosteten Online-Fragebogen vermittelte. Um Verzerrungen durch sozioökonomische Heterogenität nach Möglichkeit zu minimieren, war die Zielgruppe auf männliche Personen im Alter von über 40 Jahren mit Wohnsitz in den neuen Bundesländern beschränkt.

Insgesamt haben 94 Personen vollständig an der Umfrage teilgenommen: davon 55 Arbeitnehmer, 16 Arbeitgeber und 15 Arbeitslose.² Der Befragungszeitraum belief sich auf gut sechs Wochen. Der erste vollständige Ergebnisbogen ging am

² Acht Personen gaben an, dass sie berufstätig – und damit nicht »arbeitslos« – sind, ordneten sich aber auch weder als »Arbeitnehmer/-in« noch als »Arbeitgeber/-in« ein.

08.10.2021 ein, der letzte am 22.11.2021. Da pro Gruppe nur 15 Datensätze in die Auswertung kamen, war hier eine Selektion nötig. Die Arbeitslosen entsprachen genau der benötigten Anzahl, sodass hier keine Auswahl erforderlich war. Vor allem bei der Gruppe der Arbeitnehmer und in einem Fall auch bei jener der Arbeitgeber war es jedoch nötig, Datensätze auszuschließen. Dabei wurden nach einem möglichst einfachen Exklusionsverfahren jene 15 Datensätze bestimmt, welche die zugrundeliegende Struktur möglichst deutlich aufwiesen. Das Vorgehen ähnelt etwa dem Trimmen eines Gebüschs: Man entfernt die äußeren losen Äste, um die Gestalt des Dickichts um den Stamm besser zu erkennen. Diese Gestalt manipuliert man nicht, indem man sich lediglich schärfere Sicht darauf verschafft. Wir werden das Ausschlussverfahren gleich graphisch verfolgen und weiter unten technisch näher erläutern, wenn es daran geht, die Korrelationsmatrix zu analysieren.

4.2 Die Methode

Man könnte den Ausgangspunkt der Auswertungsmethode nun schlicht bei den 45 Datensätzen mit je 20 bewerteten Aussagen ansetzen. Damit übersähe man jedoch wohl Einsichten, die wir uns im dritten Kapitel erarbeitet haben. Was wir analysieren, sind nicht allein die 45 *realisierten* Datensätze, sondern deren Differenz zu allen anderen *möglichen*, die kontingent fehlen.

Wir können uns zu Anschauungszwecken vorstellen, dass wir die Ergebnisse wie in Tabelle 4.1 auf eine Klarsichtfolie drucken und die Auswahlkreuzchen Zeile für Zeile durch eine Linie verbinden. Die Antwortlinie für den Beispielbogen aus Tabelle 4.1 ist in Abbildung 4.1 dargestellt. Wir haben den Bogen lediglich um 90 Grad gedreht, unnötigen Text entfernt und Achsen angefügt.

Nun könnte man viele weitere solcher Linien in die Abbildung malen, indem wir die Bewertungen für je eine Person³ über die Aussagen hinweg verbinden. Man könnte die Linienverläufe auch 20 Mal mit einem neunseitigen Würfel auswürfeln. Aus diesem Bild lässt sich ablesen, dass die Zahl der möglichen Linien pro Person 9^{20} beträgt. Die realisierte Beispiellinie lässt sich demnach als kontingente Selektion aus dem Raum aller möglichen solcher Linien begreifen. Dies zeigt sich vielleicht deutlicher, wenn man einen Ergebnisbogen – wie in Abbildung 4.2 zu sehen – gemäß den im vorhergehenden Kapitel vereinbarten Regeln in Spencer-Brown-Notation darstellt. Die Form eines Ergebnisbogens liegt also in der kontingenten Unterscheidung zwischen einer *realisierten* Datenlinie und allen übrigen $9^{20} - 1$ *kontingent-fehlenden*.

3 Darin liegt auch der ausschließliche Sinn einer solchen Linie. Man kann damit Punkte identifizieren, die zu einer Person gehören. Weil eine Linie eben *eine* Linie darstellt, lässt sich diese Einheitsform auf *eine* Person beziehen.