

EMANUEL V. TOWFIGH  
NIELS PETERSEN

---

# Ökonomische Methoden im Recht

2. Auflage



MOHR SIEBECK

MOHR LEHRBUCH

*Emanuel V. Towfigh / Niels Petersen*

Ökonomische Methoden im Recht





Emanuel V. Towfigh / Niels Petersen

# Ökonomische Methoden im Recht

Eine Einführung für Juristen

mit Beiträgen von

Markus Englerth, Sebastian J. Goerg, Stefan Magen,  
Alexander Morell und Klaus Ulrich Schmolke

2., überarbeitete und aktualisierte Auflage

Mohr Siebeck

e-ISBN PDF 978-3-16-155193-2

ISBN 978-3-16-155192-5

Die Deutsche Nationalbibliothek verzeichnet diese Publikation in der Deutschen Nationalbibliographie; detaillierte bibliographische Daten sind im Internet über <http://dnb.dnb.de> abrufbar.

1. Auflage 2010

© 2017 Mohr Siebeck Tübingen. [www.mohrsiebeck.com](http://www.mohrsiebeck.com)

Dieses Werk ist lizenziert unter der Lizenz „Creative Commons Namensnennung 4.0 International“ (CC BY 4.0). Eine vollständige Version des Lizenztextes findet sich unter: <https://creativecommons.org/licenses/by/4.0/deed.de>.

Jede Verwendung, die nicht von der oben genannten Lizenz umfasst ist, ist ohne Zustimmung des Verlags unzulässig und strafbar.

Das Buch wurde von Gulde-Druck in Tübingen aus der Sabon gesetzt, auf alterungsbeständiges Werkdruckpapier gedruckt und gebunden.

## Vorwort zur zweiten Auflage

Das Interesse an der Integration ökonomischer und verhaltenswissenschaftlicher Methoden in die Rechtswissenschaft ist ungebrochen. Neben der Frage nach Auslegung, Systematisierung und Dogmatisierung des Rechts erfreut sich die Frage nach den Wirkmechanismen des Rechts auch im deutschsprachigen Rechtsraum zunehmender Aufmerksamkeit. Das hat sich im Erfolg des vorliegenden Werkes bemerkbar gemacht, dessen erfreuliche Rezeption nun eine zweite Auflage erfordert. Auch international ist die Methodenorientierung des Lehrbuchs als „Türöffner“ zu sozialwissenschaftlichen Ansätzen für Juristen wahrgenommen worden, was zwischenzeitlich zur Veröffentlichung einer englischen Fassung (*Economic Methods for Lawyers*, 2015) geführt hat.

Die eigenen Erfahrungen aus dem Einsatz des Lehrbuches in der universitären Lehre, die Rückmeldungen von Lehrenden und Studierenden und nicht zuletzt die Arbeit an der englischen Fassung haben uns wichtige Anstöße für die Überarbeitung der Neuauflage gegeben. Hervorzuheben sind die beiden völlig neu gefassten § 3 (Nachfrage, Angebot und Märkte – in dieser Auflage verfasst von *Alexander Morell*) und § 5 (Vertragstheorie und ökonomische Analyse des Vertragsrechts, jetzt verfasst von *Klaus Ulrich Schmolke*). Auch alle anderen Kapitel wurden inhaltlich aktualisiert und überarbeitet. Außerdem haben wir zur erleichterten Arbeit mit dem Buch und für klarere Verweisungen Randziffern eingefügt.

Für die wie immer kundige und sorgfältige Redaktion des Buches danken wir herzlich *Rebekka P. Herberg*, wissenschaftliche Mitarbeiterin an der EBS Universität in Wiesbaden, unterstützt von Frau *Melanie Epe* und Herrn *Luis Kleine Wortmann*; um die Überarbeitung von Glossar sowie Sachwort- und Personenverzeichnis und Korrekturen im Manuskript haben sich Frau *Henrike Boll*, Herr *Marius Kühne* und Herr *Malte Steuber* an der Westfälischen Wilhelms-Universität in Münster verdient gemacht, denen ebenso herzlicher Dank gebührt. Für die kompetente und angenehme Betreuung im Verlag danken wir einmal mehr Herrn *Franz-Peter Gillig* verbindlichst, ebenso wie Herrn *Matthias Spitzner* in der Herstellung.

Wir hoffen auf eine weiterhin wohlwollende Rezeption des Werks und bleiben dankbar für Anregungen zur Verbesserung oder Ergänzung!

Wiesbaden und Münster, im Dezember 2016

*Emanuel V. Towfigh und Niels Petersen*

## Vorwort zur ersten Auflage

Die ökonomische Methode hat in der Rechtswissenschaft in den letzten Jahrzehnten stetig steigende Aufmerksamkeit erfahren. Argumente ökonomischer Provenienz haben zunächst vor allem ins Zivilrecht, seit geraumer Zeit aber auch in anderen Rechtsgebieten Eingang gefunden. Es nimmt heute durchaus niemanden mehr wunder, wenn in einem juristischen Text von „Anreizen“ oder „Akteuren“ die Rede ist. Kaum eine Abhandlung über deliktische Haftung kommt ohne Überlegungen zur günstigsten Versicherbarkeit aus. Im Emissionshandelsrecht räsonieren Europarechtler über die Erstallokation von Zertifikaten. Und Strafrechtler diskutieren darüber, ob es nicht unter Abschreckungsgesichtspunkten sehr viel effektiver und letztlich auch effizienter wäre, daran zu arbeiten, die Entdeckungswahrscheinlichkeit gewisser Straftaten zu erhöhen, statt das Strafmaß weiter anzuheben. Weitere prominente Beispiele lassen sich beliebig für Rechtsgebiete vom Arzthaftungsrecht über das Immaterialgüterrecht, das Steuer- und Umweltrecht bis hin zum Wettbewerbs- und Kartellrecht finden. In der Rechtsvergleichung wird die ökonomische Theorie gern als *tertium comparationis*, als Vergleichsmaßstab, bemüht. Nachdem sich die Rechtsökonomie zunächst vornehmlich mit theoretischen Modellen beschäftigt hat, finden neuerdings auch vermehrt empirische Erkenntnisse Eingang in die Rechtswissenschaft.

Auch von „außen“ – etwa aus der Politik oder aus den Nachbarwissenschaften – wird zunehmend von Rechtswissenschaftlern gefordert, das geltende oder zu setzende Recht vor dem Hintergrund der Erkenntnisse über menschliches Verhalten zu rechtfertigen. Ist eine ins Auge gefasste rechtliche Maßnahme wirklich *geeignet*, ihr Regelungsziel zu erfüllen? Juristen werden so zunehmend gezwungen, sich der Grundlagen ihres eigenen Faches zu vergewissern. Um ihre gesellschaftsprägenden Einflussmöglichkeiten nicht zu verlieren, müssen sie zunehmend zu Experten für Verhaltenssteuerung durch Recht werden. Die Rechtsökonomik bietet, vor allem mit ihren verhaltenswissenschaftlichen Fortentwicklungen, hierfür einen geeigneten Rahmen.



Mit diesen Entwicklungen geht ein wachsender Bedarf an der Vermittlung von Kenntnissen sozialwissenschaftlicher Methodik im Allgemeinen und der Ökonomik im Besonderen einher. Wie findet man einen Zugang zu diesem Satz von Argumenten? Welche Einschränkungen sind zu beachten, wenn man ein ökonomisches Argument in den juristischen Diskurs einführt? Woran erkennt man ein gutes, ökonomisch fundiertes Argument? Wie entlarvt man ein schlechtes? Schließlich: Wie kann man selbst ein gutes Argument führen? Zur Beantwortung dieser Fragen möchte das vorliegende Lehrbuch einen Beitrag leisten. Es richtet sich an den juristischen Leser, der ohne jegliche sozialwissenschaftliche Vorkenntnisse eine erste Begegnung mit ökonomischen Methoden sucht und dabei auch Reiz und Stärke eines ökonomischen Arguments in ausgewählten rechtswissenschaftlichen Kontexten verstehen möchte. Die verschiedenen großen Bereiche der für das Recht relevanten ökonomischen Theorie – „Law & Economics“ – werden dabei in aller Kürze ebenso dargestellt wie neuere, stärker verhaltenswissenschaftlich orientierte Theorieansätze oder die Grundlagen der sozialwissenschaftlichen Empirie.

Das vorliegende Lehrbuch unterscheidet sich damit in seinem Zugang von konventionellen rechtsökonomischen Lehrbüchern. Es geht nicht darum, bestimmte Rechtsgebiete im Lichte ökonomischer Erkenntnisse neu zu betrachten. Es geht in erster Linie um die Vermittlung von Methode, nicht bestimmter inhaltlicher Theorien. Es zeigt nicht auf, wie bestimmte ökonomische Erkenntnisse im juristischen Kontext zu verstehen sind. Es soll vielmehr Hilfestellung geben, ökonomische Forschung selbst besser zu verstehen und auf juristische Fragestellungen anzuwenden. Ganz ohne inhaltliche Erkenntnisse kommt das Lehrbuch dabei natürlich nicht aus, so dass kurze Einführungen in einige grundlegende theoretische Konzepte der Ökonomie – von der Mikroökonomie über die öffentlichen Güter bis hin zu *Public Choice* – gegeben werden. Trotz dieser Schwerpunktsetzung haben sich die Autoren bemüht, die Bedeutung ihrer Ausführungen für das Recht anhand von Beispielen aus den verschiedenen Rechtsgebieten darzulegen. Der Fokus liegt dabei nicht – wie traditionell – allein auf dem Gebiet des Zivilrechts. Vielmehr werden Beispiele aus allen drei großen Rechtsgebieten, dem Zivil-, dem Straf- und dem öffentlichen Recht, angeführt.

Notabene: In diesem Lehrbuch werden die grundlegenden Modelle der Ökonomik präsentiert, weil es darum geht, den Juristen ökonomische Methoden näherzubringen. Wie in der Jurisprudenz herrscht auch in der Ökonomie über viele der hier als nicht weiter in Zweifel gezogen präsentierten Grundannahmen und Schlussfolgerungen bisweilen leidenschaftlicher

Streit. Zu jedem in diesem Band dargestellten Thema gibt es unzählige theoretische und empirische Variationen und Verfeinerungen, so zahlreich, dass es unmöglich ist, auch nur auf alle zu verweisen. Bei näherem Interesse sei dem geneigten Leser empfohlen, sich speziellerer Literatur zuzuwenden, die in aller Regel präzisere Modelle entwickelt hat. Entsprechende weiterführende Literaturhinweise sind am Anfang eines jeden Abschnitts angegeben. Zur Vertiefung von Spezialfragen sind auch Nachweise in den Fußnoten angegeben.

Die Autoren dieses Lehrbuchs verbindet eine Tätigkeit am Max-Planck-Institut zur Erforschung von Gemeinschaftsgütern. Hier forschen Juristen, Ökonomen und Psychologen interdisziplinär mit den verschiedensten Ansätzen aus dem verhaltenswissenschaftlichen Methodenkasten. Die jeweiligen Bearbeiter der einzelnen Abschnitte sind in der von ihnen dargestellten Materie wissenschaftlich ausgewiesen. Es war aber der Ehrgeiz der Verfasser, keinen Sammelband zur ökonomischen Methode herauszugeben, sondern ein in sich geschlossenes Lehrbuch. Die Konzeption dazu und die Vereinheitlichungsleistung am Ende haben die beiden Hauptherausgeber erbracht. Sie haben alle Beiträge sprachlich und strukturell überarbeitet, um Überschneidungen zu vermeiden, einen einheitlichen Stil sicherzustellen und Kohärenz zu gewährleisten. Dennoch wäre die Erstellung dieses Lehrbuchs nicht ohne die engagierte Hilfe einiger Institutsmitarbeiter möglich gewesen, die uns inhaltliche Anregungen gegeben und den Text am Ende Korrektur gelesen haben. Dank gebührt insbesondere *Konstantin Chatziathanasiou* und *Kristina Schönfeldt*. Wir hoffen nun auf eine wohlwollende Rezeption und sind für Anregungen zur Verbesserung oder Ergänzung des Werks dankbar.

Bonn, im Mai 2010

*Emanuel Towfigh* und *Niels Petersen*



## Inhaltsübersicht

|                                                                                        | Seite         | Rz. |
|----------------------------------------------------------------------------------------|---------------|-----|
| Vorwort zur zweiten Auflage . . . . .                                                  | V             |     |
| Vorwort zur ersten Auflage . . . . .                                                   | VII           |     |
| Verzeichnis der Abbildungen . . . . .                                                  | XXI           |     |
| Verzeichnis der Tabellen . . . . .                                                     | XXIII         |     |
| <br><b>§ 1 – Ökonomik in der Rechtswissenschaft . . . . .</b>                          | <br><b>1</b>  |     |
| <i>Niels Petersen / Emanuel V. Towfigh</i>                                             |               |     |
| I. Entwicklung der Rechtsökonomik . . . . .                                            | 2             | 1   |
| II. Normative und positive ökonomische Theorie . . . . .                               | 3             | 5   |
| III. Das Wesen sozialwissenschaftlicher Theorien . . . . .                             | 6             | 9   |
| IV. Sozialwissenschaftliche Theorie und rechtswissen-<br>schaftliche Methode . . . . . | 8             | 14  |
| V. Die relevanten Methoden der Ökonomie . . . . .                                      | 23            | 57  |
| <br><b>§ 2 – Das ökonomische Paradigma . . . . .</b>                                   | <br><b>25</b> |     |
| <i>Emanuel V. Towfigh</i>                                                              |               |     |
| I. Theoretische Grundannahmen . . . . .                                                | 25            | 61  |
| II. Wohlfahrtsanalyse und Effizienz . . . . .                                          | 39            | 87  |
| <br><b>§ 3 – Nachfrage, Angebot und Märkte . . . . .</b>                               | <br><b>45</b> |     |
| <i>Alexander Morell</i>                                                                |               |     |
| I. Einleitung . . . . .                                                                | 45            | 96  |
| II. Nachfrage . . . . .                                                                | 46            | 97  |
| III. Angebot . . . . .                                                                 | 62            | 131 |
| IV. Der Markt . . . . .                                                                | 68            | 147 |
| V. Marktversagen . . . . .                                                             | 72            | 153 |

|                                                                                              | Seite   | Rz. |
|----------------------------------------------------------------------------------------------|---------|-----|
| <b>§ 4 – Spieltheorie</b> . . . . .                                                          | 83      |     |
| <i>Stefan Magen</i>                                                                          |         |     |
| I. Spieltheorie und Recht . . . . .                                                          | 83      | 170 |
| II. Spiele in Normalform . . . . .                                                           | 86      | 177 |
| III. Typen von Spielen . . . . .                                                             | 99      | 203 |
| IV. Spiele in Extensivform . . . . .                                                         | 117     | 235 |
| V. Recht und informale Institutionen . . . . .                                               | 125     | 252 |
| <br><b>§ 5 – Vertragstheorie und ökonomische Analyse<br/>des Vertragsrechts</b> . . . . .    | <br>131 |     |
| <i>Klaus Ulrich Schmolke</i>                                                                 |         |     |
| I. Warum Verträge? . . . . .                                                                 | 132     | 260 |
| II. Marktstörungen als Begründung für Vertragsrecht . .                                      | 136     | 269 |
| III. Unvollständige Information – Problem und Lösungen .                                     | 138     | 273 |
| IV. Kognitive Beschränkungen und nichtrationales<br>Verhalten . . . . .                      | 148     | 295 |
| V. Anreizprobleme und unvollständige Information<br>nach Vertragsschluss . . . . .           | 151     | 302 |
| VI. „Verteilungsgerechtigkeit“ durch Vertragsrecht? . . . .                                  | 160     | 323 |
| <br><b>§ 6 – Public und Social Choice Theorie</b> . . . . .                                  | <br>163 |     |
| <i>Emanuel V. Towfigh / Niels Petersen</i>                                                   |         |     |
| I. Ökonomik und Staatswissenschaft . . . . .                                                 | 163     | 328 |
| II. Grundlegende Annahmen der <i>Public Choice Theory</i> . .                                | 165     | 332 |
| III. Fehlanreize in repräsentativen Systemen . . . . .                                       | 170     | 341 |
| IV. Kollektiventscheidungen durch Wahlen und<br>Abstimmungen: <i>Social Choice</i> . . . . . | 183     | 370 |
| <br><b>§ 7 – Empirische Methoden</b> . . . . .                                               | <br>195 |     |
| <i>Sebastian Goerg / Niels Petersen</i>                                                      |         |     |
| I. Grundlagen und Forschungsdesign . . . . .                                                 | 195     | 394 |
| II. Deskriptive Statistik . . . . .                                                          | 206     | 419 |
| III. Statistische Testverfahren . . . . .                                                    | 215     | 441 |
| <br><b>§ 8 – Verhaltensökonomik</b> . . . . .                                                | <br>237 |     |
| <i>Markus Englerth / Emanuel V. Towfigh</i>                                                  |         |     |

|                                                                                                    | Seite   | Rz. |
|----------------------------------------------------------------------------------------------------|---------|-----|
| I. Verhaltenstheorie in der Ökonomie . . . . .                                                     | 237     | 479 |
| II. Methodische und konzeptionelle Grundlagen . . . . .                                            | 239     | 484 |
| III. Einzelne Einsichten der Verhaltensökonomik und ihre<br>Bedeutung für das Recht . . . . .      | 243     | 493 |
| IV. <i>Nudging</i> : Verhaltenswissenschaftliche Rezepturen<br>für staatliche Steuerung? . . . . . | 266     | 545 |
| V. Offene Fragen . . . . .                                                                         | 273     | 559 |
| <br>Zu den Autoren . . . . .                                                                       | <br>277 |     |
| Glossar . . . . .                                                                                  | 279     |     |
| Sachwortverzeichnis . . . . .                                                                      | 287     |     |



## Inhaltsverzeichnis

|                                                                                      | Seite         | Rz. |
|--------------------------------------------------------------------------------------|---------------|-----|
| Vorwort zur zweiten Auflage . . . . .                                                | V             |     |
| Vorwort zur ersten Auflage . . . . .                                                 | VII           |     |
| Verzeichnis der Abbildungen . . . . .                                                | XXI           |     |
| Verzeichnis der Tabellen . . . . .                                                   | XXIII         |     |
| <br><b>§ 1 – Ökonomik in der Rechtswissenschaft . . . . .</b>                        | <br><b>1</b>  |     |
| I. Entwicklung der Rechtsökonomik . . . . .                                          | 2             | 1   |
| II. Normative und positive ökonomische Theorie . . . . .                             | 3             | 5   |
| III. Das Wesen sozialwissenschaftlicher Theorien . . . . .                           | 6             | 9   |
| IV. Sozialwissenschaftliche Theorie<br>und rechtswissenschaftliche Methode . . . . . | 8             | 14  |
| 1. Rechtsdogmatik . . . . .                                                          | 8             | 17  |
| 2. Rechtssetzung . . . . .                                                           | 16            | 40  |
| 3. Recht als soziales Phänomen . . . . .                                             | 17            | 43  |
| 4. Grenzen der ökonomischen Methode in der<br>Rechtswissenschaft . . . . .           | 19            | 47  |
| V. Die relevanten Methoden der Ökonomie . . . . .                                    | 23            | 57  |
| <br><b>§ 2 – Das ökonomische Paradigma . . . . .</b>                                 | <br><b>25</b> |     |
| I. Theoretische Grundannahmen . . . . .                                              | 25            | 61  |
| 1. Methodologischer Individualismus . . . . .                                        | 26            | 63  |
| 2. Ressourcenknappheit . . . . .                                                     | 27            | 64  |
| 3. Verhaltensmodell des <i>homo oeconomicus</i> . . . . .                            | 30            | 69  |
| 4. Grenzen des Modells . . . . .                                                     | 34            | 80  |
| II. Wohlfahrtsanalyse und Effizienz . . . . .                                        | 39            | 87  |
| 1. <i>Pareto</i> -Effizienz . . . . .                                                | 40            | 89  |
| 2. <i>Kaldor-Hicks</i> -Kriterium . . . . .                                          | 42            | 93  |
| <br><b>§ 3 – Nachfrage, Angebot und Märkte . . . . .</b>                             | <br><b>45</b> |     |
| I. Einleitung . . . . .                                                              | 45            | 96  |



|                                                                                                      | Seite     | Rz. |
|------------------------------------------------------------------------------------------------------|-----------|-----|
| II. Nachfrage . . . . .                                                                              | 46        | 97  |
| 1. Bewertung von Gütern . . . . .                                                                    | 47        | 98  |
| 2. Nutzenmaximierung . . . . .                                                                       | 53        | 114 |
| 3. Preisänderungen . . . . .                                                                         | 55        | 118 |
| 4. Nachfragefunktionen . . . . .                                                                     | 56        | 119 |
| III. Angebot . . . . .                                                                               | 62        | 131 |
| 1. Opportunitätskosten . . . . .                                                                     | 62        | 132 |
| 2. Einige weitere wichtige Kostenbegriffe . . . . .                                                  | 63        | 134 |
| 3. Spezielle Kosten und die Angebotskurve . . . . .                                                  | 66        | 140 |
| 4. Produzentenrente . . . . .                                                                        | 68        | 146 |
| IV. Der Markt . . . . .                                                                              | 68        | 147 |
| 1. Perfekter Wettbewerb . . . . .                                                                    | 69        | 148 |
| 2. Güter als Bündel von Rechten . . . . .                                                            | 71        | 152 |
| V. Marktversagen . . . . .                                                                           | 72        | 153 |
| 1. Märkte ohne Wettbewerb . . . . .                                                                  | 72        | 154 |
| 2. Asymmetrische Information und verborgene<br>Handlungen . . . . .                                  | 75        | 160 |
| 3. Externe Effekte, Transaktionskosten und das<br>Coase-Theorem . . . . .                            | 76        | 161 |
| 4. Nicht private Güter . . . . .                                                                     | 79        | 166 |
| 5. Beispiel Flughafen (2) . . . . .                                                                  | 80        | 167 |
| <b>§ 4 – Spieltheorie . . . . .</b>                                                                  | <b>83</b> |     |
| I. Spieltheorie und Recht . . . . .                                                                  | 83        | 170 |
| 1. Die Interdependenz von Interessen in juristischer<br>und spieltheoretischer Perspektive . . . . . | 84        | 171 |
| 2. Individuelles Entscheiden und strategische<br>Interdependenz . . . . .                            | 85        | 174 |
| 3. Spieldefinition, Normalform und Extensivform . . . . .                                            | 85        | 175 |
| II. Spiele in Normalform . . . . .                                                                   | 86        | 177 |
| 1. Das Kartelldilemma . . . . .                                                                      | 86        | 177 |
| 2. Lösungskonzepte für individuell rationales Verhalten . . . . .                                    | 89        | 182 |
| 3. Soziale Wohlfahrt und politische Gemeinwohlziele . . . . .                                        | 97        | 198 |
| III. Typen von Spielen . . . . .                                                                     | 99        | 203 |
| 1. Einfache Motive . . . . .                                                                         | 100       | 204 |
| 2. Gemischte Motive . . . . .                                                                        | 103       | 210 |
| 3. Kooperation . . . . .                                                                             | 107       | 218 |
| 4. Wiederholte Spiele . . . . .                                                                      | 114       | 231 |

|                                                                                     | Seite   | Rz.     |
|-------------------------------------------------------------------------------------|---------|---------|
| IV. Spiele in Extensivform . . . . .                                                | 117     | 235     |
| 1. Definition eines Spiels in Extensivform . . . . .                                | 117     | 235     |
| 2. Teilspielperfektion . . . . .                                                    | 118     | 239     |
| 3. Imperfekte Informationen und Informationsmengen . . . . .                        | 121     | 244     |
| 4. Unvollständige Informationen . . . . .                                           | 122     | 246     |
| V. Recht und informale Institutionen . . . . .                                      | 125     | 252     |
| 1. Recht als Preis oder Brennpunkt . . . . .                                        | 125     | 252     |
| 2. Recht und soziale Normen . . . . .                                               | 126     | 254     |
| <br>§ 5 – Vertragstheorie und ökonomische Analyse<br>des Vertragsrechts . . . . .   | <br>131 |         |
| I. Warum Verträge? . . . . .                                                        | 132     | 260     |
| 1. Austauschgeschäfte in einer idealen Welt:<br>Das <i>Coase</i> -Theorem . . . . . | <br>132 | <br>261 |
| 2. Verträge als Instrument der (Vorab-)Bindung<br>und Koordination . . . . .        | <br>134 | <br>266 |
| II. Marktstörungen als Begründung für Vertragsrecht . . . . .                       | 136     | 269     |
| III. Unvollständige Information – Problem und Lösungen . . . . .                    | 138     | 273     |
| 1. Das Problem adverser Selektion . . . . .                                         | 138     | 274     |
| 2. <i>Signaling</i> . . . . .                                                       | 140     | 277     |
| 3. <i>Screening</i> . . . . .                                                       | 142     | 283     |
| 4. Marktmacht und unvollständige Information . . . . .                              | 146     | 292     |
| IV. Kognitive Beschränkungen und nichtrationales<br>Verhalten . . . . .             | <br>148 | <br>295 |
| 1. Kognitive Schranken als Ursache unvollständiger<br>Information . . . . .         | <br>148 | <br>295 |
| 2. Staatliche Intervention durch paternalistisches<br>Vertragsrecht . . . . .       | <br>150 | <br>299 |
| V. Anreizprobleme und unvollständige Information<br>nach Vertragsschluss . . . . .  | <br>151 | <br>302 |
| 1. <i>Moral hazard</i> . . . . .                                                    | 151     | 302     |
| 2. Langzeitverträge, Opportunismus und die<br>Kostenabwägung der Parteien . . . . . | <br>155 | <br>312 |
| VI. „Verteilungsgerechtigkeit“ durch Vertragsrecht? . . . . .                       | 160     | 323     |
| <br>§ 6 – Public und Social Choice Theorie . . . . .                                | <br>163 |         |
| I. Ökonomik und Staatswissenschaft . . . . .                                        | 163     | 328     |
| II. Grundlegende Annahmen der <i>Public Choice Theory</i> . . . . .                 | 165     | 332     |

|                                                                                               | Seite | Rz. |
|-----------------------------------------------------------------------------------------------|-------|-----|
| 1. Politiker . . . . .                                                                        | 166   | 333 |
| 2. Wähler . . . . .                                                                           | 167   | 334 |
| 3. Bürokraten . . . . .                                                                       | 169   | 337 |
| III. Fehlanreize in repräsentativen Systemen . . . . .                                        | 170   | 341 |
| 1. Das Medianwähler-Theorem . . . . .                                                         | 171   | 342 |
| 2. Sonderinteressen bei Wählern und Politikern –<br><i>rent seeking</i> . . . . .             | 177   | 335 |
| 3. Budgetmaximierung bei den Bürokraten . . . . .                                             | 179   | 360 |
| IV. Kollektiventscheidungen durch Wahlen und<br>Abstimmungen: <i>Social Choice</i> . . . . .  | 183   | 370 |
| 1. Probleme bei Wahlen und Abstimmungen . . . . .                                             | 183   | 372 |
| 2. Das <i>Arrows</i> -Theorem . . . . .                                                       | 189   | 384 |
| 3. Das <i>Ostrogorski</i> -Paradox . . . . .                                                  | 191   | 390 |
| 4. Bewertung und juristische Implikationen . . . . .                                          | 192   | 392 |
| <b>§ 7 – Empirische Methoden</b> . . . . .                                                    | 195   |     |
| I. Grundlagen und Forschungsdesign . . . . .                                                  | 195   | 394 |
| 1. Forschungsdesign und Kausalität . . . . .                                                  | 196   | 397 |
| 2. Die Messung von Daten . . . . .                                                            | 201   | 409 |
| 3. Validität der Ergebnisse . . . . .                                                         | 204   | 414 |
| II. Deskriptive Statistik . . . . .                                                           | 206   | 419 |
| 1. Statistische Variable . . . . .                                                            | 207   | 420 |
| 2. Histogramme und Verteilungen . . . . .                                                     | 208   | 423 |
| 3. Kennzahlen . . . . .                                                                       | 211   | 428 |
| III. Statistische Testverfahren . . . . .                                                     | 215   | 441 |
| 1. Grundbegriffe statistischer Tests . . . . .                                                | 216   | 442 |
| 2. Auswahl des Testverfahrens . . . . .                                                       | 217   | 446 |
| <b>§ 8 – Verhaltensökonomik</b> . . . . .                                                     | 237   |     |
| I. Verhaltenstheorie in der Ökonomie . . . . .                                                | 237   | 479 |
| II. Methodische und konzeptionelle Grundlagen . . . . .                                       | 239   | 484 |
| 1. Verhaltenswissenschaftliche Komponente . . . . .                                           | 240   | 485 |
| 2. Ökonomische Komponente . . . . .                                                           | 241   | 488 |
| 3. Juristische Komponente . . . . .                                                           | 242   | 490 |
| III. Einzelne Einsichten der Verhaltensökonomik<br>und ihre Bedeutung für das Recht . . . . . | 243   | 493 |
| 1. Begrenztes Eigeninteresse . . . . .                                                        | 243   | 494 |
| 2. Begrenzte Rationalität . . . . .                                                           | 247   | 503 |

|                                                             | Seite   | Rz. |
|-------------------------------------------------------------|---------|-----|
| 3. Begrenzte Selbstdisziplin . . . . .                      | 263     | 539 |
| IV. <i>Nudging</i> : Verhaltenswissenschaftliche Rezepturen |         |     |
| für staatliche Steuerung? . . . . .                         | 266     | 545 |
| 1. Konzept . . . . .                                        | 267     | 548 |
| 2. Instrumente . . . . .                                    | 268     | 550 |
| 3. Kritik . . . . .                                         | 269     | 552 |
| 4. Rhetorisches Mittel? . . . . .                           | 271     | 556 |
| V. Offene Fragen . . . . .                                  | 273     | 559 |
| <br>Zu den Autoren . . . . .                                | <br>277 |     |
| Glossar . . . . .                                           | 279     |     |
| Sachwortverzeichnis . . . . .                               | 287     |     |



## Verzeichnis der Abbildungen

|                                                                                                                                                                                                                                                                     |     |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| Abbildung 2.1: Abnehmender Grenznutzen . . . . .                                                                                                                                                                                                                    | 73  |
| Abbildung 3.1: Indifferenzkurve . . . . .                                                                                                                                                                                                                           | 106 |
| Abbildung 3.2: Indifferenzkurve für perfekte Substitute . . . . .                                                                                                                                                                                                   | 109 |
| Abbildung 3.3: Indifferenzkurve für perfekte Komplemente<br>(zwei Paar Schuhe) . . . . .                                                                                                                                                                            | 109 |
| Abbildung 3.4: Vollständigkeit – es gibt unendlich viele Indifferenz-<br>kurven; „Mehr ist besser“ – das Nutzenniveau steigt mit größerem<br>Konsum; Kontinuität – zwischen jedem Paar von Indifferenzkurven<br>kann eine weitere Indifferenzkurve liegen . . . . . | 111 |
| Abbildung 3.5: Transitivität – Indifferenzkurven einer Person schneiden<br>sich nicht . . . . .                                                                                                                                                                     | 111 |
| Abbildung 3.6: Nutzenmaximierung mit begrenztem Budget . . . . .                                                                                                                                                                                                    | 116 |
| Abbildung 3.7: Reaktion des Verbrauchs auf Preisveränderung . . . . .                                                                                                                                                                                               | 118 |
| Abbildung 3.8: Substitutions- und Einkommenseffekt . . . . .                                                                                                                                                                                                        | 118 |
| Abbildung 3.9: Nachfrage eines Individuums . . . . .                                                                                                                                                                                                                | 121 |
| Abbildung 3.10: Aggregierte Nachfrage (zwei Individuen) . . . . .                                                                                                                                                                                                   | 121 |
| Abbildung 3.11: Vollständig inelastische Nachfrage – die nachgefragte<br>Menge reagiert nicht auf Preisveränderungen . . . . .                                                                                                                                      | 122 |
| Abbildung 3.12: Elastische Nachfrage – eine kleine Preisänderung<br>verändert die nachgefragte Menge stark . . . . .                                                                                                                                                | 122 |
| Abbildung 3.13: Konsumentenrente . . . . .                                                                                                                                                                                                                          | 126 |
| Abbildung 3.14: Durchschnittliche Kosten . . . . .                                                                                                                                                                                                                  | 138 |
| Abbildung 3.15: Grenzkosten und Angebot . . . . .                                                                                                                                                                                                                   | 138 |
| Abbildung 3.16: Produzentenrente . . . . .                                                                                                                                                                                                                          | 144 |
| Abbildung 3.17: Profit . . . . .                                                                                                                                                                                                                                    | 144 |
| Abbildung 3.18: Preis unter perfektem Wettbewerb . . . . .                                                                                                                                                                                                          | 150 |
| Abbildung 3.19: Preis unter Monopol . . . . .                                                                                                                                                                                                                       | 150 |
| Abbildung 3.20: Angebot und soziale Grenzkosten . . . . .                                                                                                                                                                                                           | 168 |
| Abbildung 3.21: <i>Pigou</i> 'sche Steuer . . . . .                                                                                                                                                                                                                 | 168 |
| Abbildung 4.1: Spielbaum des sequentiellen Standardisierungs-Spiels . . .                                                                                                                                                                                           | 236 |
| Abbildung 4.2: Gleichgewichtsstrategien des Gleichgewichts ( $Y$ , $Y/Y$ ) . .                                                                                                                                                                                      | 240 |
| Abbildung 4.3: Beispiel für <i>tree pruning</i> . . . . .                                                                                                                                                                                                           | 243 |
| Abbildung 4.4: Beispiel für ein Spiel mit imperfekter Information . . . .                                                                                                                                                                                           | 244 |
| Abbildung 4.5: Beispiel für ein Markteintrittsspiel . . . . .                                                                                                                                                                                                       | 249 |
| Abbildung 4.6: Markteintrittsspiel nach Rückwärtsinduktion . . . . .                                                                                                                                                                                                | 251 |
| Abbildung 6.1: Position des Medianwählers bei unterschiedlicher<br>Verteilung des politischen Spektrums . . . . .                                                                                                                                                   | 344 |

|                                                                                                                                                                                        |     |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| <b>Abbildung 6.2:</b> Wettbewerb der Parteien um den Medianwähler<br>im zweidimensionalen Raum . . . . .                                                                               | 350 |
| <b>Abbildung 6.3:</b> Wohlfahrtsoptimale Budget . . . . .                                                                                                                              | 364 |
| <b>Abbildung 6.4:</b> Verschwendung durch Budgetmaximierung seitens<br>der Bürokraten . . . . .                                                                                        | 366 |
| <b>Abbildung 6.5:</b> Abhängigkeit des Abstimmungsergebnisses<br>von der Agenda . . . . .                                                                                              | 374 |
| <b>Abbildung 6.6:</b> Agenda-Paradox – „eigentlich“ ziehen alle Wähler<br>Option „x“ der Option „y“ vor . . . . .                                                                      | 377 |
| <b>Abbildung 6.7:</b> <i>Condorcet-Paradox</i> . . . . .                                                                                                                               | 379 |
| <b>Abbildung 6.8:</b> <i>Inverted Order Paradox</i> . . . . .                                                                                                                          | 382 |
| <b>Abbildung 7.1:</b> Pfaddiagramm, das die Kausalverläufe eines<br>empirischen Modells darstellt . . . . .                                                                            | 406 |
| <b>Abbildung 7.2:</b> Unterdrückungseffekt . . . . .                                                                                                                                   | 407 |
| <b>Abbildung 7.3:</b> Beispiel einer Urne mit blauen und grauen Kugeln,<br>Anzahl der gezogenen grauen Kugeln in einer Simulation . . . . .                                            | 415 |
| <b>Abbildung 7.4:</b> Beispielhafte Darstellung des qualitativen Merkmals<br><i>Bundesland</i> . . . . .                                                                               | 422 |
| <b>Abbildung 7.5:</b> Histogramme mit unterschiedlichen Intervallgrößen<br>über das Bruttoerwerbseinkommen . . . . .                                                                   | 423 |
| <b>Abbildung 7.6:</b> Skizze einer Gleichverteilung (links) und einer<br>Normalverteilung (rechts) . . . . .                                                                           | 425 |
| <b>Abbildung 7.7:</b> Eigenschaften von Verteilungen – rechtsschief (links),<br>linksschief (mittig) sowie bimodal (rechts) . . . . .                                                  | 426 |
| <b>Abbildung 7.8:</b> Geschätzte Verteilung des Bruttoerwerbseinkommens . . . . .                                                                                                      | 434 |
| <b>Abbildung 7.9:</b> Histogramme zweier imaginärer Lohnverteilungen<br>mit niedriger (links) und hoher (rechts) Varianz . . . . .                                                     | 437 |
| <b>Abbildung 7.10:</b> Verteilung aus der repräsentativen Stichprobe des<br>sozio-ökonomischen Panels (durchgezogen) im Vergleich zu einer<br>Normalverteilung (gestrichelt) . . . . . | 455 |
| <b>Abbildung 7.11:</b> Punktwolke ohne Korrelation (links), mit positiver<br>Korrelation (mittig) und negativer Korrelation (rechts) . . . . .                                         | 459 |
| <b>Abbildung 7.12:</b> Arbeitslosenquote und die Anzahl rechtsextremer<br>Gewalttaten in den Bundesländern als Balkendiagramm . . . . .                                                | 461 |
| <b>Abbildung 7.13:</b> Arbeitslosenquote und die Anzahl rechtsextremer<br>Gewalttaten in den Bundesländern als Streudiagramm . . . . .                                                 | 461 |
| <b>Abbildung 7.14:</b> Geburtenrate beim Menschen und die Anzahl<br>von Storchpaaren in 17 Ländern . . . . .                                                                           | 464 |
| <b>Abbildung 7.15:</b> Streudiagramm über die gefahrenen Kilometer und<br>den Kaufpreis von Gebrauchtwagen . . . . .                                                                   | 470 |

## Verzeichnis der Tabellen

|                                                                                                |     |
|------------------------------------------------------------------------------------------------|-----|
| Tabelle 4.1: Spielmatrix des Kartell-Dilemmas . . . . .                                        | 181 |
| Tabelle 4.2: Vorzugsrelation für $D_1$ . . . . .                                               | 184 |
| Tabelle 4.3: Dominante Strategie für $D_1$ . . . . .                                           | 185 |
| Tabelle 4.4: Vorzugsrelation für $D_2$ . . . . .                                               | 186 |
| Tabelle 4.5: Dominante Strategie für $D_2$ . . . . .                                           | 186 |
| Tabelle 4.6: Abweichungsdiagramm . . . . .                                                     | 187 |
| Tabelle 4.7: Beispiel für schwache Dominanz . . . . .                                          | 189 |
| Tabelle 4.8: Spielmatrix des Standardisierungs-Spiels . . . . .                                | 192 |
| Tabelle 4.9: Spielmatrix des Diskoordinierungsspiels . . . . .                                 | 196 |
| Tabelle 4.10: Beispiel für <i>Pareto</i> -Optimalität . . . . .                                | 199 |
| Tabelle 4.11: Beispiel für <i>Kaldor-Hicks</i> -Effizienz . . . . .                            | 200 |
| Tabelle 4.12: Beispiel für ein Harmonie-Spiel . . . . .                                        | 205 |
| Tabelle 4.13: Beispiel für ein reines Konfliktspiel . . . . .                                  | 206 |
| Tabelle 4.14: Beispiel für ein reines Koordinationsspiel . . . . .                             | 208 |
| Tabelle 4.15: Beispiel für ein Falke-Taube-Spiel . . . . .                                     | 212 |
| Tabelle 4.16: Beispiel für ein Hirschjagd-Spiel . . . . .                                      | 215 |
| Tabelle 4.17: Beispiel für ein Gefangen-Dilemma . . . . .                                      | 219 |
| Tabelle 4.18: Private und kollektive Güter . . . . .                                           | 222 |
| Tabelle 4.19: Beispiel für Nicht-Dominanz von Defektion . . . . .                              | 226 |
| Tabelle 4.20: Beispiel für ein gemischtes Kooperations- und Konfliktspiel                      | 228 |
| Tabelle 4.21: Beispiel für ein Kooperationsspiel im weiteren Sinn . . . . .                    | 230 |
| Tabelle 4.22: Spielmatrix des sequentiellen Standardisierungs-Spiels . . . .                   | 239 |
| Tabelle 4.23: Beispiel für geänderte Auszahlungen durch Sanktionen . . . .                     | 252 |
| Tabelle 5.1: Selbstselektion im Versicherungsmarkt ( <i>separating equilibrium</i> ) . . . . . | 285 |
| Tabelle 5.2: Selbstselektion im Versicherungsmarkt ( <i>pooling equilibrium</i> )              | 288 |
| Tabelle 6.1: Nutzen oder Kosten je Wähler im jeweiligen Wahlbezirk . . . .                     | 358 |
| Tabelle 6.2: <i>Voting cycle</i> . . . . .                                                     | 386 |
| Tabelle 6.3: <i>Ostrogorski</i> -Paradox . . . . .                                             | 391 |
| Tabelle 7.1: Fehler beim Testen von Hypothesen . . . . .                                       | 445 |
| Tabelle 7.2: Reaktionszeit mit und ohne Benutzung eines Mobiltelefones .                       | 450 |
| Tabelle 7.3: Nettoarbeitslöhne von zufälligen ausgewählten Frauen<br>und Männern . . . . .     | 452 |
| Tabelle 7.4: Fiktive Nettoarbeitslöhne von Frauen und Männern . . . . .                        | 453 |
| Tabelle 7.5: Abgefüllte Menge Joghurt mit Verfahren 1 und 2 . . . . .                          | 454 |
| Tabelle 7.6: Bruttoerwerbseinkommen 14 zufällig ausgewählter Personen                          | 456 |



Tabelle 7.7: Preise von Gebrauchtwagen . . . . . 467

Tabelle 7.8: Lineare Regression Preis Gebrauchtwagen . . . . . 473

Tabelle 7.9: Lineare Regression Anzahl Geburten pro Jahr . . . . . 477

## § 1 – Ökonomik in der Rechtswissenschaft

**Literatur:** A. van Aaken, „Rational Choice“ in der Rechtswissenschaft, 2003; *dies.*, Vom Nutzen der ökonomischen Theorie für das öffentliche Recht, in: M. Bungenberg *et al.* (Hg.), Recht und Ökonomik, 2004, 1–31; R.H. Coase, The Problem of Social Cost, *Journal of Law & Economics* 3 (1960), 1–44; *dies.*, The Nature of the Firm, *Economica* 4 (1937), 386–405; R. Dworkin, Why Efficiency?, *Hofstra Law Review* 8 (1980), 563–590; H. Eidenmüller, Effizienz als Rechtsprinzip, 4. Aufl. 2015; C. Engel/M. Morlok (Hg.), Öffentliches Recht als ein Gegenstand ökonomischer Forschung, 1998; C. Engel, Rationale Rechtspolitik und ihre Grenzen, *JZ* 2005, 581–590; K. Grechenig/M. Gelter, Divergente Evolution des Rechtsdenkens – Von amerikanischer Rechtsökonomie und deutscher Dogmatik, *RabelsZ* 72 (2008), 513–561; S. Grundmann, Methodenpluralismus als Aufgabe, *RabelsZ* 61 (1997), 423–453; J. C. Harsanyi, Cardinal Utility in Welfare Economics and in the Theory of Risk Taking, *Journal of Political Economics* 61 (1953), 434–435; G. Janson, Ökonomische Theorie im Recht, 2004; G. Kirchgässner, Homo Oeconomicus: Das ökonomische Modell individuellen Verhaltens und seine Anwendung in den Wirtschafts- und Sozialwissenschaften, 4. Aufl. 2013; C. Kirchner, Ökonomische Theorie des Rechts, 1997; L. Kornhauser, A Guide to the Perplexed Claims of Efficiency in the Law, *Hofstra Law Review* 8 (1980), 591–639; O. Lepsius, Sozialwissenschaften im Verfassungsrecht – Amerika als Vorbild?, *JZ* 2005, 1–13; J. Lüdemann, Netzwerke, Öffentliches Recht und Rezeptionstheorie, in: S. Boysen *et al.* (Hg.), Netzwerke, 2007, 266–285; K. Mathis, Effizienz statt Gerechtigkeit?, 3. Aufl. 2009; F. Müller, Ökonomische Theorie des Rechts, in: S. Buckel/R. Christensen/A. Fischer-Lescano (Hg.), Neue Theorien des Rechts, 2006, 323–344; N. Petersen, Braucht die Rechtswissenschaft eine empirische Wende?, *Der Staat* 2010, 435–455; R. Posner, *Economic Analysis of Law*, 9. Aufl. 2014; H.-P. Schwintowski, Ökonomische Theorie des Rechts, *JZ* 1998, 581–588; S. Tontrup, Ökonomik in der dogmatischen Jurisprudenz, in: C. Engel (Hg.), Methodische Zugänge zu einem Recht der Gemeinschaftsgüter, 1998, 41–120; E. V. Towfigh, Empirical arguments in public law doctrine: Should empirical legal studies make a “doctrinal turn”?, *International Journal of Constitutional Law* 12 (2014), 670–691.

## I. Entwicklung der Rechtsökonomik

- 1 Die gemeinsamen Bemühungen von Juristen und Ökonomen um Erkenntnisgewinn haben eine wechselhafte Geschichte. In den Universitäten waren beide Disziplinen oft zu einer „Staatswissenschaftlichen Fakultät“ zusammengefasst. Dennoch gab es vor allem in der ersten Hälfte des 20. Jahrhunderts kaum gleichgerichtete Forschung. Das hing vor allem damit zusammen, dass die Nationalökonomie ihren Blick immer stärker auf die Verteilung von Gütern konzentrierte. Gleichzeitig wurden die ökonomischen Methoden immer exakter und der Rückgriff auf mathematische Ausdrucksformen immer stärker. Dies erschwerte zum einen die Rezeption ökonomischer Erkenntnisse durch andere Disziplinen, verringerte zugleich aber auch die Relevanz ökonomischer Forschung für konkrete wirtschaftspolitische Forderungen.
- 2 Daneben spielte auch die Sorge vor Autonomieverlusten der Fächer eine nicht zu unterschätzende Rolle. Juristen und Ökonomen warfen sich wechselseitig disziplinären Imperialismus vor.<sup>1</sup> Durch diese Trennung konnten sich die Disziplinen zwar unabhängig voneinander methodisch entwickeln und profilieren, verloren aber auch das große Potenzial gemeinschaftlicher, interdisziplinärer Ansätze aus dem Blick.<sup>2</sup>
- 3 Etwa seit den 1960er-Jahren befassen sich Ökonomen wieder verstärkt auch mit Fragen des Rechts. Aus verschiedenen theoretischen Ansätzen – etwa der Theorie der Verfügungsrechte,<sup>3</sup> der *Principal-Agent-Theorie*<sup>4</sup> und der *Neuen Politischen Ökonomie*<sup>5</sup> – entwickelte sich die *Neue Institutionenökonomik*. Mit der Fortentwicklung dieser Strömung löste sich die Volkswirtschaftslehre zunehmend von ihrem klassischen Gegenstand. Waren einst allein volkswirtschaftliche Prozesse, in erster Linie das Geschehen am Markt, im Blickfeld ihrer Forschung, werden nun auch Entscheidungen in anderen institutionellen Gefügen mit den eigenen Methoden ins Visier genommen – ursprünglich in Unternehmungen<sup>6</sup> und „industriellen Organisationen“, später aber etwa auch politische Prozesse. Das Proprium ökonomischer Forschung ergibt sich fortan nicht mehr aus dem Betrachtungsgegenstand der Ökonomie, sondern aus ihrer Methodik, der

<sup>1</sup> *Kirchner*, Ökonomische Theorie des Rechts, 1997, S. 12; *Kirchgässner*, *Homo Oeconomicus*, 2013, S. 153.

<sup>2</sup> *Janson*, Ökonomische Theorie, 2004, S. 20.

<sup>3</sup> Dazu unten Rz. 152.

<sup>4</sup> Dazu unten Rz. 304 ff.

<sup>5</sup> Hierzu Rz. 328 ff.

<sup>6</sup> Epochal: *Coase*, The Nature of the Firm, *Economica* 4 (1937), 386 ff.

„Ökonomik“. Dabei meint **Ökonomie** jene Wissenschaft, die sich mit wirtschaftlichen Zusammenhängen befasst, die Wirtschaft also zum Gegenstand hat; die **Ökonomik** dagegen wendet die *Methodik* der Ökonomie, ihr Instrumentarium, auch auf nicht-volkswirtschaftliche Fragestellungen an.<sup>7</sup> Wie wir im Einzelnen sehen werden, untersucht die Ökonomik ganz allgemein menschliches Entscheidungs- oder Wahlverhalten unter der Annahme knapper Ressourcen.

Mit der Hinwendung zur Ökonomik bedurfte es nunmehr nur eines kurzen Stück Weges hin zum Recht. Denn die Neue Institutionenökonomik meint mit „Institution“ ein System miteinander verknüpfter formaler und informeller Normen zur Verhaltenssteuerung einschließlich der Vorkehrungen zu deren Durchsetzung<sup>8</sup> – und behandelt damit u. a. wichtige Bereiche des Rechts. Die Nähe zum weiten Feld der auch von Juristen und Politikern betriebenen „Steuerungswissenschaften“, bei denen es um die Erforschung der Steuerungswirkung von Normen geht, sieht man dieser Definition unmittelbar an. Vor allem in den USA, aber auch in Kontinentaleuropa hat sich mit den Jahren eine als „Law & Economics“ bekannte Schnittmengendisziplin etabliert,<sup>9</sup> in der Juristen und Ökonomen gleichermaßen aktiv sind.

## II. Normative und positive ökonomische Theorie

Gegen die ökonomische Analyse des Rechts bestanden in der deutschen Rechtswissenschaft lange große Vorbehalte, die auch heute noch nicht vollständig ausgeräumt sind. So schrieb *Karl-Heinz Fezer* einst, dass „[ö]konomische Rechtsanalyse und freiheitliches Rechtsdenken“ schlechthin „unvereinbar“ seien.<sup>10</sup> Diese Bedenken gründen sich vor allem auf dem Selbstverständnis der Ökonomie als Wissenschaft und dem normativen Anspruch vieler Vertreter der Rechtsökonomie. Während die Ökonomie früher vor allem der Analyse wirtschaftlicher Zusammenhänge vorbehalten war, versteht sie sich immer mehr als umfassende Sozialwissenschaft, die menschliches Verhalten in allen Lebensbereichen zu erklären beansprucht – von der Kriminalität über die Demokratie bis hin zur Sexuali-

<sup>7</sup> Janson, *Ökonomische Theorie*, 2004, S. 21.

<sup>8</sup> Richter/Furubotn, *Neue Institutionenökonomik*, 4. Aufl. 2010, S. 7.

<sup>9</sup> Grechenig/Gelter, *Divergente Evolution des Rechtsdenkens – Von amerikanischer Rechtsökonomie und deutscher Dogmatik*, *RabelsZ* 72 (2008), 513 ff.

<sup>10</sup> Fezer, *Aspekte einer Rechtskritik an der economic analysis of law und am property rights approach*, *JZ* 1986, 817 (823).

tät.<sup>11</sup> Kritiker sehen darin einen Wissenschaftsimperialismus, der versucht, die ökonomische Rationalität des Strebens nach dem Eigennutz auf alle Gesellschaftsbereiche zu übertragen. Nahrung bekommt diese Kritik durch Aussagen von Exponenten der ökonomischen Analyse des Rechts, die „weiterhin glauben, dass die Wohlstandsmaximierung Leitlinie für die staatliche Politik in allen Gesellschaftsbereichen sein soll“.<sup>12</sup> Wie so oft bei hitzig geführten Debatten liegt die Wahrheit wohl in der Mitte. Selbstverständlich kann die Wohlfahrtssteigerung nicht alleiniges Ziel der Politik, Effizienzdenken nicht alleiniger normativer Maßstab der Rechtswissenschaft sein. Auf der anderen Seite bieten die Sozialwissenschaften im Allgemeinen und die Ökonomie im Besonderen jedoch nützliche Analyseinstrumente, die durchaus auch dem Rechtswissenschaftler die Arbeit erleichtern können. Dies wird deutlich, wenn wir die beiden unterschiedlichen Formen der Rechtsökonomie auseinanderhalten – die positive und die normative ökonomische Theorie des Rechts.

- 6 Unter dem Begriff der **positiven ökonomischen Theorie des Rechts** werden all jene Herangehensweisen verstanden, die das Recht analytisch oder empirisch betrachten, die sich mithin „positiv“ mit der Beschreibung, Erklärung und Prognose menschlichen Verhaltens im Hinblick auf das Recht befassen. Für die positive ökonomische Analyse ist das Recht ein soziales Phänomen, es geht ihr darum, das Wissen über die soziale Welt zu verbessern. Auf der Grundlage der positiven ökonomischen Theorie sollen prognostische Aussagen zum Verhalten der Rechtssubjekte formuliert werden, die etwa von veränderten Anreizen bei der Einführung neuer Normen ausgehen können. Recht wird dabei als Mechanismus verstanden, der Handlungsalternativen verbilligt oder verteuert.<sup>13</sup> Wenn das Recht also Diebstahl unter Strafe stellt, dann hat jeder Mensch grundsätzlich noch die tatsächliche Möglichkeit, Dinge zu stehlen. Das Strafrecht erlegt ihm jedoch für dieses Verhalten Kosten in Form von drohenden Geld- oder Gefängnisstrafen auf. Wenn diese Kosten den aus dem Diebstahl gezogenen Nutzen übersteigen, führt dies nach der ökonomischen Theorie dazu, dass ein potentieller Dieb vom Stehlen absehen wird.

<sup>11</sup> Vgl. nur *Becker*, Ökonomische Erklärung menschlichen Verhaltens, 2. Aufl. 1993.

<sup>12</sup> *Posner*, A Reply to some Recent Criticism of the Efficiency Theory of the Common Law, *Hofstra Law Review* 9 (1981), 775 (780).

<sup>13</sup> *Van Aaken*, Vom Nutzen der ökonomischen Theorie für das öffentliche Recht: Methode und Anwendungsmöglichkeiten, in: *Bungenberg et al.* (Hg.), *Recht und Ökonomie*, 2004, 1 (6).

Diese Kenntnis positiver Folgen bestimmter Interventionen ist erforderlich, wenn steuernd in das Sozialgefüge eingegriffen werden soll. Nur so können funktionelle Zusammenhänge erkannt und im Rahmen von Normwirkungsanalysen genutzt werden. Gleichzeitig gibt es Bemühungen, *ex post* die Entwicklung oder Erforderlichkeit rechtlicher Institutionen zu erklären, beispielsweise des Strafrechts<sup>14</sup> oder des Eigentums- und Immaterialgüterschutzes.<sup>15</sup> Für Juristen ist die positive ökonomische Theorie in vielerlei Hinsicht ein nützliches Werkzeug – in der Rechtstatsachenforschung, bei der Abschätzung von Gesetzesfolgen (für den Gesetzgeber) und der Folgenorientierung in der Rechtsanwendung (insbesondere für die Rechtsprechung), aber auch in der Rechtsvergleichung.

Der zweite Ansatz, die **normative oder präskriptive ökonomische Theorie des Rechts**, versucht, ausgehend von theoretisch wohl definierten Prämissen, zu bewerten, welche von verschiedenen denkbaren Normgestaltungen für die Lösung eines betrachteten Problems vorzugswürdig ist: Wie sollte das Recht gestaltet werden? Welche Ziele sollte es haben? Sie steht damit Rechtspolitik und Rechtsphilosophie nahe. Für derlei normative Aussagen bedarf es eines Maßstabes, anhand dessen die Bewertung vorgenommen wird. Die Ökonomie ist dabei eng mit der **Philosophie des Utilitarismus** verknüpft.<sup>16</sup> Dem Utilitarismus zufolge ist es das Ziel einer jeden Gesellschaftsordnung, das Glück aller Menschen zu maximieren. Ersetzen wir nun Glück durch Präferenzen, sehen wir, wie aus der sozialwissenschaftlichen ökonomischen Theorie plötzlich normative Schlussfolgerungen gezogen werden können. Schließlich geht es der Ökonomie gerade darum, die menschlichen Präferenzen so zu befriedigen, dass der Gesamtnutzen maximiert wird. Die effizienteste Gesellschaftsordnung ist nach utilitaristischer Vorstellung die vorzugswürdige. Die normativen Forderungen des Utilitarismus können jedoch mit den Forderungen konkurrierender normativer Systeme, wie etwa der Rechtsordnung, kollidieren. Allerdings ist die normative Verknüpfung der Ökonomie mit dem Utilitarismus nur für die normative ökonomische Theorie notwendig. Dagegen kann die Ökonomie der Rechtswissenschaft in ihrer positiven, beschreibenden Funktion zu einem beträchtlichen **Rationalitätsgewinn** verhelfen,

<sup>14</sup> Etwa *Becker*, Ökonomische Erklärung menschlichen Verhaltens, 2. Aufl. 1993, S. 39 ff. und unten § 8.

<sup>15</sup> Vgl. etwa *Schäfer/Ott*, Lehrbuch der ökonomischen Analyse des Zivilrechts, 5. Aufl. 2012, S. 549 ff.

<sup>16</sup> *Eidenmüller*, Effizienz, 2015, S. 173.

ohne dabei einen ökonomischen Imperialismus zu betreiben oder die Normativität des Rechts zu überlagern.<sup>17</sup>

### III. Das Wesen sozialwissenschaftlicher Theorien

- 9 Die Ökonomie ist nicht die einzige Wissenschaft, die sich mit menschlichem Verhalten beschäftigt. Vielmehr gibt es andere Sozialwissenschaften, die sich mit menschlichem Handeln aus anderer Perspektive beschäftigen, etwa die Soziologie, die Anthropologie oder die Psychologie. Obwohl dieses Buch sich in erster Linie mit der ökonomischen Methode beschäftigt, nimmt es auf andere Sozialwissenschaften Bezug, wenn dies geboten erscheint. Dies gilt insbesondere für die Psychologie, soweit diese Kritik am Bild des Menschen als rational handelndem Akteur übt.<sup>18</sup> Trotz aller Differenzen im Detail sind sich die Sozialwissenschaften in ihren Methoden sehr ähnlich, weswegen wir im Folgenden zunächst allgemein auf das Wesen sozialwissenschaftlicher Theoriebildung eingehen werden.
- 10 Sozialwissenschaftliche Theorien beschäftigen sich mit menschlichem Verhalten und gesellschaftlichen Prozessen. Sie haben dabei vor allem zwei Funktionen: Zum einen sollen sie bestimmte Phänomene erklären, zum anderen generelle Gesetzmäßigkeiten aufzeigen und es so ermöglichen, Vorhersagen zu treffen. Halten sich Menschen an rechtliche Normen? Warum halten sie sich an rechtliche Normen? Unter welchen Bedingungen halten sie sich an rechtliche Normen? Das Problem ist dabei, dass soziale Prozesse oft wesentlich komplexer sind als viele naturwissenschaftliche Zusammenhänge, so dass es schwieriger ist, Gesetzmäßigkeiten zu identifizieren, und diese auch meistens eine geringere Vorhersagekraft haben.
- 11 Auch naturwissenschaftliche Gesetzmäßigkeiten hängen immer von bestimmten Voraussetzungen ab. Wenn ich einen Körper aus einer bestimmten Höhe fallen lasse, kann ich grundsätzlich vorhersagen, wann und wo dieser Körper auf dem Boden auftrifft. Allerdings finden solche Vorgänge selten in einem Vakuum statt. Vielmehr hängt der Zeitpunkt des Auftreffens vom Luftwiderstand und von den Windverhältnissen ab und daher gleichzeitig auch von Volumen und Beschaffenheit des Körpers, den ich fallen lasse, so dass eine Vorhersage unter realen Umweltbedingungen nur theoretisch möglich ist. Ähnlich sieht es auch mit sozialwissenschaftlichen

<sup>17</sup> Van Aaken, Vom Nutzen der ökonomischen Theorie für das öffentliche Recht: Methode und Anwendungsmöglichkeiten, in: Bungenberg *et al.* (Hg.), Recht und Ökonomik, 1 (31).

<sup>18</sup> Vgl. zu *Behavioral Law and Economics* ausführlich § 8.

Theorien aus. Man wird in den seltensten Fällen strikte Wenn-dann-Beziehungen feststellen können. Vielmehr können allenfalls Faktoren identifiziert werden, die das Auftreten eines bestimmten Phänomens wahrscheinlicher machen.

Ob ein bestimmter Mensch sich an eine Rechtsnorm hält, hängt nicht nur davon ab, um was für eine Rechtsnorm es sich handelt, was für eine soziale Vorprägung er hat, in was für einer Stimmung er sich befindet, sondern auch von anderen, ähnlichen Umweltfaktoren. So gibt es Menschen, die bestimmte Rechtsnormen, wie etwa das Verbot von Mord und Totschlag, nie brechen würden, gegen andere Normen aber teilweise bewusst verstoßen. So ist es für einige Menschen nicht unüblich, die Straße trotz einer roten Ampel zu überqueren. Doch auch dieser potentielle Normverstoß hängt von verschiedenen Umweltbedingungen ab. Nachts um zwei mag die Wahrscheinlichkeit größer sein als morgens um halb neun, wenn auf der anderen Seite der Straße eine Gruppe von Schulkindern wartet. Ökonomen beschäftigen sich zwar nicht mit dem Verhalten einzelner Personen, sondern dem aggregierten Verhalten mehrerer Individuen. Dennoch spielen hier ähnliche Überlegungen eine Rolle wie beim einzelnen Individuum. Die beeinflussenden Faktoren werden lediglich komplexer.

Trotz der Bedingtheit menschlichen Verhaltens können wir bestimmte allgemeine Gesetzmäßigkeiten formulieren. So können wir etwa sagen, dass Staaten mit einem hohen wirtschaftlichen Entwicklungsstand eher zur Demokratie als Staatsform neigen als Staaten mit geringer wirtschaftlicher Entwicklung. Das bedeutet jedoch nicht, dass dieser Zusammenhang monokausal und deterministisch ist, dass ein Staat zwingend umso demokratischer ist, je höher sein Bruttoinlandsprodukt ist. Vielmehr gibt es auch andere Faktoren, die in diesem Kontext eine Rolle spielen – etwa die soziale Struktur der Gesellschaft, die politische und wirtschaftliche Machtkonzentration innerhalb des Staates oder gar historische Zufälligkeiten. Dennoch sind auch probabilistische Aussagen über die soziale Realität wichtige Erkenntnisse, solange wir uns vor Augen führen, dass wir sie nicht monokausal interpretieren dürfen. Ökonomische Theorien treffen also meistens nicht Aussagen der Form, dass X *immer* zu Y führt. Vielmehr haben sie eine weniger absolute Form. Sie können beispielsweise aussagen, dass X die Wahrscheinlichkeit von Y *erhöht*, wenn bestimmte Rahmenbedingungen konstant gehalten werden.



#### IV. Sozialwissenschaftliche Theorie und rechtswissenschaftliche Methode

- 14 In diesem Abschnitt wollen wir konkreter darauf eingehen, welche Rolle die Sozialwissenschaften für die Rechtswissenschaft spielen können. In der Rechtswissenschaft gibt es – vereinfacht – drei unterschiedliche Perspektiven, aus denen Forschungsfragen gestellt werden können. Erstens wird nach dem Inhalt bestehender rechtlicher Normen gefragt – was ist das Recht? Diese Beschäftigung mit Recht aus der Innenperspektive des Richters oder Rechtsanwenders ist die in Deutschland am weitesten verbreitete Art der wissenschaftlichen Beschäftigung mit Recht. Sie zielt sowohl auf die Auslegung einzelner Normen als auch die Systematisierung ganzer Rechtsgebiete ab. Die Funktion ist die wissenschaftliche Vorbereitung von Gerichtsentscheidungen sowie deren Kritik und Einordnung in das bestehende dogmatische System.
- 15 Auch wenn die dogmatische Perspektive in Deutschland bisher dominiert, ist sie nicht die einzige Art, *Rechtswissenschaft* zu betreiben. Vielmehr kann rechtswissenschaftliche Forschung sich zweitens über den **optimalen Inhalt rechtlicher Normen** Gedanken machen – wie sollte das Recht sein? Adressat entsprechender Studien sind nicht die Gerichte, sondern ist der Gesetzgeber. Derartige Arbeiten sind in der deutschen Rechtswissenschaft weniger verbreitet, genießen aber vor allem in den USA große Popularität.
- 16 Schließlich kann man sich, drittens, mit **Recht als sozialem Phänomen** beschäftigen und seine **Wirkungsweise** oder sein Verhältnis zur Gerechtigkeit untersuchen. Hier setzt man sich als Forscher nicht an die Stelle eines bestimmten Akteurs, sondern nimmt vielmehr die Beobachterperspektive ein und versucht, Recht aus einer Außenperspektive zu betrachten. Je nachdem, welche dieser drei Perspektiven man einnimmt, haben sozialwissenschaftliche Methoden im Allgemeinen und die ökonomische Theorie im Besonderen ein anderes Forschungsfeld und einen anderen Nutzen. Daher sollen die drei unterschiedlichen Perspektiven im Folgenden getrennt betrachtet werden. In einem vierten Schritt wird dann auf die Grenzen der ökonomischen Analyse in der Rechtswissenschaft eingegangen.

##### 1. Rechtsdogmatik

- 17 In der Rechtsdogmatik scheint die Anwendung sozialwissenschaftlicher Methoden auf den ersten Blick fern zu liegen. Folgt man dem klassischen Schema der Gewaltenteilungslehre, dann beschäftigt sich die Rechtswis-

senschaft mit der **Anwendung und Auslegung** des durch den Gesetzgeber gesetzten Rechts. Die Gesetzesauslegung wird dabei nach einem festen, vorher definierten Methodenkanon vorgenommen. Sie richtet sich nach dem Wortlaut und der Systematik der Norm sowie deren Entstehungsgeschichte und ihrem Sinn und Zweck. Normauslegung ist dabei ein Verfahren, das eine **normative Schlussfolgerung** zu begründen sucht, während sich die Sozialwissenschaften vornehmlich *positiv* mit der Beschreibung und Erklärung der Wirklichkeit beschäftigen. Die Wirklichkeit kommt nach der traditionellen juristischen Methode allerdings nicht auf der Stufe der Normauslegung, sondern erst bei der Subsumtion unter die entsprechende Norm ins Spiel.

Allerdings ist diese Darstellung zu einfach. Auslegung einer Norm und positive Beschreibung der Wirklichkeit lassen sich nicht so trennscharf auseinander halten, wie es das dargestellte holzschnittartige Schema suggeriert. Vielmehr basieren normative Konzepte oft auf **tatsächlichen Annahmen**. Teilweise setzen einzelne Auslegungsmethoden einen starken Wirklichkeitsbezug voraus. Insbesondere in drei Fällen spielt der sozialwissenschaftliche Bezug bei der Normauslegung eine entscheidende Rolle: bei der teleologischen Auslegung, bei der Abwägung konkurrierender Rechtspositionen in der Grundrechtsprüfung und bei der Konkretisierung von Normen, die dem Richter einen beträchtlichen **Interpretationsspielraum** lassen. 18

#### a. Teleologische Auslegung

Die teleologische Auslegung, die Orientierung am **Sinn und Zweck** einer Norm, zählt zu den zentralen Auslegungsmethoden im Zivilrecht. An dieser Auslegungsmethode wird jedoch oft Kritik geübt, da sie der Beliebigkeit des Interpreten Tür und Tor zu öffnen scheint.<sup>19</sup> Eine sozialwissenschaftliche Einhegung der teleologischen Auslegung kann die Argumentation mit dieser Interpretationsfigur jedoch rationalisieren. Die Sozialwissenschaften helfen dabei nicht bei der Identifizierung des Zwecks, handelt es sich bei dieser doch um eine normative und nicht um eine positive Frage. Ist der Zweck jedoch identifiziert, kann die sozialwissenschaftliche Forschung dabei helfen, die Auslegung der Norm zu identifizieren, die diesem Ziel am ehesten gerecht wird. Ist durch vertragsrechtliche Regelungen also eine möglichst effiziente Güterallokation beabsichtigt, dann sind ökonomische 19

<sup>19</sup> Vgl. Müller/Christensen, Juristische Methodik I, 11. Aufl. 2013, Rz. 364.

Effizienzgesichtspunkte bei der Auslegung der Norm als Leitlinien heranzuziehen.<sup>20</sup>

- 20 Ein Beispiel bieten die Zuteilungsregeln im Emissionshandelsrecht.<sup>21</sup> Beim Emissionshandel geht es darum, dass Unternehmen Berechtigungen für CO<sub>2</sub>-Emissionen erwerben müssen. Am Beginn jeder Handelsrunde wird den Unternehmen dabei eine gewisse Zahl an Berechtigungen zugeteilt, mit denen diese dann untereinander handeln können, wenn sie überschüssige oder nicht ausreichende Berechtigungen haben. Diese Zuteilungsregeln dienen dazu, den beteiligten Unternehmen zu Beginn einer Handelsperiode Emissionsrechte zuzuteilen. Zweck des Emissionshandels ist es, Anreize zu einer möglichst kostengünstigen Modernisierung von Altanlagen zu setzen, um eine Reduktion der Treibhausgasemissionen zu erreichen. Unter mehreren Anlagenbetreibern wird durch den Emissionshandel derjenige Betreiber die größten Modernisierungsanreize haben, der diese Modernisierung am kostengünstigsten vornehmen kann. Nun könnte man versucht sein, die Zuteilungsregelungen, mit denen die ursprünglichen Emissionsberechtigungen verteilt werden, als Anreizregelungen zu verstehen. Je weniger Berechtigungen einer bestimmten Anlage zugeteilt werden, desto stärkere Anreize hat diese zur Modernisierung.
- 21 Eine solche Auslegung widerspräche jedoch der ökonomischen Logik des Emissionshandelsrechts. Anreize zur Modernisierung entstehen nämlich dann, wenn die Emissionsreduktion durch Modernisierung günstiger ist als die Preise für entsprechende Emissionszertifikate. Besitzt der Anlagenbetreiber die Emissionszertifikate bereits, dann kann er diese verkaufen und damit seine Modernisierung finanzieren. Besitzt er sie noch nicht, spart er sich die Kosten für den Zukauf.<sup>22</sup> Die Zuteilungsregeln haben also auf die Modernisierungsanreize keinen Einfluss. Sie sind reine Verteilungsregeln, so dass bei ihrer Auslegung eher Gleichheits- und Wettbewerbsaspekte zu berücksichtigen sind als Umwelterwägungen. Insofern wird man bei der Auslegung der Zuteilungsregeln kaum zu adäquaten Ergebnissen kommen, wenn man den ökonomischen Kontext des Emissionshandelsystems nicht berücksichtigt.
- 22 Ein weiteres Beispiel, bei dem die teleologische Auslegung durch sozialwissenschaftliche Erkenntnisse informiert werden kann, findet sich im

---

<sup>20</sup> Grundmann, Methodenpluralismus als Aufgabe, *RabelsZ* 61 (1997), 423 (434).

<sup>21</sup> Vgl. Magen, Rechtliche und ökonomische Rationalität im Emissionshandelsrecht, in: Towfigh *et al.* (Hg.), *Recht und Markt*, 2009, 9 ff.

<sup>22</sup> Dies ist die Logik des *Coase*-Theorems; zu diesem s. unten Rz. 161 ff.

Glücksspielrecht.<sup>23</sup> Die Veranstaltung von Glücksspielen ist gem. § 284 StGB unter Strafe gestellt. Bei der Konkretisierung dieser Vorschrift unterscheiden Rechtsprechung und Literatur dabei zwischen Glücks- und Geschicklichkeitsspielen. Nur erstere sind strafbewährt, während die Veranstaltung von Geschicklichkeitsspielen straffrei ist. Die Unterscheidung zwischen beiden Konzepten ist jedoch nicht immer ganz einfach und insbesondere dort umstritten, wo Spiele sowohl Glücks- als auch Geschicklichkeitsargumente enthalten. So wird insbesondere bei der Einordnung von Sportwetten diskutiert, ob diese als Glücks- oder als Geschicklichkeitsspiele zu qualifizieren sind.

Als Ziel von § 284 StGB wird überwiegend die Eindämmung der mit Glücksspielen einhergehenden Suchtgefahr angesehen. Legt man § 284 StGB also teleologisch aus, müsste man ermitteln, ob bei Sportwetten, die Glücks- und Geschicklichkeitselemente miteinander verbinden die Suchtgefahr ähnlich groß ist wie bei reinen Glücksspielen. Empirische Untersuchungen deuten darauf hin, dass die Suchtgefahr bei gemischten Spielen sogar größer ist als bei reinen Glücksspielen, da die Spieler einer Kontrollillusion unterliegen.<sup>24</sup> Eine teleologische Auslegung legt demnach nahe, auch Sportwetten in den Anwendungsbereich des Glücksspielverbots des § 284 StGB miteinzubeziehen. 23

#### b. Verhältnismäßigkeitsprüfung bei den Grundrechten

Ein zweites Feld, in dem sozialwissenschaftliche Methoden eine zentrale Rolle spielen können, ist die **Grundrechtsdogmatik**. Dies liegt an der Struktur der Grundrechtsprüfung, die von der Auslegung und Anwendung anderer Normen in einigen Punkten deutlich abweicht. Der Kern jeder Grundrechtsprüfung ist die Prüfung der Verhältnismäßigkeit im engeren Sinne, bei der **konkurrierende normative Ziele** gegeneinander abgewogen werden müssen. Bei dieser Abwägung vermischen sich faktische und normative Elemente der Auslegung miteinander, so dass der sozialwissenschaftlichen Wirklichkeitsbeschreibung in diesem Bereich ein besonderes Gewicht zukommt. 24

Dies soll im Folgenden am Beispiel des Apothekenurteils des Bundesverfassungsgerichts verdeutlicht werden.<sup>25</sup> In diesem wollte sich der Beschwer- 25

<sup>23</sup> Dazu ausführlich *Glöckner/Towfigh*, Geschicktes Glücksspiel. Die Sportwette als Grenzfall des Glücksspielrechts, JZ 2010, 1027 ff.

<sup>24</sup> *Towfigh/Glöckner*, Game Over. Empirical support for soccer bets regulation, Psychology, Public Policy, and Law 17 (2011), 475 ff.

<sup>25</sup> BVerfGE 7, 377 (Apothekenurteil [1958]).

deführer, ein approbierter Apotheker, in einer bayrischen Gemeinde mit einer Apotheke niederlassen. In Bayern stand die Errichtung einer Apotheke zu jener Zeit jedoch unter einem Erlaubnisvorbehalt. Die Erlaubnis wurde dem Apotheker durch die Regierung von Oberbayern versagt, weil es in der Gemeinde, in der er sich niederzulassen beabsichtigte, bereits eine Apotheke gebe und angesichts der Größe der Gemeinde nicht mehr als eine Apotheke lebensfähig sei. Daraufhin wandte sich der Apotheker an das Bundesverfassungsgericht und argumentierte, dass der Erlaubnisvorbehalt ihn in seiner Berufsfreiheit verletze.

- 26 Sicherlich kann allein der Schutz der bestehenden Apotheke vor Konkurrenz kein Argument zur Beschränkung der Berufsfreiheit eines Dritten sein, da Konkurrenz gerade ein essentielles Element einer marktorientierten Wirtschaftsordnung ist. Der Schutz vor Konkurrenz muss vielmehr einem übergeordneten Ziel dienen. Die Regierung von Oberbayern argumentierte, dass dies im konkreten Fall der Schutz der Volksgesundheit sei. Nun kann sich das Bundesverfassungsgericht in seiner Entscheidung nicht damit begnügen, abstrakt die Volksgesundheit gegen die Berufsfreiheit abzuwägen. Vielmehr muss es sich im Rahmen der Erforderlichkeit und der Angemessenheit die Frage stellen, wie groß die Gefahr für die Volksgesundheit ist, und was für alternative Schutzmechanismen es gäbe.
- 27 Für den Erhalt einer Zulassungsbeschränkung lassen sich im Wesentlichen zwei Argumente anführen: Zum einen könnte Konkurrenz zum Zusammenbruch des gesamten Apothekensystems führen, weil sich die Apotheken gegenseitig ruinieren. Zum anderen könnte sie Anreize setzen, gegen gesetzliche Vorschriften zu verstoßen, indem etwa, um wirtschaftlich zu überleben, verschreibungspflichtige Medikamente auch ohne Rezept an Kunden abgegeben werden. Bei beiden Fragen handelt es sich jedoch um empirische Argumente, die zu ihrer Überprüfung sozialwissenschaftlicher Methoden bedürfen: Wie hoch ist die Gefahr, dass die Apothekenversorgung zusammenbricht, wenn man das System der Konkurrenz freigibt? Wie wahrscheinlich ist es, dass Apotheker gegen bestimmte Ordnungsvorschriften verstoßen?
- 28 Zudem kann man sich Gedanken darüber machen, ob es nicht andere Mittel gibt, diesen Gefahren vorzubeugen. Man könnte etwa an einen staatlichen Subventionsschirm denken, um ein Zusammenbrechen des Apothekennetzes zu verhindern – oder an stärkere staatliche Kontrollen, um einem Verstoß gegen Ordnungsvorschriften vorzubeugen. Beide Maßnahmen kosten natürlich Geld. In diesem Zusammenhang kann eine vergleichende ökonomische Analyse der beiden Systeme nicht die normative Entscheidung abnehmen, ob letztlich die individuelle Berufsfreiheit höher

zu veranschlagen ist oder doch die Kosten alternativer Maßnahmen stärker ins Gewicht fallen. Die ökonomische Analyse kann jedoch die Entscheidungsgrundlage geben, da die normative Bewertung entscheidend davon abhängen wird, wie hoch die Kosten für die jeweiligen Alternativenmaßnahmen sind und wie hoch die Gefahr eines Zusammenbruchs des Systems überhaupt ist.

Ähnliche Überlegungen wie hier im Rahmen der Berufsfreiheit kann man auch in anderen Bereichen anstellen, in denen letztlich über einen Konflikt konkurrierender Prinzipien zu entscheiden ist. Ein Beispiel ist etwa die Überprüfung des deutschen Notariatssystems vor dem Maßstab der europäischen Niederlassungsfreiheit. Dadurch, dass die Zahl der Notare durch den Staat streng begrenzt ist, wird eine objektive Zulassungsschranke zum Notariatsberuf errichtet. Diese Beschränkung wird durch ökonomische Erwägungen gerechtfertigt:<sup>26</sup> Durch die Unparteilichkeit des Notars werde Informationsasymmetrien vorgebeugt und zudem würden positive externe Effekte internalisiert, da dem Notar eine Beweisfunktion zukomme, die letztlich auch Dritte und somit den Rechtsverkehr schütze.

Es ist nun Gegenstand ökonomischer Analysen herauszufinden, ob diese Ziele möglicherweise auch durch alternative Mechanismen mit geringeren Eingriffen in die Niederlassungsfreiheit erreicht werden können – oder ob das System aufgrund der ihm eigenen Beschränkung des Wettbewerbs möglicherweise sogar zu Ineffizienzen statt zur Wohlfahrtssteigerung führt, oder die Effizienzgewinne möglicherweise so gering sind, dass sie einen Eingriff in individuelle Freiheiten nicht rechtfertigen können. Auch in diesem Fall steht am Ende eine normative Abwägung, deren Ergebnis durch die ökonomische Analyse nicht determiniert ist. Die ökonomische Analyse liefert dabei jedoch die notwendige Tatsachengrundlage, aufgrund derer eine informierte Abwägung überhaupt erst vorgenommen werden kann.

### c. Normkonkretisierung

Eine letzte Auslegungsart, bei der sozialwissenschaftliche Erkenntnisse nützlich sein können, ist die Normkonkretisierung. Damit sind alle Arten der Normauslegung gemeint, die außerhalb des klassischen Interpretationskanons stehen. Dies kann bei Generalklauseln der Fall sein, aber auch bei Konzepten mit generalklauselartiger Weite. Der Gesetzgeber hat bei

<sup>26</sup> Vgl. dazu *van den Bergh/Montangie*, Competition in professional services markets: Are latin notaries different?, *Journal of Competition Law & Economics* 2 (2006), 189 ff.

der Regelung bestimmter Materien die Wahl zwischen engen und weiten, generalklauselartigen Tatbeständen. Je enger der Tatbestand ist, desto stärker werden richterliche Entscheidungen durch den Gesetzgeber vorherbestimmt. Bei weiten Tatbeständen haben Richter dagegen oft einen erheblichen Konkretisierungsspielraum.

- 32 Orientiert man sich am *Montesquieu*'schen Gewaltenteilungsmodell, dann scheinen enge Tatbestände grundsätzlich vorzugswürdig. Je enger der Tatbestand, desto geringer ist der Einschätzungsspielraum des Richters. Allerdings setzt dies voraus, dass der Gesetzgeber jeden möglichen Fall, der in der Wirklichkeit auftritt, genau vorhersagen und entsprechend normativ regeln kann. Je komplexer jedoch die zu regelnde Realität ist und je schwieriger es ist, Prognosen über zukünftige Fallkonstellationen zu machen, desto eher wird der Gesetzgeber auf Generalklauseln oder konkretisierungsbedürftige Konzepte zurückgreifen.
- 33 Ein Beispiel aus dem Zivilrecht ist etwa der Fahrlässigkeitsmaßstab, der als die im Verkehr erforderliche Sorgfalt definiert ist (§ 276 Abs. 2 BGB). Welche Sorgfalt im Verkehr erforderlich ist, wird jedoch nicht gesagt. Hier hätte der Gesetzgeber zumindest theoretisch die Möglichkeit gehabt, alle möglichen Handlungen, die er als fahrlässig ansieht, genau vorherzubestimmen. Ein solches System wäre allerdings äußerst unflexibel gewesen, und zudem wären mit Sicherheit mit zunehmender Zeit (und technologischer Entwicklung) Konstellationen aufgetreten, die der Gesetzgeber zum Zeitpunkt der Kodifikation des Zivilrechts noch nicht vorhergesehen hatte.
- 34 Der Ausfüllung des Fahrlässigkeitsmaßstabes kann man nun unterschiedliche Erwägungen zugrunde legen – zum einen solche **ökonomischer Effizienz**, zum anderen Erwägungen der **Verteilungsgerechtigkeit**. Ökonomische Effizienzmaßstäbe zielen vor allem darauf ab, dass das Haftungs- und Deliktsrecht eine optimale Abschreckungswirkung hat. Sozialschädliche Tätigkeiten sollen durch die Abschreckungswirkung des Deliktsrechts verhindert werden, während dieses gleichzeitig keinen *chilling effect* gegenüber Tätigkeiten mit einem insgesamt positiven Nutzen erzeugen soll.
- 35 Ökonomische Effizienz und Verteilungsgerechtigkeit können dabei in einem Spannungsverhältnis zueinander stehen. Soweit es jedoch möglich ist, die ökonomisch effizienteste Regelung mit Gesichtspunkten der Verteilungsgerechtigkeit in Einklang zu bringen, sollte diese Lösung angestrebt werden. So gibt es in der ökonomischen Theorie etwa die Regel, dass derjenige das Schadensrisiko trägt, der den Schaden am kostengünstigsten vermeiden oder versichern kann. Aus der Perspektive der Verteilungsge-

rechtigkeit hat diese Regelung den Nachteil, dass eine Partei die gesamten Kosten der Schadensvermeidung trägt bzw. voll haftet, soweit sie keine Vorkehrungen zur Schadensvermeidung trifft. Soweit dies den Parteien aber bekannt ist (etwa, weil es eine gefestigte Rechtsprechung in dieser Hinsicht gibt), können sie entsprechende vertragliche Regelungen treffen, um den Vertragspartner, der sich um die Schadensvermeidung kümmert, zu kompensieren.

Ein Rechtsgebiet, in dem konkretisierungsbedürftige Konzepte eher die Regel als die Ausnahme sind, ist das Verfassungsrecht. Dies liegt daran, dass die Verfassung nur einen ausfüllungsbedürftigen Rahmen vorgibt, der auch in Zeiten gesellschaftlichen Wandels flexibel bleiben soll. So schreibt etwa Art. 20 Abs. 1 GG vor, dass der deutsche Staat demokratisch organisiert sein müsse. Außer der Tatsache, dass die Staatsgewalt vom Volke auszugehen habe (Art. 20 Abs. 2 S. 1 GG), wird im Grundgesetz nicht näher geregelt, was unter Demokratie zu verstehen ist. Mit zunehmender Übertragung von Entscheidungsgewalt auf die Exekutive oder internationale Institutionen sind diese Entwicklungen am Maßstab der Demokratie zu messen, wobei uns die simple Formel von der vom Volke ausgehenden Staatsgewalt kaum Entscheidungshilfen gibt. 36

Die deutsche Staatsrechtslehre hat zu diesem Zweck die Theorie der Legitimationskette entwickelt.<sup>27</sup> Danach muss jede Ausübung von Hoheitsgewalt auf das Parlament als dem einzigen direkt vom Volk legitimierten Organ rückführbar sein. Die Theorie entwickelt ein hierarchisches Modell, in dem das Parlament alle Staatsgewalt zum einen durch die Gesetzgebung programmiert und zum anderen durch die Verantwortlichkeit der Regierung als Spitze der Exekutive kontrolliert. Die Verwaltungsbeamten selbst unterliegen einem inhaltlichen Weisungsrecht, das sich bis zur Spitze der Exekutive zurückführen lässt. Sinn und Zweck dieses Modells soll es sein, durch die Staatsgewalt den jeweiligen „Volkswillen“ durchzusetzen. 37

Abgesehen von der Frage, ob es einen solchen Volkswillen überhaupt gibt bzw. geben kann,<sup>28</sup> hängt die Überzeugungskraft des Modells aber von mehreren empirischen Prämissen ab. So setzt sie zunächst voraus, dass der Volkswille durch das Parlament tatsächlich repräsentiert wird. Gerade hieran meldet jedoch die ökonomische *Public-Choice*-Forschung erhebliche Zweifel an.<sup>29</sup> Des Weiteren setzt die Legitimationskettentheorie vo 38

<sup>27</sup> Böckenförde, in: Isensee/Kirchhof (Hg.), Handbuch des Staatsrechts der Bundesrepublik Deutschland, Band II: Verfassungsstaat, 3. Aufl. 2004, § 24.

<sup>28</sup> Kritisch dazu etwa Bryde, Die bundesrepublikanische Volksdemokratie als Irrweg der Demokratietheorie, Staatswissenschaft & Staatspraxis 5 (1994), 305 ff.

<sup>29</sup> Hierzu unten § 6.



raus, dass parlamentarische Gesetzgebung und Weisungen innerhalb der Exekutive eine ausreichende Steuerungswirkung für Verwaltungshandeln haben.<sup>30</sup> Sollte man empirisch feststellen, dass diese Mechanismen nur eine eingeschränkte **Steuerungswirkung** haben, müsste man sich über zusätzliche Legitimationsfaktoren Gedanken machen.

- 39 Das letzte Beispiel verdeutlicht auch die Funktion, die sozialwissenschaftliche Erkenntnisse im Verfassungsrecht bei der Konkretisierung normativer Konzepte übernehmen können. Als positive Wissenschaften können sie keine normativen Theorien oder Konzepte entwickeln. Sie können jedoch dabei helfen, die impliziten und expliziten Prämissen normativer Theorien zu überprüfen und zu hinterfragen. Ein normatives Modell mag noch so elegant und innerlich kohärent sein. Wenn es keine Basis in der Wirklichkeit hat, ist es für die rechtswissenschaftliche Theoriebildung nur von geringem Nutzen.

## 2. Rechtssetzung

- 40 In der US-amerikanischen Rechtswissenschaft beziehen sich, wie erwähnt, die meisten wissenschaftlichen Beiträge nicht auf die Auslegung des geltenden Rechts, die *lex lata*, sondern vielmehr auf eine Bewertung des Rechts *de lege ferenda*, also die Reform des geltenden Rechts. Es werden bestehende Regelungen auf ihre Sinnhaftigkeit überprüft und Vorschläge für Reformen gemacht. In der deutschen Rechtswissenschaft spielt die Gesetzgebungslehre<sup>31</sup> eine deutlich geringere Rolle.<sup>32</sup> Gerade bei der Rechtssetzung ist eine **Folgenorientierung** jedoch von zentraler Bedeutung.<sup>33</sup> Der Gesetzgeber verfolgt mit einer Regelung zumeist konkrete regulatorische Ziele. Er möchte etwa den Schadstoffausstoß von Autos senken oder mittels Bankenregulierung das Risiko von Bankeninsolvenzen eindämmen. Zur Erreichung dieser Ziele stehen ihm grundsätzlich mehrere Regelungsmöglichkeiten zur Verfügung. Daher muss er eine Prognose anstellen, welches Mittel zu einer **optimalen Erreichung des Ziels** bei möglichst ge-

<sup>30</sup> Möllers, Braucht das öffentliche Recht einen neuen Methoden- und Rechtsstreit?, VerwArch 90 (1999), 187 (203).

<sup>31</sup> Vgl. etwa Karpén, Gesetzgebungslehre, 2. Aufl. 2008.

<sup>32</sup> Für ein Plädoyer für eine stärkere Konzentration der Rechtswissenschaft auf die Rechtssetzung s. Eidenmüller, Effizienz, 2015, S. 490. Eine Kritik *de lege ferenda* an einer Figur des geltenden Rechts findet sich bspw. bei Schmolke, Der Grundsatz der Nichtübertragbarkeit beschränkter persönlicher Dienstbarkeiten aus rechtsvergleichender und rechtsökonomischer Perspektive – Eine kritische Betrachtung der §§ 1092, 1090 Abs. 2 i. V. m. § 1061, AcP 208 (2008), 515 ff.

<sup>33</sup> Dazu ausführlich van Aaken, Rational Choice, 2003, S. 156 ff.

ringen Nebenwirkungen führt. Solche Prognosen sind empirischer Art und bedürfen sozialwissenschaftlicher, insbesondere ökonomischer, Methoden und Theorien, um sinnvoll durchgeführt werden zu können.

Wenn etwa bei bestimmten strafrechtlichen Delikten eine Reform des Strafmaßes diskutiert wird, dann ist das Strafmaß an den Zwecken zu orientieren, die mit dem Strafrecht erreicht werden sollen. Diese sind in erster Linie die Spezial- und die Generalprävention. Mit einer Erhöhung des Strafmaßes wird dabei vor allem eine größere Abschreckungswirkung beabsichtigt. Eine die Gesetzgebung unterstützende Wissenschaft müsste daher empirische Studien durchführen, ob und in welchem Umfang eine Erhöhung des Strafmaßes tatsächlich eine größere Abschreckungswirkung hat. 41

Ein weiteres Beispiel betrifft die Schuldrechtsmodernisierung, die Anfang 2002 in Kraft getreten ist. Im Rahmen dieser Reform ist unter anderem das Verjährungsrecht überarbeitet worden. Geht man davon aus, dass das Zivilrecht in erster Linie der effizienten **Ressourcenallokation** dient, kommt man bei der Frage, wie das Verjährungsrecht optimal auszugestalten ist, um die ökonomische Theorie nicht herum, um zu bestimmen, welche Verjährungsregeln die effizientesten sind.<sup>34</sup> Die Notwendigkeit der sozialwissenschaftlichen Absicherung von Gesetzgebung wird nicht nur von der Wissenschaft, sondern durchaus auch im politischen Prozess gesehen. So werden Gesetzgebungsvorhaben in zunehmendem Maße von Untersuchungen zur Folgenabschätzung begleitet, die teilweise explizit auf ökonomischen oder sozialwissenschaftlichen Sachverstand zurückgreifen. 42

### 3. Recht als soziales Phänomen

Ein dritter Forschungsbereich der Rechtswissenschaft beschäftigt sich schließlich mit Recht als sozialem Phänomen: Welche Auswirkung hat Recht auf die Gesellschaft? Welche Auswirkung haben Gesellschaft und Kultur auf das Recht und auf die Rechtsauslegung? Bei diesen Problemen handelt es sich um empirische Fragen, die sich nur mit sozialwissenschaftlichen Methoden bearbeiten lassen. Der rechtswissenschaftliche Methodenkanon versagt in dieser Hinsicht, ist er doch allein darauf ausgerichtet, das bestehende Recht auszulegen, nicht jedoch seine Folgen auf menschliches Verhalten vorherzusagen. Die sozialwissenschaftlichen Methoden können dabei ganz unterschiedlichen Fachrichtungen entlehnt 43

<sup>34</sup> Dazu etwa *Eidenmüller*, Zur Effizienz der Verjährungsregeln im geplanten Schuldrechtsmodernisierungsgesetz, JZ 2001, 283 ff.

werden, wobei jeweils auch eine andere Perspektive in den Blick genommen wird – die der Soziologie, der Anthropologie, der Ökonomie oder der Psychologie.

- 44 Ein Beispiel für eine derartige Forschung ist die Debatte um die Effektivität des Völkerrechts. Das Völkerrecht ist das Rechtsgebiet, das den Umgang von Staaten miteinander koordiniert. Es geht dabei sowohl um die Koordination von Verhalten – etwa die Frage, wann militärische Gewalt gegen dritte Staaten angewendet werden darf oder wie man die Diplomaten eines Gaststaates behandeln muss – als auch um die Kooperation in internationalen Organisationen, wie den Vereinten Nationen oder der Welthandelsorganisation (WTO). Spezifikum des Völkerrechts ist das Fehlen einer zentralen Sanktionsinstanz. Anders als im nationalstaatlichen Kontext gibt es keine globale Exekutive, die Rechtsverletzungen mit Sanktionen ahndet oder gerichtliche Urteile mit Zwangsgewalt durchsetzt. Ist das Völkerrecht deswegen bedeutungslos?<sup>35</sup>
- 45 Um diese Frage gibt es einen heftigen Streit. Insbesondere die ökonomische Literatur hat sich in den letzten Jahren dieses Themas angenommen und versucht, Faktoren zu identifizieren, aufgrund derer Staaten möglicherweise auch beim Fehlen von zentralen Sanktionen Anreize haben, völkerrechtliche Normen zu beachten.<sup>36</sup> Solche Anreize könnten etwa in der Gefahr dezentraler Sanktionen oder in der Schädigung des guten Rufes eines Staates liegen, die es diesem in der Zukunft schwerer macht, mit anderen Staaten zu kooperieren. Auch gibt es empirische Studien, die untersuchen, inwieweit sich beispielsweise das Unterzeichnen von menschenrechtlichen Verträgen auf die Beachtung von Menschenrechten durch einen Staat auswirkt.<sup>37</sup>
- 46 Ein weiterer, mittlerweile sehr differenzierter Forschungszweig, der das Recht aus einer Außenperspektive betrachtet, beschäftigt sich mit richterlichem Entscheidungsverhalten. So untersucht beispielsweise die Psychologie kognitive Mechanismen, die richterliches Entscheidungsverhalten beeinflussen und dabei zu verzerrten Wahrnehmungen führen können. Demgegenüber setzen sich Ökonomie und Politikwissenschaft mit Anreizen und institutionellen Zwängen auseinander, denen Richter bei ihren Ent-

<sup>35</sup> So die provokante These von *Goldsmith/Posner*, *The Limits of International Law*, 2005.

<sup>36</sup> Dazu insbesondere *Guzman*, *How International Law Works*, 2008; *Trachtman*, *The Economic Structure of International Law*, 2008.

<sup>37</sup> Siehe *Hathaway*, *Do Human Rights Treaties make a Difference?*, *Yale Law Journal* 111 (2002), 1935 ff.; *Neumayer*, *Do International Human Rights Treaties Improve Respect for Human Rights?*, *Journal of Conflict Resolution* 49 (2005), 925 ff.

scheidungen unterliegen. Die Leitfrage dieser Forschung ist dabei, wie Richter die ihnen offen stehenden Entscheidungsspielräume ausfüllen. Als mögliche Faktoren werden dabei Karriereerwägungen der einzelnen Richter,<sup>38</sup> politische Einstellungen<sup>39</sup> oder institutionelle Zwänge<sup>40</sup> untersucht. Dabei ist wichtig zu betonen, dass diese Untersuchungen nicht davon ausgehen, dass rechtliche Normen bei richterlichen Entscheidungen keine Rolle spielen. Vielmehr wird angenommen, dass sie nicht die *einzigsten* entscheidungsleitenden Faktoren sind.

#### 4. Grenzen der ökonomischen Methode in der Rechtswissenschaft

Die Anwendung der ökonomischen Methode in der Rechtswissenschaft unterliegt jedoch gewissen Grenzen, insbesondere soweit aus ökonomischen Studien normative Schlussfolgerungen gezogen werden sollen. Hier sollen insbesondere zwei Limitierungen herausgegriffen werden. Zum einen findet die Effizienzorientierung der normativen Ökonomie im Verfassungsstaat Grenzen in konkurrierenden normativen Zielen. Zum anderen ist bei jedem ökonomischen Modell und bei jeder empirischen Studie die Übertragbarkeit auf juristische Kontexte genau zu überprüfen. 47

##### a. Effizienzorientierung und Umverteilung

Wie bereits angesprochen kann die Ökonomie positiv und normativ verstanden werden. Positive Studien versuchen zu zeigen, welche Regelungsmöglichkeit unter mehreren Alternativen die effizienteste ist. Die normative Ökonomie würde daraus die Schlussfolgerung ziehen, dass die effizienteste Regelung auch immer die beste Regelung ist, da sie den kumulierten Präferenzen der Betroffenen am ehesten entspricht. Diese Annahme geht auf die Philosophie des **Utilitarismus** zurück. Nach utilitaristischer Vorstellung ist es das Ziel einer politischen Ordnung, jedem Menschen das größtmögliche Glück zuteilwerden zu lassen. Jeder Mensch soll also seine Präferenzen so weit wie möglich durchsetzen können. 48

Nach dem **Pareto-Kriterium** ist eine staatliche Maßnahme dann gerechtfertigt, wenn sie mindestens einen Menschen besser stellt, aber nie- 49

<sup>38</sup> Vgl. Posner, *How Judges Think*, 2010.

<sup>39</sup> Hierzu bspw. Segall/Cover, *Ideological Values to the Votes of U.S. Supreme Court Justices*, *American Political Science Review* 83 (1989), 557 ff.; Richards/Kritzer, *Jurisprudential Regimes in Supreme Court Decision Making*, *American Political Science Review* 96 (2002), 305 ff.

<sup>40</sup> Vgl. z.B. zum BVerfG Vanberg, *The Politics of Constitutional Review in Germany*, 2005; Petersen, *Verhältnismäßigkeit als Rationalitätskontrolle*, 2015.

mand anders schlechter gestellt wird.<sup>41</sup> Dieses Prinzip erscheint unmittelbar einsichtig – eine solche Maßnahme wird selten auf Widerstand stoßen. Das Problem ist, dass es in der Praxis selten Maßnahmen geben wird, von denen niemand einen Nachteil erleidet. Selbst staatliche Subventionen, die einer bestimmten Gruppe zuteilwerden, verbrauchen öffentliche Ressourcen und nehmen damit jeden einzelnen Steuerzahler in Anspruch. Aus diesem Grund entwickelten die Ökonomen *Nicholas Kaldor* und *John Hicks* ein Prinzip, wonach Maßnahmen gerechtfertigt sind, wenn diejenigen, die Vorteile aus dieser Maßnahme ziehen, theoretisch diejenigen kompensieren könnten, die einen Nachteil erleiden.<sup>42</sup> Der Saldo aus den positiven und den negativen Effekten einer Maßnahme muss also insgesamt positiv sein. Das Problem dieses sog. *Kaldor-Hicks-Kriteriums* ist, dass die Kompensation nur theoretisch möglich sein, nicht aber tatsächlich erfolgen muss. Es wäre demnach auch eine Maßnahme gerechtfertigt, die den reichsten 20 % der Bevölkerung auf Kosten der übrigen 80 % der Bevölkerung einen Zuwachs bringt, solange dieser Zuwachs insgesamt die Kosten für den Rest der Bevölkerung übersteigt. Diese Lesart des Effizienzprinzips bringt natürlich erhebliche Verteilungsprobleme mit sich.

50 Dennoch haben die Utilitaristen versucht, die reine Effizienzorientierung zu rechtfertigen. Den wohl interessantesten Rechtfertigungsversuch finden wir in einem Gedankenspiel des ökonomischen Nobelpreisträgers *John Harsanyi*.<sup>43</sup> *Harsanyi* nimmt einen Urzustand an, in dem über die Ausgestaltung der Gesellschaft abgestimmt wird, ohne dass jemand weiß, welche Position er oder sie später in dieser Gesellschaft einnehmen wird (*Schleier des Nichtwissens*). Die Menschen wissen nicht, ob sie reich oder arm, dumm oder intelligent, schön oder hässlich sein werden. *Harsanyi* nimmt an, dass die Menschen sich bei der Ausgestaltung der Gesellschaft am *Erwartungswert* orientieren werden. Dieser Erwartungswert berechnet sich, indem man den Nutzen einer Position in der Gesellschaft mit der Wahrscheinlichkeit multipliziert, dass man diese Position tatsächlich erreicht. Eine Maßnahme, die die Reichen besser stellt, erhöht diesen Erwartungswert dann, wenn die Verbesserung für die Reichen größer ist als die Verschlechterung für die Armen.

51 Ein Beispiel mag dieses verdeutlichen. Nehmen wir eine Gesellschaft von fünf Personen an. Eine bestimmte soziale Ausgestaltung würde allen Mitgliedern der Gesellschaft ein Vermögen von jeweils 2 Einheiten zuteil-

<sup>41</sup> Hierzu Rz. 89 ff.

<sup>42</sup> Hierzu Rz. 93 ff.

<sup>43</sup> *Harsanyi*, Cardinal Utility in Welfare Economics and in the Theory of Risk-taking, *Journal of Political Economy* 61 (1953), 434 (434 f.).

werden lassen. Der Erwartungswert im Urzustand wäre also 2. Gehen wir jetzt davon aus, dass es eine Maßnahme gäbe, die es uns erlaubte, einer bestimmten Person innerhalb dieser Gesellschaft einen Zuwachs von 8 Einheiten zukommen zu lassen, wofür die übrigen Mitglieder der Gesellschaft auf jeweils eine Einheit verzichten müssten. Es gäbe anschließend also eine Person mit 10 Einheiten, während die übrigen vier Personen jeweils eine Einheit hätten. Für die Personen im Urzustand würde dadurch der Erwartungswert steigen. Jeder hätte eine 20 %ige Chance, in die Rolle der reichen Person zu schlüpfen, so dass der Erwartungswert bei 2,8 Einheiten liegen würde. Nach *Harsanyis* Modell wäre letztere Variante also vorzuziehen.

Orientieren wir uns in der Realität aber tatsächlich ausschließlich am Erwartungswert? Die empirische Evidenz spricht dagegen. So verweisen Experimentalökonominnen und Psychologen darauf, dass Menschen in der Regel risikoavers sind.<sup>44</sup> Wenn die Chance auf den Hauptgewinn gering ist, dann entscheiden sie sich lieber für die sicherere Variante mit einem geringeren Erwartungswert. Zudem haben die meisten Menschen eine Aversion gegen zu starke Ungleichheit.<sup>45</sup> Insofern geht *Harsanyis* Rechtfertigungsversuch von Prämissen – der ausschließlichen Orientierung am Erwartungswert – aus, die in der Realität zumeist nicht gegeben sind. Insofern kann die Maximierung ökonomischer Effizienz im Sinne der normativen Ökonomie weder für die Gesetzgebung noch für richterliche Entscheidungen der einzige Referenzpunkt sein. 52

In diesem Zusammenhang ist noch auf den häufig formulierten Einwand einzugehen, der dem *homo oeconomicus* zugrundeliegende Utilitarismus propagiere ein **Menschenbild**, das mit jenem des Rechts unvereinbar sei.<sup>46</sup> Dieser Vorwurf geht ins Leere. Tatsächlich beansprucht die Ökonomie nämlich gar nicht, ein alternatives Menschenbild entworfen zu haben. Sie liefert vielmehr ein **Modell menschlichen Verhaltens**. Modelle aber verfolgen, wie wir sogleich ausführlicher sehen werden, lediglich den Zweck, anhand sparsamer Annahmen Voraussagen über die Wirklichkeit zu treffen. Sie sind auf Vereinfachung zwingend angewiesen, um ihren Zweck erfüllen zu können. Gegen ein Modell menschlichen Verhaltens kann man folglich nicht einwenden, es sei ethisch nicht vertretbar oder raube dem Menschen seinen Personencharakter.<sup>47</sup> Man kann allerdings 53

<sup>44</sup> Zu Risikopräferenzen s. unten Rz. 76, zur *Prospect Theory* Rz. 529ff.

<sup>45</sup> *Magen*, Gerechtigkeit als Proprium des Rechts, 2009.

<sup>46</sup> So etwa *Fezer*, Aspekte einer Rechtskritik an der *economic analysis of law* und am *property rights approach*, JZ 1986, 817 (822).

<sup>47</sup> Zu den problematischen Folgen auch einer lediglich deskriptiven Verhaltenstheo-

bemängeln – und damit kommen wir zum nächsten Einwand –, dass seine Voraussagen in der Realität nicht zutreffen.

### b. Die Übertragbarkeit der Ergebnisse sozialwissenschaftlicher Studien auf Rechtsfragen

- 54 Sozialwissenschaftler arbeiten für Erklärungen und Prognosen in der Regel mit **Modellen**. Ähnlich wie ein Modell von einem Segelschiff oder Landkarten versuchen auch sozialwissenschaftliche Modelle das, was wir beobachten, zu repräsentieren, dabei aber gleichzeitig Komplexität zu reduzieren.<sup>48</sup> Jedes sozialwissenschaftliche Modell ist daher eine Vereinfachung der Wirklichkeit. Modelle können gar nicht alle Faktoren berücksichtigen, da die Realität zu komplex ist. Die Kunst besteht darin, die wichtigen Faktoren, die soziale Prozesse erklären, zu identifizieren und einzubinden, so dass das Modell ein strukturell möglichst ähnliches Abbild dessen zeichnet, was wir beobachten. Die Annahme, dass Menschen rational handeln, ist also dann unschädlich, wenn es zwar Abweichungen gibt, diese Abweichungen aber nicht systematisch sind, so dass zumindest im Mittel rationales Handeln beobachtet werden kann. Kann man dagegen systematische Abweichungen feststellen, müssen diese in einem Verhaltensmodell berücksichtigt werden, soll dieses nicht an Erklärungskraft verlieren.<sup>49</sup> Insofern ist für die Übernahme von Erkenntnissen aus der ökonomischen Theorie wichtig, die Plausibilität der den entsprechenden ökonomischen Modellen zugrunde gelegten **Annahmen zu überprüfen**.
- 55 Gleiches gilt auch für die **empirische Forschung**. Bei jeder empirischen Studie ist der Kontext der Studie zu berücksichtigen und zu fragen, ob die Erkenntnisse auch auf andere Konstellationen übertragen werden können. So zeigt die psychologische Forschung, dass Menschen zum Vermeiden von Verlusten ein deutlich höheres Risiko eingehen als zum Erzielen von Gewinnen. Dieses psychologische Phänomen lässt sich durchaus für die Gestaltung von Gesetzen nutzbar machen. So wird in einigen Staaten – wie etwa in Deutschland – die Einkommenssteuer an der Quelle, d. h. bei Auszahlung des Gehalts erhoben, während die in anderen Staaten – etwa

---

rie vgl. *Towfigh*, Das Parteien-Paradox. Ein Beitrag zur Bestimmung des Verhältnisses von Demokratie und Parteien, S. 151 ff.

<sup>48</sup> In der Ökonomie werden Modelle immer mit mathematischen Formeln beschrieben, was in anderen Sozialwissenschaften, wie der Soziologie oder der Politikwissenschaft, nicht notwendig der Fall sein muss.

<sup>49</sup> Zu den systematischen Abweichungen menschlichen Verhaltens vom Rationalmodell, das in der psychologischen Forschung beobachtet worden ist, s. noch unten § 8.

Frankreich – erst mit der Steuererklärung fällig wird. Die empirische Evidenz legt nahe, dass das erstere System eine geringere Steuerhinterziehungsrate zur Folge hat, da es in Deutschland bei der Steuererklärung von Arbeitnehmern zumeist nur noch um die Höhe der Rückzahlung (= Gewinn) geht, während es in Frankreich um die Höhe der noch zu zahlenden Steuer (= Verlust) geht. Die Sinnhaftigkeit dieses Vorschlags setzt jedoch voraus, dass die im Experiment in einer kontrollierten Umgebung gewonnenen Ergebnisse auf die komplexeren Fälle der Steuerhinterziehung übertragbar sind und das Phänomen nicht möglicherweise durch andere Effekte überlagert wird.

Diese beiden Warnungen sprechen mitnichten gegen die Anwendbarkeit sozialwissenschaftlicher Forschung in der Rechtswissenschaft. Sie sollen allerdings davor bewahren, sozialwissenschaftliche Erkenntnisse vorschnell aus ihrem Kontext zu lösen, zu abstrahieren und auf Situationen zu übertragen, in denen andere Grundvoraussetzungen herrschen, was im schlimmsten Fall dazu führen kann, dass die Erkenntnisse im übertragenen Kontext keine Anwendung mehr finden. 56

## V. Die relevanten Methoden der Ökonomie

In der Ökonomie gibt es verschiedene methodische Richtungen. Die wichtigste ist die **ökonomische Theorie**, auf die in diesem Buch der Schwerpunkt gelegt wird. Die ökonomische Theorie versucht, bestimmte in der Realität ablaufende Prozesse zu modellieren und zu erklären. Zudem sollen die ökonomischen Modelle Prognosen für zukünftiges Verhalten erlauben. Im vergangenen Jahrhundert ist die ökonomische Theorie zunehmend mathematisiert worden. Die Mathematik wird dabei als eine Sprache verstanden, von der man sich eine größere Exaktheit verspricht als von der verbalen Darstellung von Modellen. Begriffliche Konzepte sind immer zu einem gewissen Grade dehnbar. Sie können gleichzeitig verschiedene Phänomene bezeichnen und sind gerade an den begrifflichen Rändern oft unscharf. Die Mathematik hat hier den Vorteil, dass sie dieser sprachlichen Ambiguität nicht unterliegt. Da dieses Lehrbuch jedoch an Juristen gerichtet ist, werden wir versuchen, die Konzepte, soweit dies möglich ist, verbal darzustellen. 57

Ökonomische Modelle stellen bestimmte **tatsächliche Annahmen** auf und versuchen, aus diesen Annahmen Folgerungen für gesellschaftliche Prozesse abzuleiten. Die Überzeugungskraft eines Modells hängt also von der Robustheit der Annahmen ab. Wenn ein ökonomisches Modell bei- 58



spielsweise annimmt, dass die Beförderung von Richtern von der Qualität richterlicher Entscheidungen abhängt, und daraus bestimmte Folgerungen für die Effizienz des Justizsystems ableitet,<sup>50</sup> dann sind diese Folgerungen nur überzeugend, wenn Beförderungen tatsächlich auf der Urteilsqualität und nicht vielmehr auf der Zahl der erledigten Streitfälle oder der politischen Vernetzung des Richters beruhen.

59 An dieser Stelle kommt die **empirische Ökonomie** ins Spiel, die versucht, Annahmen über die Wirklichkeit empirisch zu testen. In der Empirie gibt es dabei zwei unterschiedliche Richtungen – zum einen die **Experimental-ökonomie**, zum anderen die **Ökonometrie**. Empirische Forschung versucht im Wesentlichen den Zusammenhang von unterschiedlichen tatsächlichen Faktoren herauszufinden. Wir beobachten, wie sich eine bestimmte Tatsache (abhängige Variable) verändert, wenn wir die Umwelt (unabhängige Variable) modifizieren. Wie wirken sich also die Qualität des Urteils, die Zahl der erledigten Streitfälle und die Vernetzung eines Richters (unabhängige Variablen) auf dessen Beförderung (abhängige Variable) aus?

60 Experimente haben den Vorteil, dass wir eine große Kontrolle über die Umwelt haben, da wir sichergehen können, dass bestimmte beobachtete Effekte nicht auf dritte Faktoren zurückzuführen sind. Gleichzeitig hat dies jedoch den Nachteil, dass die Ergebnisse experimenteller Forschung sich möglicherweise nicht auf die Realität übertragen lassen, da dort eben noch eine ganze Reihe weiterer Faktoren eine Rolle spielen können, die wir im Experiment nicht berücksichtigen konnten. Daher geht die **Ökonometrie** den umgekehrten Weg. Sie sucht sich tatsächliche Daten und versucht diese zu analysieren, muss dabei jedoch sicherstellen, dass nicht andere Faktoren für den Effekt sehr viel wichtiger sind als der untersuchte Zusammenhang. Die zu wählende Methode hängt dabei von der Fragestellung ab. Darauf werden wir näher im vorletzten Kapitel dieses Buches eingehen, in dem wir uns mit empirischen Methoden beschäftigen.

---

<sup>50</sup> So etwa *Palumbo/Sette*, Career Concerns and Excessive Signaling in Courts, Working Paper 2009.

## § 2 – Das ökonomische Paradigma

**Literatur:** A. van Aaken, „Rational Choice“ in der Rechtswissenschaft, 2003; *dies.*, Vom Nutzen der ökonomischen Theorie für das öffentliche Recht, in: M. Bungenberg *et al.* (Hg.), Recht und Ökonomik, 2004, 1–31; M. Adams, Ökonomische Theorie des Rechts, 2. Aufl. 2004; M. Allingham, Choice Theory: A Very Short Introduction, 2002; H.-D. Assmann/C. Kirchner/E. Schanze, Ökonomische Analyse des Rechts, 1993; D. G. Baird/R. H. Gertner/R. C. Picker, Game Theory and the Law, 3. Aufl. 2003; P. Behrens, Die ökonomischen Grundlagen des Rechts, 1986; R. H. Coase, The Problem of Social Cost, Journal of Law & Economics 3 (1960), 1–44; *ders.*, The Nature of the Firm, Economica 4 (1937), 386–405; R. Cooter/T. Ulen, Law and Economics, 6. Aufl. 2012; A. Hindmoor/B. Taylor, Rational Choice, 2. Aufl. 2015; H. E. Jackson/L. Kaplow/S. M. Shavell/W. K. Viscusi/D. Cope, Analytical Methods for Lawyers, 2003; G. Jansson, Ökonomische Theorie im Recht, 2004; G. Kirchgässner, Homo Oeconomicus: Das ökonomische Modell individuellen Verhaltens und seine Anwendung in den Wirtschafts- und Sozialwissenschaften, 4. Aufl. 2014; C. Kirchner, Ökonomische Theorie des Rechts, 1997; K. Mathis, Effizienz statt Gerechtigkeit?, 3. Aufl. 2009; J. Noll, Rechtsökonomie, 2005, 13–29; R. Posner, Economic Analysis of Law, 9. Aufl. 2014; R. Richter/E. Furubotn, Neue Institutionenökonomik, 4. Aufl. 2010; M. Rodi, Ökonomische Analyse des Öffentlichen Rechts, 2014; H. Schäfer/C. Ott, Lehrbuch der ökonomischen Analyse des Zivilrechts, 5. Aufl. 2012; W. Weigel, Rechtsökonomik, 2003.

### I. Theoretische Grundannahmen

Das Charakteristikum der Ökonomik, die sich ja nicht mehr über ihren Forschungsgegenstand definieren kann, ist ihr Forschungsansatz, das heißt eine Reihe von grundlegenden Ideen, Konzepten und Annahmen, die als **ökonomisches Paradigma** bezeichnet werden, und von denen auch die ökonomische Theorie des Rechts ausgeht. Zunächst einmal nehmen Ökonomen bei ihren Untersuchungen, aus einer methodischen Perspektive, individuelles Verhalten in den Blick – und beispielsweise nicht Systeme (die zwar aus Individuen bestehen mögen, bei denen individuelles Verhalten aber nicht der Hauptuntersuchungsgegenstand ist) oder neuro-kognitive

61

Mechanismen (also jene Prozesse, die individuelles Verhalten erzeugen). Diese Perspektive wird **methodologischer Individualismus** genannt. Außerdem gehen Ökonomen davon aus, dass zwar die menschlichen Bedürfnisse und Wünsche unbegrenzt sind, die dafür zur Verfügung stehenden Ressourcen allerdings knapp – aufgrund dieser **Ressourcenknappheit** sind wir gezwungen, Entscheidungen zu treffen. Wie treffen wir nun unsere Wahl, wie lösen wir Entscheidungsprobleme? Hier kommen das Eigennutztheorem und die Rationalitätsannahme ins Spiel: Ökonomen nehmen ausgehend von methodologischem Individualismus und von der Ressourcenknappheit an, dass individuelle Akteure unter mehreren Optionen jene wählen, die ihren individuellen Nutzen maximiert. Die beiden letztgenannten Elemente des Paradigmas sind zum Verhaltensmodell des *homo oeconomicus* verdichtet worden, das vielfach Kritik und zwischenzeitlich einige Modifikationen und Präzisierungen erfahren hat,<sup>1</sup> das aber nach wie vor auch im Hinblick auf die empirische Fundierung und mangels ganzheitlicher theoretischer Alternativen Grundlage auch der ökonomischen Analyse des Rechts ist.

- 62 Die gemeinsamen Bemühungen von Juristen und Ökonomen um Erkenntnisgewinn haben eine wechselhafte Geschichte. In den Universitäten waren beide Disziplinen oft zu einer „Staatswissenschaftlichen Fakultät“ zusammengefasst. Dennoch gab es vor allem in der ersten Hälfte des 20. Jahrhunderts kaum gleichgerichtete Forschung. Das hing vor allem damit zusammen, dass die Nationalökonomie ihren Blick immer stärker auf die Verteilung von Gütern konzentrierte. Gleichzeitig wurden die ökonomischen Methoden immer exakter und der Rückgriff auf mathematische Ausdrucksformen immer stärker. Dies erschwerte zum einen die Rezeption ökonomischer Erkenntnisse durch andere Disziplinen, verringerte zugleich aber auch die Relevanz ökonomischer Forschung für konkrete wirtschaftspolitische Forderungen.

### 1. Methodologischer Individualismus

- 63 Die erste Grundannahme der Ökonomik ist der **methodologische Individualismus**, demzufolge allein die Handlungen von Individuen Gegenstand der wissenschaftlichen Analyse sind. Kollektiventscheidungen – etwa durch Unternehmen oder Staaten – folgen danach nicht der Eigenlogik eines „Kollektivwillens“, sondern können auf das Zusammenwirken individueller Entscheidungsträger zurückgeführt und durch dieses erklärt wer-

---

<sup>1</sup> Dazu insbesondere § 8.

den.<sup>2</sup> Trotz der Annahme des methodischen Individualismus stellt jedenfalls die positive ökonomische Theorie nicht auf das Verhalten des Einzelnen in dem Sinne ab, dass sie beanspruchen würde, das Verhalten jedes einzelnen konkreten Akteurs erklären oder vorhersagen zu können. Sie stellt vielmehr auf im Aggregat identifizierbare Verhaltensmuster ab, also auf eine Art „durchschnittlichen“ Verhaltens einer größeren Zahl der gleichen Entscheidungssituation ausgesetzter Akteure, auf einen durch die Aggregation vertypen Normmenschen. Die Rechtswissenschaft kennt einen vergleichbaren methodologischen Ansatz etwa bei der Rechtsfolgenanalyse, die gerichtliche Entscheidungspraxis hinsichtlich ihrer Folgen nicht nur auf den untersuchten Einzelfall, sondern hinsichtlich der Gesamtheit der gleichgelagerten Fälle beurteilt. Damit lassen sich Hypothesen hinsichtlich der bewirkten tatsächlichen Rechtsänderung aufstellen.<sup>3</sup>

## 2. Ressourcenknappheit

Die fundamentale wirtschaftswissenschaftliche Annahme knapper Güter schlägt sich im ökonomischen Paradigma bei der Betrachtung des Verhältnisses der Bedürfnisse zu den für ihre Befriedigung verfügbaren Mitteln nieder. Während die menschlichen Bedürfnisse grundsätzlich unbegrenzt sind, sind die zur Verfügung stehenden Mittel prinzipiell begrenzt. Dabei muss es sich bei Bedürfnissen und Mitteln nicht um materielle Güter handeln, auch immaterielle Güter wie Sicherheit, Wissen oder die Rechtsordnung können im ökonomischen Sinne „knapp“ sein. Zur Befriedigung der Bedürfnisse sind wegen der Ressourcenknappheit menschliche Wahlhandlungen erforderlich: Aus der ökonomischen wird damit eine **entscheidungstheoretische Betrachtung**.<sup>4</sup> Das Individuum, welches im Zentrum dieser Wahlhandlung steht, wird bei dieser seiner Entscheidung durch seine **Präferenzen** geleitet und durch **Restriktionen** beschränkt. 64

<sup>2</sup> Vom methodologischen Individualismus zu unterscheiden ist der normative Individualismus, nach dem über den methodologischen Individualismus hinaus auch normativ gelten soll, dass alles staatliche Handeln nur unter alleiniger Berücksichtigung individueller Interessen erfolgen darf, andere Präferenzen aber keine Rolle spielen dürfen. Die individuelle Wertentscheidung ist die einzig normativ akzeptierte Währung, daneben soll keine weitere Instanz berücksichtigt werden. Die Annahme des normativen Individualismus ist hoch umstritten und kein notwendiger Bestandteil des ökonomischen Paradigmas. Vgl. *Kirchner*, *Ökonomische Theorie*, 1997, S. 21.

<sup>3</sup> *Kirchner*, *Ökonomische Theorie*, 1997, S. 19 m. w. N.

<sup>4</sup> *Janson*, *Ökonomische Theorie*, 2004, S. 26.

### a. Präferenzen

- 65 Präferenzen beschreiben innere Motive des Akteurs. Sie sind unabhängig von den aktuellen Handlungsmöglichkeiten und enthalten die Wertvorstellungen des Individuums. So kann auf ein Alltagsbeispiel heruntergebrochen eine Präferenz für Schokoladeneis gegenüber Vanilleeis und wiederum von Vanilleeis gegenüber Erdbeereis bestehen:

|                |   |            |
|----------------|---|------------|
| Schokoladeneis | > | Vanilleeis |
| Vanilleeis     | > | Erdbeereis |

Die vorhandenen Optionen können **transitiv** geordnet werden, so dass aus den beiden Beispielsrelationen gefolgert werden kann, dass

|                |   |            |
|----------------|---|------------|
| Schokoladeneis | > | Erdbeereis |
|----------------|---|------------|

Dabei sind die Präferenzen **ordinal** geordnet (das heißt, sie stehen in einer fixen Folge),

1. Schokoladeneis
2. Vanilleeis
3. Erdbeereis

so dass der **Präferenzordnung** keine Aussage über den „Abstand“ zwischen zwei Optionen oder einen quantifizierbaren „Wert“ der Optionen entnommen werden kann. Es kann mithin etwa nicht gesagt werden, man möge Schokoladeneis „doppelt so gern“ wie Vanilleeis. Allerdings kann eine Person hinsichtlich zweier Optionen **indifferent** sein, also beide Optionen gleichermaßen mögen. Ferner wird angenommen, dass die Akteure hinsichtlich jeder Option eine Präferenz angeben können, ihre Präferenzordnung mithin vollständig ist.

- 66 Zwei weitere wichtige Annahmen finden sich in ökonomischen Modellen: Zum einen, dass Präferenzen **inkommensurabel** (d.h. interpersonal oder intersubjektiv nicht vergleichbar) seien, dass in unserem Beispiel also keine Akteurin behaupten kann, sie möge Vanilleeis mehr als ein anderer Akteur – denn Präferenzen enthalten immer auch Werturteile und sind daher subjektiv. Zum anderen wird in der Regel angenommen, dass Präferenzen **konstant** sind (jedenfalls über den im jeweiligen Modell beobachteten Zeitraum hinweg) und sich allenfalls über sehr lange Zeiträume verändern. Das führt auch dazu, dass Präferenzen zwar für die Erklärung konkreten menschlichen Verhaltens eine Rolle spielen, aber nicht in erster Linie als Erklärung für Verhaltensänderungen bei Akteuren herangezogen werden. Aus der Perspektive der Rechtsökonomik ist dies auch deshalb

sinnvoll, weil es wenig Raum für Interventionen des Rechts gäbe, wenn Verhaltensänderungen ausschließlich oder überwiegend auf eine Veränderung der Präferenzen zurückzuführen wären.<sup>5</sup>

## b. Restriktionen und Anreize

Restriktionen beschreiben die äußeren Begebenheiten, die der Akteur in seiner Umwelt vorfindet und die seinen Handlungsspielraum begrenzen. Diese Begrenzung erfolgt nicht allein durch die Knappheit der Ressourcen des Einzelnen, auch das Verhalten anderer Individuen oder institutionelle, informelle, zeitliche und informationelle Beschränkungen stecken den Handlungsrahmen des Einzelnen ab.<sup>6</sup> Dabei ist der Begriff der Restriktionen weit zu verstehen: Jede die Knappheit betrachteter Güter – z. B. Geld, Zeit, Sicherheit – betreffende Verteuerung verschärft die Restriktionen, jede Vergünstigung lockert sie (die Lockerung der Restriktionen wird gemeinhin als „Anreiz“ bezeichnet): Beispielsweise könnte die Etablierung einer Institution „Schokoladeneissteuer“ eine Restriktion darstellen, da sie das durch ein mageres Gehalt auferlegte Knappheits-Problem (monetäres Budget) verschärft, ohne das Budget selbst weiter zu verknappen; eine Reduzierung des Preises für Erdbeereis erhöht möglicherweise den Anreiz, diese Sorte zu wählen. Restriktionen und Anreize können eine Entscheidung des Akteurs also „verteuern“ oder „vergünstigen“. So können auch das Recht oder Sitten und Gebräuche Restriktionen im ökonomischen Sinne darstellen (wie etwa die beispielhaft erwähnte Schokoladeneissteuer). Besonders eindrücklich verteuert auch das Strafrecht Entscheidungsoptionen.<sup>7</sup> 67

In der Vorstellungswelt der Ökonomen sind nur die Restriktionen veränderlich (während, wie gesagt, die Präferenzen stabil sind). Das hängt in erster Linie damit zusammen, dass sie leichter zu ermitteln und zu beeinflussen sind. Wenn Ökonomen also überlegen, wie menschliches Verhalten in eine gewünschte Richtung beeinflusst werden kann, setzen sie bei einer Änderung der Restriktionen an, nicht aber bei einer Änderung der Präferenzen. Umgekehrt werden deshalb auch realiter beobachtete Verhaltensänderungen auf Veränderungen der Restriktionen zurückgeführt. Wenn etwa beobachtet wird, dass der Konsum von Speiseeis in einer Population 68

<sup>5</sup> Vgl. dazu aber *Lüdemann*, *Edukatorisches Staatshandeln. Steuerungstheorie und Verfassungsrecht am Beispiel der staatlichen Förderung von Abfallmoral*, 2004.

<sup>6</sup> *Van Aaken*, *Vom Nutzen der ökonomischen Theorie für das öffentliche Recht: Methode und Anwendungsmöglichkeiten*, in: *Bungenberg et al.* (Hg.), *Recht und Ökonomik*, 1 (5); *Mathis*, *Effizienz*, 2009, S. 27.

<sup>7</sup> Dazu auch unten § 8 IV.

zurückgeht, wird der Ökonom nach erhöhten Kosten (monetär oder auch bekanntgewordene Gesundheitsrisiken, die den Eiskonsum „verteuern“) Ausschau halten, sich aber nicht fragen, ob sich die zugrundeliegende Präferenzordnung geändert hat.

### 3. Verhaltensmodell des *homo oeconomicus*

- 69 Das Verhaltensmodell des *homo oeconomicus* besteht aus zwei Annahmen: Einerseits, dass Akteure ihre Entscheidungsoptionen nach deren Nutzen beurteilen (**Eigennutztheorem**), andererseits, dass sie immer die Option wählen, die ihnen den höheren Nutzen verschafft (**Rationalitätsannahme**).

#### a. Eigennutztheorem

- 70 Die ökonomische Theorie nimmt an, dass die Akteure die ihnen offen stehenden Optionen stets anhand ihres **Nutzens** bewerten. In der populärwissenschaftlichen Literatur wird dieses Wort oft mit der großen Vokabel „Glück“ übersetzt. Ob diese Übersetzung philosophisch zutreffend ist, soll hier dahinstehen, sie zeigt aber, dass prinzipiell jedwede Art von Nutzen (monetär, zeitlich, emotional, ethisch, geschmacklich usw.) Berücksichtigung finden kann. Vereinzelte Stimmen in der Rechtsökonomie vertreten, jeder Nutzen müsse sich in einer *pekuniären* Währung ausdrücken lassen;<sup>8</sup> sie finden heute jedoch kaum mehr Anhänger, auch wegen des oben geschilderten paradigmatischen Wandels von Ökonomie zu Ökonomik und der damit einhergehenden Erweiterung der Betrachtungsgegenstände. Eine am Eigennutz orientierte Entscheidung muss folglich auch nicht zwangsläufig naiv-egoistisch sein und zu Lasten des Wohles anderer gehen. Der Eigennutz kann sich vielmehr beispielsweise auch in altruistischem Verhalten äußern, etwa wenn dieses Verhalten für den Entscheidenden einen hohen ethischen Wert hat und daher eine altruistisch orientierte Entscheidung seinen individuellen Nutzen erhöht. Die Bewertung der Handlungsalternativen anhand des durch sie verschafften Nutzens beinhaltet also **keine moralische Dimension**, sondern trifft nur die wertneutrale Aussage, dass sich die Akteure gemäß ihrer Präferenzen verhalten und selbständig beurteilen, was in diesem Sinne „gut“ für sie ist.
- 71 Dem **Eigennutztheorem** zufolge haben alle Akteure eine subjektive, „innere“ Nutzenfunktion, die letztlich jeder Option einen eindeutigen Nutzen

<sup>8</sup> Vor allem etwa von *Posner*, vgl. etwa *ders.*, Utilitarianism, Economics, and Legal Theory, *Journal of Legal Studies* 8 (1979), 103 ff.

zuweist. Da der Nutzen wie geschildert nicht monetär sein muss, ist auch die Auszahlung nicht auf eine pekuniäre Währung festgelegt, sondern kann je nach Akteur in einer anderen oder auch in mehreren Währungen gleichzeitig Ausdruck finden. Bei Individuen geht die ökonomische Theorie im Regelfall davon aus, dass sich die Auszahlung in der **Zahlungsbe-reitschaft** ausdrückt: Bin ich bereit, für den großen Eisbecher € 10,- zu zahlen, so sei dies auch meine Auszahlung und damit der Nutzen, den ich aus dem Eisbecher ziehe. Bei Firmen setzt die Ökonomik auf gleiche Weise die Auszahlung mit dem Gewinn gleich. Die Nutzenfunktion lässt sich mathematisch formulieren – dass man die Nutzenfunktion (und damit die Nutzungsoptimierungsaufgabe, dazu sogleich) mathematisch präzise formulieren kann, darf jedoch nicht darüber hinwegtäuschen, dass uns freilich selbst die eigene konkrete Nutzenfunktion völlig unbekannt ist. Es handelt sich also wieder lediglich um eine Annahme oder Fiktion; und da wir mit dieser Fiktion gut arbeiten können und sie brauchbare Ergebnisse liefert, geben wir vor, dies sei der Mechanismus, mit dessen Hilfe Menschen ihre Entscheidungen treffen (**Als-ob-Annahme**<sup>9</sup>).

Ferner geht die Ökonomik in ihrem Grundmodell davon aus, dass unter der Bedingung, dass alles andere gleich bleibt (**Ceteris-paribus-Annahme**), eine zusätzliche Einheit eines Gutes grundsätzlich positiv bewertet wird: Zwei Kugeln Eis sind besser als eine. Man spricht von einem grundsätzlich **positiven Grenznutzen**, und wir nehmen diesen unabhängig davon an, wie viele Einheiten eines Gutes wir bereits besitzen. Das bedeutet, dass wir die Möglichkeit eines negativen Grenznutzens in aller Regel ignorieren – beispielsweise, dass die elfte Kugel Eis ein gewisses körperliches Unwohlsein auslösen könnte (wir könnten sie ja verkaufen, statt sie selbst zu konsumieren); außerdem operiert die Ökonomie meist mit pekuniären Einheiten und es ist selten schädlich, einen Euro mehr zu haben.

Allerdings kann der Nutzen eines Gutes von der Menge abhängen, die wir konsumieren können, selbst wenn wir keinen negativen Nutzen erleiden. So ist der Nutzen der ersten Kugel Eis am größten, während bei der zweiten und dritten Kugel der zusätzliche Nutzen kleiner wird (auch wenn der absolute Nutzen noch steigt). Und es macht einen großen Unterschied, ob Sie € 1.000 im Monat verdienen oder € 2.000, während derselbe Unterschied eher unerheblich erscheinen mag, wenn Sie in derselben Zeit € 1.001.000 oder € 1.002.000 verdienen. In vielen Modellen wird daher von einem **abnehmenden Grenznutzen** ausgegangen, was besagt, dass der

<sup>9</sup> Dazu grundlegend und mit illustrativen Erklärungen *Friedman*, *The Methodology of Positive Economics*, 1953, S. 10ff.



zusätzliche Nutzen einer zusätzlich konsumierten Einheit mit fortlaufendem Konsum abnimmt.

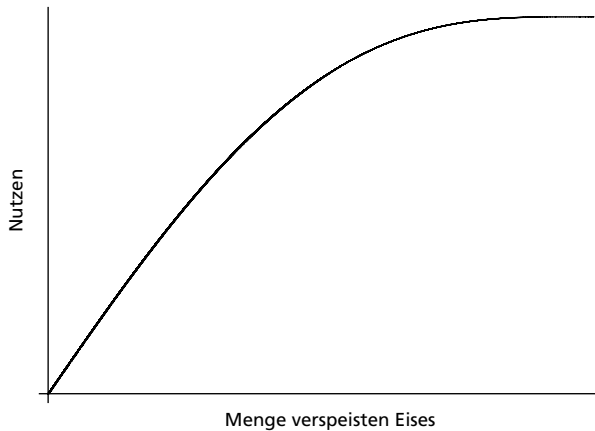


Abbildung 2.1: Abnehmender Grenznutzen

- 74 Nun tritt nicht jeder Nutzen mit Sicherheit ein. Der Nutzen des Schokoladeneises materialisiert sich nur, wenn es wohlschmeckend ist – eine Eigenschaft, die man dem Eis nicht unbedingt ansieht. Bei derlei Unsicherheit kann sich der Entscheider aber mit dem sog. **Erwartungsnutzen** (auch: **erwarteter Nutzen**) behelfen, dem rechnerischen Produkt von Nutzen und Eintrittswahrscheinlichkeit: Wenn ich bei einem Münzwurf-Glücksspiel mit 50 % Wahrscheinlichkeit € 2 gewinnen kann und mit 50 % Wahrscheinlichkeit € 0, dann ist der Erwartungsnutzen des Spiels € 1 (vorausgesetzt, dass für die Gewinnchance kein Entgelt zu bezahlen ist).
- 75 Das Konzept des Erwartungsnutzens hat einen weiteren nützlichen Effekt: Man kann mit seiner Hilfe auch bei einer ordinalen Präferenzordnung zu (**kardinalen**) Aussagen über die Abstände zwischen den Optionen gelangen. Etwa für unser oft bemühtes Eis-Beispiel indem man fragt, ob der Entscheider es vorzieht, sicher – also mit 100 % Wahrscheinlichkeit – Vanilleeis (2. Wahl) zu bekommen oder ob er es stattdessen vorzieht, mit 50 % Wahrscheinlichkeit Schokolade (1. Wahl) und mit 50 % Wahrscheinlichkeit Erdbeere (3. Wahl) zu bekommen. Durch Veränderungen dieser „Lotterie“ und der Eintrittswahrscheinlichkeiten lassen sich Abwägungen (*trade-offs*) zwischen mittleren, besten und schlechtesten Auszahlungen ablesen.

Die Krümmung der Nutzenfunktion (das heißt die Art des Grenznutzens) erlaubt uns zusammen mit diesen Lotterien ferner, **Risikopräferenzen** in unserem Verhaltensmodell zu berücksichtigen. Nehmen Sie an, es stünden folgende Wetten zur Wahl: Zu gewinnen wären einerseits entweder € 100 oder € 0 mit je 50 % Wahrscheinlichkeit („echtes Spiel“) oder andererseits eine Auszahlung von € 50 mit 100 % Wahrscheinlichkeit („sichere Wette“, *safe bet*). **Risikoneutrale** Individuen wären indifferent hinsichtlich beider Optionen, da sie denselben Erwartungsnutzen versprechen; ihre Nutzenfunktion verläuft linear. Individuen, die die „sichere Wette“ wählen, können als **risikoavers** bezeichnet werden; sie erfahren einen abnehmenden Grenznutzen, ihre Nutzenfunktion ist konkav (wie in Abbildung 2.1). Schließlich können Individuen, die das echte Spiel wählen, als **risikosuchend** angesehen werden; ihre Nutzenfunktion ist konvex und weist damit einen steigenden Grenznutzen aus. Die meisten ökonomischen Modelle operieren mit Risikoaversion. 76

## b. Rationalitätsannahme

Durch die theoretische Formulierbarkeit der individuellen Nutzenfunktion kann auch die Wahlentscheidung als analytische Optimierungsaufgabe ausgedrückt werden: Die letztendliche Wahl einer Option hängt davon ab, welche Entscheidungsalternative dem Entscheidenden den größten individuellen Nutzen bringt, wo, mathematisch gesprochen, das Maximum der Funktion liegt. Die **Rationalitätsannahme** – der Schlussstein des ökonomischen Paradigmas – besagt, dass Menschen rational handeln, also stets diejenige Option wählen, die ihnen den größten Nutzen verschafft. 77

Um dergestalt rational handeln zu können, muss das Individuum auf Änderungen der Umweltbedingungen reagieren: Die Entscheidung ist dann von den (konstanten) Präferenzen des Akteurs abhängig, welche zusammentreffen mit den veränderlichen (variablen) Restriktionen, und die er mit Hilfe seiner individuellen Nutzenfunktion bewertet. Das setzt prinzipiell voraus, dass der Akteur über alle erforderlichen Informationen zu allen Handlungsalternativen und deren Nutzen verfügt, d.h. **vollständig informiert** ist. Diese, in älteren verhaltensökonomischen Modellen oftmals formulierte, Annahme wird heute jedenfalls nicht mehr allgemein postuliert.<sup>10</sup> Die Realität zeigt, dass dem Akteur in der Regel vollständige Informationen gerade nicht zur Verfügung stehen, so dass die Entscheidung unter **Unsicherheit** und **Risiko** getroffen werden muss. Die Lockerung der Annahme vollständiger Information hat einen Preis: Sie kann dazu führen, 78

<sup>10</sup> Hierzu Rz. 246 ff.

dass die Rationalität der Entscheidung nicht mehr ‚objektiv‘ messbar ist, sondern sich allein aus der Sicht des handelnden Akteurs bewerten lässt. Eine vollständige Informationsbeschaffung wäre unter Umständen mit unverhältnismäßigen Kosten verbunden, die wiederum mit dem zu erwartenden Nutzen abgewogen werden müssten.

- 79 Die Rationalitätsannahme als ökonomisches Verhaltensmodell sieht sich inzwischen einer gewissen Skepsis ausgesetzt, da eine Reihe von vorwiegend verhaltenswissenschaftlichen Experimenten nachweisen konnte, dass Akteure in bestimmten Situationen gerade nicht rational handeln und nicht als rationale Nutzenmaximierer im o.g. Sinne agieren.<sup>11</sup> Diesen Einsichten und neuen Fragen geht § 8 (*Verhaltensökonomik*) in diesem Buch genauer nach. Dennoch behält die Rationalitätsannahme ihre Berechtigung, solange berücksichtigt bleibt, dass unter bestimmten Umständen differenzierte(re) Annahmen erforderlich sind,<sup>12</sup> zumal die Ökonomik inzwischen den Abweichungen der Akteure vom streng rationalen Verhalten und den sich daraus ergebenden Anomalien Raum gibt und diese eigener Forschung unterwirft.<sup>13</sup> Das ökonomische Paradigma, so wie es hier vorgestellt wird, bildet das Fundament einer **ökonomischen Entscheidungstheorie** (*Rational Choice* Theorie), deren spezielle Anwendungen (etwa in der Mikroökonomie oder in der Neuen Politischen Ökonomie) und Weiterentwicklungen (beispielsweise von der Entscheidungs- zur Spieltheorie oder zur Verhaltensökonomie) wir in den folgenden Kapiteln eingehender beleuchten wollen.

#### 4. Grenzen des Modells

- 80 Der *homo oeconomicus*, gleichsam das „Musterkind“ der Ökonomik, setzt die dargestellten Verhaltensannahmen vollendet um – und zwar in jedem Kontext. Ökonomen lassen ihn in die verschiedensten Rollen schlüpfen, ohne dass sich dabei seine Persönlichkeit verändern würde; so begegnet uns der *homo oeconomicus* als Konsument oder Produzent, als Arbeitgeber oder Arbeitnehmer, als Beamter, Unternehmer, Wähler. Diese Universalität macht das **ökonomische Verhaltensmodell** auch für Juristen überaus interessant. Man hat es mit einigem Fug und Recht als den wichtigsten Einzelbeitrag der *Law & Economics*-Bewegung zur Rechtswissen-

<sup>11</sup> Dazu sogleich Rz. 81 ff.

<sup>12</sup> *Kirchner*, Ökonomische Theorie, 1997, S. 14; *van Aaken*, *Rational Choice*, 2003, S. 84 f.

<sup>13</sup> Vgl. hierzu ausführlich *Janson*, Ökonomische Theorie, 2004, S.43 f.; *van Aaken*, *Rational Choice*, 2003, S. 82.

schaft bezeichnet.<sup>14</sup> Freilich wird es selten helfen, klassisch dogmatische Fragen zu klären. Sobald die Rechtswissenschaft aber ihre Binnenwahrnehmung aufgibt und zur angewandten Sozialwissenschaft wird, sieht sie sich zwangsläufig mit der Frage nach dem erwartbaren Verhalten der Normadressaten, aber auch des Normanwenders konfrontiert. *Rational Choice* dient also der Rechtswissenschaft in ihrer Eigenschaft als **Steu- rungswissenschaft**. Dennoch hat das Modell auch Grenzen, um die es in diesem Lehrbuch ebenfalls gehen soll. Eine der großen Leistungen des ökonomischen Paradigmas ist, dass die damit verbundene positive Theorie es erlaubt, empirisch überprüfbare Hypothesen abzuleiten. Und die Beschreibung der Grenzen des ökonomischen Paradigmas beginnt mit der einfachen Frage: Verhalten sich Menschen denn nun, wie es das ökonomische Verhaltensmodell vorhersagt?

### a. Empirische Herausforderungen

Der Befund zu dieser Frage ist nicht ganz eindeutig: Einerseits haben em- 81  
pirische und vor allem experimentelle Studien zeigen können, dass das Verhalten beobachteter Studienteilnehmerinnen und Studienteilnehmer in bestimmten Situationen – etwa auf Märkten oder in Auktionen, oder allgemeiner in Wettbewerbsumgebungen – in der Tat vom *Homo-oeconomicus*-Modell zutreffend vorhergesagt wurde. Das heißt freilich nicht, dass sich jedes Individuum so verhalten hat, wie die Theorie es beschreibt, sondern lediglich, dass keine **systematischen Abweichungen** zu verzeichnen waren: Vereinfacht gesprochen heben sich die Abweichungen von den Modellvorhersagen gegenseitig auf. Damit können wir sagen, dass das beobachtete Verhalten in diesen Situationen im Durchschnitt dem der Theorie nach zu erwartenden Verhalten recht genau entspricht.

Andererseits haben vor allem Psychologen ganze Batterien von Verhal- 82  
tensexperimenten durchgeführt, die zeigen konnten, dass in vielen anderen Bereichen menschliches Verhalten in wesentlichen Aspekten von den theoretischen Vorhersagen abweicht. Wir werden uns diese Abweichungen in § 8 genauer anschauen, aber um Ihnen schon jetzt ein Gefühl zu vermitteln, wollen wir uns hier zwei Beispiele anschauen:

– Im **Ultimatum-Spiel**<sup>15</sup> kann eine Spielerin entscheiden, welchen Teil eines Geldbetrags sie (Sender) und ihr Mitspieler (Empfänger) jeweils be-

<sup>14</sup> Ulen, *Firmly Grounded: Economics in the Future of the Law*, Wisconsin Law Review 3 (1997), 433 ff.

<sup>15</sup> Güth/Schmittberger/Schwarze, *An Experimental Analysis of Ultimatum Bargaining*, Journal of Economic Behavior and Organization 3 (1982), 367 ff.

kommen, sofern der Mitspieler dieser Aufteilung zustimmt (*Take-it-or-leave-it-Angebot*); stimmt der Empfänger nicht zu, ist der gesamte Betrag verloren. Die Versuchspersonen auf der Empfängerseite lehnen die Angebote durchweg ab, wenn sie unter einem Drittel der Gesamtsumme liegen, während die ökonomische Theorie vorhersagen würde, dass sie sich auch mit dem kleinsten Betrag zufriedenstellen müssten, weil selbst dieser immer noch besser ist als nichts. Sie zeigen damit eine Präferenz für nicht-strategische Bestrafungen der Sender. Spiegelbildlich bieten die Sender den Empfängern durchweg einen viel höheren Anteil am zu verteilenden Geldbetrag, an als die klassische ökonomische Theorie vorhersagen würde: Die durchschnittliche Offerte ist eine Aufteilung im Verhältnis 63:37, das häufigste Angebot war eine hälftige Teilung (möglicherweise, weil sie die nicht-rationalen Präferenzen der Empfänger vorhersehen). Selbst wenn die Beträge, um die gespielt wird, einem Monatseinkommen entsprechen, verändert sich das Verhalten der Spieler nicht. Offenbar empfinden also die allermeisten Angebotsempfänger das Bedürfnis, ein knausriges Gegenüber zu bestrafen, selbst wenn dies für sie selbst Nachteile bedeutete. Beachten Sie: Es ist das Verhalten des Empfängers, welches das Eigeninteresse-Postulat verletzt. Die Sender dagegen verhalten sich rational, wenn sie im Bewusstsein dessen ein großzügiges Angebot abgeben, um sich wenigstens ein Stück des „Kuchens“ zu erhalten.

– In *Framing-Experimenten* konnte gezeigt werden, dass die bloße Beschreibung bzw. der Kontext einer Situation – der nach der ökonomischen Standard-Theorie unmaßgeblich sein müsste – das Verhalten der Versuchspersonen fundamental beeinflusst:<sup>16</sup> Wenn eine Situation als „Gemeinschafts-Spiel“ bezeichnet wird, kooperieren die Teilnehmerinnen und Teilnehmer, das heißt, sie verhalten sich weniger rational im Sinne der ökonomischen Theorie, nehmen dadurch aber mehr Geld mit nach Hause, als wenn alle rational agiert hätten. Bezeichnet man dieselbe Situation dagegen als „Wall-Street-Spiel“, dann bricht die Kooperation zusammen und eigennütziges Verhalten wird zur Regel.<sup>17</sup> Die Versuchspersonen handeln dann zwar (häufiger) im Einklang mit den Vorhersagen der ökonomischen Theorie – rationaler in diesem Sinne –, aber sie verdienen weniger.

<sup>16</sup> Vgl. zum *Framing* Rz. 535 ff.

<sup>17</sup> *Liberman/Samuels/Ross*, The Name of the Game: Predictive Power of Reputations versus Situational Labels in Determining Prisoner's Dilemma Game Moves, *Personality and Social Psychology Bulletin* 30 (2004), 1175 ff.

## b. Urteilsverzerrungen und nicht-rationales Verhalten

Die empirisch beobachteten, systematischen Abweichungen von der Rationaltheorie werden als **Urteilsverzerrungen** oder *biases* bezeichnet; manche Autoren sprechen von „irrationalen Verhalten“. Bisweilen ist diese Qualifikation gerechtfertigt, weil menschliche Entscheidungsprozesse unter typischen und systematischen Fehlern leiden. Kognitionspsychologen und Neurowissenschaftler haben diese Fehlleistungen in einigen Fällen mit dem kognitiven Apparat in Verbindung bringen können: So wie unser Hirn also Fehler macht, die mit optischen Täuschungen sichtbar gemacht oder ausgenutzt werden können, so gibt es auch typische Fehler, denen wir unterliegen, wenn wir Entscheidungen treffen. 83

Trotzdem sollten wir vorsichtig sein, solche Fehler als „irrational“ zu bezeichnen, nur weil sie zu Entscheidungen führen, die nicht den Vorhersagen der Rationaltheorie entsprechen. Wenn wir die empirischen Ergebnisse des Ultimatum-Spiels betrachten, könnten Fairness- oder Verteilungserwägungen eine wichtige Rolle spielen: Vielleicht sind wir **ungleichheitsavers**, und eine gleichmäßige Aufteilung des Geldes, das uns zufällig in den Schoß fiel, scheint uns angemessener als jene, die der *homo oeconomicus* wählen würde. Oder wir spüren intuitiv, dass ein zu schlechtes Angebot den Empfänger dazu verleiten könnte, es abzulehnen, so dass auch wir als Sender schlechter dastünden. Wenn Ihnen jemand von einem Guthaben von € 20 nur einen Groschen abzugeben gewillt wäre, und € 19,90 für sich behielte – wäre der Spaß, ihm die € 19,90 aus der Hand zu schlagen, nicht auch 10 Cent wert? Wäre es „irrational“, die Person durch die Ablehnung des Angebots zu bestrafen, damit sie solche Angebote in Zukunft nicht mehr unterbreitet (selbst wenn wir annehmen, *dieser* Person nie wieder zu begegnen)? Und wenn wir an das *Framing*-Experiment denken: Können wir Versuchspersonen, die sich weigern, wie von der Rationaltheorie vorhergesagt eigennützig-rational zu handeln, aber dadurch mehr Geld verdienen als rationale Spieler, wirklich als „irrational“ im herkömmlichen Sinne des Wortes bezeichnen? Eine angemessenere, präzisere Beschreibung des Verhaltens, das den Vorhersagen des klassischen ökonomischen Paradigmas nicht entspricht, ist daher **nicht-rationales Verhalten**. 84

## c. Gelockerte Annahmen: Die verhaltenswissenschaftliche Wende

Was machen wir nun aus alledem? Sollen wir das ökonomische Paradigma aufgeben? Sicher nicht – dann würde das Buch hier enden. Was also dann? Sollen wir das Modell modifizieren? Sollen wir die empirischen Einsichten ignorieren und einfach bei unserem schlichten Modell bleiben? Die schlich- 85

teste Konsequenz ist, dass wir uns dieser Grenzen unserer theoretischen und empirischen Bemühungen und der Bedeutung des jeweiligen Kontexts (**ökologische Validität**) bewusst sein müssen – gerade als Juristen, weil im Recht der spezifische Kontext, der konkrete Sachverhalt eine besonders wichtige Rolle spielt. Andere Konsequenzen sind weniger augenfällig, oder erzeugen ihrerseits Probleme. So erscheint es wenig sinnhaft, das in sich geschlossene, einfache und in vielen Situationen wirkmächtige Modell des *homo oeconomicus* ersatzlos aufzugeben. Und bisher verfügen wir über keine Theorie, die menschliches Verhalten besser zu beschreiben und vorherzusagen vermag, und die uns ebenfalls erlaubt, empirisch überprüfbare Hypothesen zu generieren. Natürlich können wir das Modell mit zusätzlichen Parametern anreichern, etwa mit Risikopräferenzen (die sich, wie wir gesehen haben, schon im klassischen Modell abbilden lassen) oder mit (Un-)Gleichheitspräferenzen. Aber mit solchen Ergänzungen müssen wir sehr zurückhaltend verfahren: Wenn wir beobachten, dass Menschen morgens vor der Arbeit einen starken Kaffee trinken, und wenn wir deshalb in die allgemeine Nutzenfunktion eine Präferenz für starken Kaffee vor der Arbeit am Morgen einfügen, dann ist die Vorhersage, dass Menschen morgens vor der Arbeit Kaffee trinken, banal; eine solche Theorie hätte keinen sinnvollen Erklärungswert. Außerdem gilt: Je mehr Variablen wir berücksichtigen, je mehr Freiheitsgrade wir uns statistisch gesprochen erlauben, desto wahrscheinlicher wird es auch, aus einer falschen Theorie eine richtige Vorhersage abzuleiten – deswegen lag das geozentrische ptolemäische Weltbild so erstaunlich oft richtig.

- 86 Um die Einsichten empirischer Forschung verwerten zu können, hat die Ökonomik eine **verhaltenswissenschaftliche Wende** (*behavioral turn*) vollzogen. Heute befasst sich ein Teil der (Rechts-)Ökonomik, die Verhaltens(rechts)ökonomik (*Behavioral [Law &] Economics*), mit kritischen Anfragen an das klassische ökonomische Verhaltensmodell. Eine ganze Reihe von Modellannahmen sind auf vorsichtige Weise gelockert und daraufhin empirisch überprüft worden. Ein ganzer Wissenschaftszweig ist bemüht, mit Verhaltensexperimenten dazu beizutragen, jene (kognitiven) Mechanismen besser zu verstehen, die zum Auseinanderfallen von tatsächlich beobachtetem menschlichen Verhalten und den Vorhersagen des ökonomischen Verhaltensmodells führen. Die Hoffnung ist, dass wir eine ausgeklügeltere Theorie menschlichen Verhaltens entwickeln können. Die psychologisch informierte Verhaltenstheorie wird durch drei Begrenzungen charakterisiert, die das traditionelle ökonomische Modell modifizieren: **begrenzttes Eigeninteresse** (*bounded self-interest*), **begrenzte Rationalität** (*bounded rationality*) und **begrenzte Selbstdisziplin** (*bounded*

*self-control/willpower*).<sup>18</sup> Allerdings hat die Verhaltensökonomik bis heute den Status einer Mikrodisziplin; sie kann kein in sich geschlossenes, allgemeines Modell menschlichen Verhaltens vorweisen, sondern ist bislang auf situationsbezogene Aussagen beschränkt. Das ist dort unschädlich, wo wir detaillierte Kenntnisse vom Kontext und von situativen Parametern haben, die Entscheidungen beeinflussen können. Wo wir die Spezifika einer Situation nicht kennen oder einzuschätzen vermögen, bleibt eine allgemeine Verhaltenstheorie vonnöten. Aus diesem Gründen koexistieren heute klassische und verhaltenswissenschaftlich informierte Ansätze in der Ökonomie harmonisch nebeneinander.

## II. Wohlfahrtsanalyse und Effizienz

Auch wenn der Ausgangspunkt der ökonomischen Forschung das Verhalten des Einzelnen ist, geht es der Ökonomie im Wesentlichen darum, gesellschaftliche Phänomene und Prozesse zu beschreiben und zu erklären. Dabei bleibt sie jedoch nicht immer bei der rein empirischen Bestandsaufnahme stehen. Vielmehr gibt es auch eine **normative** Richtung der Ökonomie, der es darum geht, Institutionen so zu gestalten, dass die allgemeine Wohlfahrt gefördert wird. Die Entwicklung entsprechender Maßstäbe ist Gegenstand der **Wohlfahrtsökonomik**. Diese beschäftigt sich mit der Frage, auf welche Weise gesamtgesellschaftliche (also nicht nur individuelle) Wohlfahrtssteigerungen oder – wenn möglich – soziale Optima erzielt werden können. Ihr Kernbegriff ist die **Effizienz**: Je effizienter eine Gesellschaft organisiert ist, desto größer die von ihr erreichte Wohlfahrt. Was aber sind angemessene **Effizienzkriterien** zur Beurteilung wirtschaftspolitischer Maßnahmen oder wirtschaftlicher Ordnungssysteme? Die Wohlfahrtsökonomik hat sich zur Aufgabe gemacht, solche Maßstäbe zu entwickeln und zur Bewertung politischer Entscheidungen und rechtlicher Interventionen anzuwenden. 87

Der Ausgangspunkt ist dabei ein **deskriptiver**, indem die Auswirkungen solcher Maßnahmen und Ordnungssysteme beschrieben und vorhergesagt werden. Ziel ist letztlich jedoch eine **normative Bewertung der sozialen Gesamtzustände**. Gerade letzterer Aspekt ist nicht unproblematisch: Es müssen Maßstäbe gefunden werden, die es ermöglichen, die gesamtgesellschaftliche Wohlfahrtssteigerung zu bewerten. Dies wird in der Regel durch einen Rückgriff auf die Gesamtsumme individueller Wohlfahrt er- 88

<sup>18</sup> Vgl. dazu Rz. 493 ff.



reicht. Dabei entsteht das Problem, dass subjektive Präferenzen verschiedener Individuen verglichen und addiert werden müssten: Die Addition des jeweiligen Nutzens unterschiedlicher Individuen setzt jedoch voraus, dass die Präferenzen verschiedener Individuen vergleichbar sind. Wir haben aber gesehen, dass Präferenzen Werturteile enthalten, deshalb subjektiv und **inkommensurabel** sind und Nutzenvergleiche damit gerade nicht möglich. In der älteren Wohlfahrtsökonomik wurde das Problem ignoriert – und einfach ein **kardinal** messbarer und **interpersonell vergleichbarer Nutzen** zugrunde gelegt, bevor der oben dargestellte ordinale Nutzenbegriff auch in diesem Forschungszweig allgemein anerkannt wurde. Um die Probleme, die die Inkommensurabilität individueller Nutzenvergleiche und der daraus folgende ordinale Nutzenbegriff aufwirft, zu lösen, wurden, wie bereits kurz erwähnt, zunächst das **Pareto-** und später das **Kaldor-Hicks-Kriterium** entwickelt. Sie berücksichtigen die Problematik der Unvergleichbarkeit subjektiver Präferenzen und wollen der Wohlfahrtsökonomik zugleich ein wissenschaftlich-objektives Fundament sein. Wichtig ist, schon an dieser Stelle festzuhalten, dass die Effizienz-Kriterien unterschiedlichen Möglichkeiten der Güterverteilung (Allokation) weitgehend neutral gegenüberstehen. Was den Juristen oftmals in erster Linie umtreibt – die „gerechte Verteilung“ der Güter – ist für den Ökonomen völlig belanglos, solange und soweit eine andere Güterallokation keine höhere Effizienz zur Folge hat.

### 1. Pareto-Effizienz

- 89 Der Effizienzbegriff nach *Vilfredo Pareto*<sup>19</sup> beruht auf den Prinzipien der Konsumentensouveränität, des Nonpaternalismus und der Einstimmigkeit. Danach sind die subjektiven Präferenzen der Individuen autonom zu betrachten und werden als solche respektiert, ohne dass eine Unterscheidung in „gute“ oder „schlechte“ Präferenzen erfolgt (**Konsumentensouveränität**). Der einzig für die Gesellschaft bedeutende Nutzen ist der Nutzen des Individuums, einen Selbstzweck des Staates gilt es nicht zu berücksichtigen (**Nonpaternalismus**). **Einstimmigkeit** meint, dass Änderungen der Allokation von Gütern der Zustimmung aller bedürfen, jeder also ein Vetorecht hat, das aber nur im Falle einer Schlechterstellung ausgeübt wird, nicht bei Indifferenz.

---

<sup>19</sup> Pareto, Manuel d'Économie Politique, 1909, Kap. VI, Nr. 33 sowie Appendice, N° 88, 89.

Vor diesem Hintergrund postuliert das **Pareto-Kriterium**, dass eine Situation A besser ist als eine Situation B, wenn es nach den verglichenen, individuellen subjektiven Präferenzen in A mindestens einem Individuum besser und keinem Individuum schlechter geht als in Situation B: 90

- **Pareto-superior** ist ein Zustand im Vergleich zu einem anderen Zustand, wenn sich der Nutzen *eines* Akteurs erhöht, während kein anderer schlechter gestellt wird.
- Besteht ein Zustand, gegenüber dem noch eine *Pareto-Verbesserung* möglich ist, so bezeichnet man ihn als **pareto-inferior**.
- Sind keine *Pareto-Verbesserungen* mehr möglich, kann also kein Akteur mehr besser gestellt werden, ohne dass ein anderer schlechter gestellt würde, ist ein Zustand **pareto-optimal**.

Dem *pareto-optimalen* Zustand liegen stets eine effiziente Produktion, ein effizienter Konsum und eine effiziente Produktionsstruktur zugrunde. Für eine effiziente Produktion heißt dies, dass die Produktion eines Gutes nicht erhöht werden kann, ohne dass dadurch zumindest diejenige eines anderen Gutes eingeschränkt werden müsste. **Effizienter Konsum** ist erreicht, wenn ohne neue Produktion, allein durch Tauschgeschäfte, ein Zustand erreicht ist, bei dem weitere Tauschgeschäfte zum gegenseitigen Vorteil nicht mehr möglich sind. Die **effiziente Produktionsstruktur** schließlich setzt voraus, dass Konsum und Produktion optimal aufeinander abgestimmt sind. 91

Auf Kritik stößt das Kriterium der *Pareto-Effizienz* vor allem deshalb, weil es die Anfangsverteilung der Güter nicht berücksichtigt und deshalb zu evidenten Ungerechtigkeiten führen kann. Dies gilt insbesondere, wenn schon die Verteilung der Güter in der Ausgangssituation ungleichmäßig ist. Jedenfalls ist eine Umverteilung von begüterten zu weniger begüterten Akteuren bei Anwendung des *Pareto-Kriteriums* nie effizient möglich, da eine solche Umverteilung unweigerlich zu einer Schlechterstellung der begüterten Akteure führen würde. Die *Pareto-Effizienz* gibt also keinerlei Aufschluss darüber, welcher Zustand vorzugswürdig ist, sondern festigt lediglich den Status quo. Dies ist allerdings als Bewertungskriterium für eine sozialpolitische Maßnahme oft wenig hilfreich, da es selten eine Maßnahme geben wird, die allein Menschen bevorzugt, aber niemandem zum Nachteil gereicht. Als reines Marktkriterium ist das *Pareto-Kriterium* somit auch für das Recht nur eingeschränkt nützlich. 92

## 2. Kaldor-Hicks-Kriterium

- 93 Das Effizienzkriterium nach *Nicholas Kaldor* und *John Hicks*<sup>20</sup> beansprucht, die Mängel des *Pareto*-Kriteriums zu überwinden. Nach diesem Kompensationskriterium wird die Entscheidung zwischen zwei Zuständen zwar weiterhin auf Grundlage des individuellen Nutzens getroffen, jedoch muss nicht mehr jeder einzelne Akteur einer Maßnahme zumindest indifferent gegenüberstehen, vielmehr dürfen Kosten und Nutzen einer sozialen Entscheidung bei unterschiedlichen Akteuren anfallen. Wohlfahrtssteigernd ist demnach schon eine solche Änderung, bei der aus den Gewinnen der bessergestellten Akteure die Verluste der Verlierer kompensiert werden können, und zusätzlich zumindest ein Akteur nach dieser Kompensation besser steht als im alten Zustand. Diese Entschädigung muss dabei nicht tatsächlich geleistet werden, es genügt, dass sie möglich wäre, ohne dass alle Wohlfahrtszuwächse bei den Gewinnern aufgezehrt würden.
- 94 Unter anderem an dieser nur hypothetischen Entschädigung der Verlierer entzündet sich die Kritik am *Kaldor-Hicks*-Kriterium. Es sei fraglich, ob ein Zustand wohlfahrtsökonomisch vorzugswürdig sein kann, bei dem die „Verlierer“ nicht tatsächlich entschädigt würden und sich somit an der Ausgangsverteilung nichts ändere. Eine solche Umverteilung ist jedoch durch das *Kaldor-Hicks*-Kriterium gar nicht angestrebt, sieht es doch seine Aufgabe nicht darin, die Probleme der Einkommensverteilung zu lösen.<sup>21</sup> Gegen eine tatsächliche Kompensation spricht auch, dass unter Umständen der Kreis der zu entschädigenden Personen unüberschaubar ist und sich nicht eingrenzen lässt. Der mit einer tatsächlichen Kompensation verbundene Aufwand würde im Gegenteil zu Effizienzverlusten führen und die Wohlfahrtsgewinne reduzieren.
- 95 Allerdings bedeutet die wegen der nur hypothetischen Entschädigung fehlende Umverteilung einen weiteren Schwachpunkt des *Kaldor-Hicks*-Kriteriums: Misst man die Zuteilung der Gewinne und Verluste, und damit auch die notwendige Entschädigung, an der Zahlungsbereitschaft der Akteure, so sind reichere Akteure aufgrund des ihnen faktisch zur Verfügung stehenden größeren Vermögens und des abnehmenden Grenznutzens ihres Einkommens in der Lage, einen höheren Preis zu zahlen und demnach im Vorteil. Eine höhere Zahlungsbereitschaft zeigt dem *Kaldor-Hicks*-Krite-

<sup>20</sup> *Kaldor*, Welfare Propositions of Economics and Interpersonal Comparisons of Utility, *Economic Journal* 49 (1939), 549ff. und *Hicks*, The Foundation of Welfare Economics, *Economic Journal* 49 (1939), 696ff.

<sup>21</sup> *Kaldor*, Welfare Propositions of Economics and Interpersonal Comparisons of Utility *Economic Journal* 49 (1939), 549 (550f.).

rium zufolge einen höheren Nutzen, der wohlhabendere Akteur ist damit in der Regel in der Position desjenigen, der den Verlierer entschädigen müsste; da diese Entschädigung aber nicht tatsächlich erfolgt, mehrt sich sein Vermögen und er ist für die Zukunft weiterhin in der Position desjenigen, dessen Nutzen sich erhöht, ohne dass er die Verluste der Schlechtergestellten kompensieren müsste. Ein zusätzliches Problem stellt die Messbarkeit der potenziellen Entschädigung dar. Da die Entschädigung nicht in Nutzen, sondern in Geld gemessen wird, wird dadurch auch nicht der gesellschaftliche Nutzen, sondern nur die **hypothetische Zahlungsbereitschaft** gemessen. Dies stellt gerade bei immateriellen Gütern ein Problem dar.



## § 3 – Nachfrage, Angebot und Märkte

**Literatur:** *G. A. Akerlof*, The Market for ‘Lemons’: Quality Uncertainty and the Market Mechanism, *Quarterly Journal of Economics* 84 (1970), 488–500; *K. J. Arrow/G. Debreu*, Existence of an Equilibrium for a Competitive Economy, *Econometrica* 22 (1954), 265–290; *R. H. Coase*, The Nature of the Firm, *Economica* 4 (1937), 386–405; *ders.*, The Problem of Social Cost, *Journal of Law & Economics* 3 (1960), 1–44; *R. D. Cooter/T. S. Ulen*, *Law and Economics*, 6. Aufl. 2012; *H. Demsetz*, Toward a Theory of Property Rights, *American Economic Review* 57 (1967), 347–399; *N. G. Mankiw/M. P. Taylor*, *Grundzüge der Volkswirtschaftslehre*, 6. Aufl. 2016; *J. v. Neumann/O. Morgenstern*, *Theory of Games and Economic Behavior*, 1953; *W. Nicholson/C. Snyder*, *Microeconomic Theory*, 10. Aufl. 2008; *E. Ostrom*, *Governing the Commons: The Evolution of Institutions for Collective Action*, 1990; *A. C. Pigou*, Divergences Between Marginal Social Net Product and Marginal Private Net Product, in: *A. C. Pigou* (Hg.), *The Economics of Welfare II*, 4. Aufl. 1932, Chapter IX; *R. S. Pindyck/D. L. Rubinfeld*, *Mikroökonomie*, 8. Aufl. 2013; *A. Smith*, *An Inquiry into the Nature and Causes of the Wealth of Nations: Two Volumes in One*, 1994; *H. R. Varian*, *Grundzüge der Mikroökonomik*, 8. Aufl. 2011.

### I. Einleitung

Dieses Kapitel behandelt die drei vielleicht praktisch relevantesten Konzepte der Wirtschaftswissenschaften: Nachfrage, Angebot und Märkte. Diese Konzepte können im Recht auf viele verschiedene Sachverhalte angewandt werden. Stellen Sie sich etwa vor, Sie seien Lokalpolitiker, Richter oder Beamter und müssten einen Fall entscheiden, in dem Anwohner ein Nachtflugverbot eines nahegelegenen Flughafens durchsetzen wollen. In dem Konflikt geht es um knappe Ressourcen. Sie können nicht zugleich ruhige Nächte für die Anwohner und einen nachts aktiven Flughafen erreichen. Jede juristische Lösung des Streits verteilt Rechte und somit Ressourcen an die Parteien (Nachtruhe für Anwohner bzw. für den Flughafen die Möglichkeit, Geld zu verdienen). Im ökonomischen Paradigma<sup>1</sup> haben Sie

---

<sup>1</sup> Vgl. oben Rz. 64.

gesehen, dass sich die Wirtschaftswissenschaften mit der Verteilung knapper Ressourcen befassen. Nach dem Effizienzmaßstab sollte das Recht, über die Aktivität des Flughafens zu entscheiden, an diejenigen gehen, der es am meisten wertschätzt. Der folgende Abschnitt über die Nachfrage (II.) wird sich nun damit auseinandersetzen, wie Ökonomen über Wert und Wertschätzung von Gütern nachdenken – zum Beispiel über den Wert einer ruhigen Nacht. Der Begriff des Nutzens sowie der Umstand, dass der Grenznutzen abnimmt, wurden bereits in § 2 eingeführt.<sup>2</sup> Hier werden Sie erfahren, wie Nutzen aus Abwägungen konstruiert wird, wie der Begriff der Nachfrage auf dem des Nutzens beruht, und inwiefern der Begriff des Wertes, ökonomisch gesehen, fundamental relativ ist. Im darauffolgenden Abschnitt III. werden Sie sehen, dass das Konzept des Angebots auf dem der Kosten beruht, und warum Kosten im ökonomischen Sinne als entgangene Möglichkeiten beschrieben werden. Abschnitt IV. wird erklären, wie Märkte dazu beitragen, Kosten und Nutzen gegeneinander abzuwägen, wie Märkte den Preis knapper Güter bestimmen und wie sie deren effiziente Zuteilung fördern können. Abschnitt V. befasst sich schließlich mit dem Scheitern von Märkten. Es wird gezeigt, wie ein Unternehmen, das auf Kosten Dritter (etwa schlafender Anwohner) ein Gut oder eine Dienstleistung anbietet (z.B. nächtliches Fliegen), eine effiziente Zuteilung von Gütern gefährden kann und wie rechtliche Interventionen dafür sorgen können, dass Firmen die vollen Kosten ihrer Aktivitäten berücksichtigen.

## II. Nachfrage

- 97 Die Nachfrage beschreibt die Wertschätzung von Gütern, indem sie die konsumierte Menge mit dem Preis des Gutes in Beziehung setzt. Sie bildet damit ab, inwieweit Käufer finden, eine Sache sei „ihren Preis wert“. In der Ökonomie wird der Begriff „Wert“ meist in zwei unterschiedlichen Bedeutungen gebraucht. Zum einen bezeichnen Ökonomen hiermit den Tauschwert eines Gutes, also den „Wert“ im Sinne des Preises. Mit dieser Deutung werden wir uns aber erst im Abschnitt über Märkte befassen. In diesem Abschnitt über Nachfrage setzen wir Preise als gegeben voraus und behandeln „Wert“ im Sinne von Wertschätzung. Wir fragen also, wie Einzelpersonen Güter bewerten. Zwischen Tauschwert und Wertschätzung eines Gutes besteht aber natürlich ein Zusammenhang. Der Tauschwert eines Gutes wird letztlich davon abhängen, wie sehr Menschen das Gut schätzen.

---

<sup>2</sup> Hierzu Rz. 70 ff., 73.

## 1. Bewertung von Gütern

Da Menschen Dinge in der Regel sehr unterschiedlich bewerten, verstehen Ökonomen die Wertschätzung zutiefst relativ. Der Relativismus der Wertschätzung geht in der Ökonomie so weit, dass es keinen absoluten Maßstab für Wertschätzung gibt. Wie wir sehen werden, ist es ein Irrglaube, dass Geld den Wert eines Gutes objektiv messen kann. Die Wertschätzung eines Gutes durch eine Person wird in der Ökonomie deshalb anhand der Wertschätzung anderer Güter durch dieselbe Person gemessen. Dem ökonomischen Denken über Wertschätzung liegt eine Auswahltheorie zugrunde, die ihrerseits auf den folgenden Annahmen beruht. 98

### a. Eine rationale Präferenzordnung als Grundbaustein der Nachfrage

Im Ausgangspunkt stellen sich Ökonomen die Bewertung von Gütern als Rangfolge vor: Jeder Mensch sortiert alle ihm zur Verfügung stehenden Auswahloptionen von der erwünschtesten zu der am wenigsten erwünschten. Diese Theorie verlangt, dass man zum Beispiel bei der Wahl zwischen einer Flasche Wein und einem Stück Käse sagen kann, ob man den Wein dem Käse vorzieht, oder umgekehrt Käse dem Wein, oder ob man beides gleich ansprechend fände. Die Rangordnung von Optionen kann man **Präferenzordnung** oder Nutzenskala nennen: Höher eingestufte Optionen verschaffen dem Individuum mehr Nutzen als niedriger einsortierte. 99

Damit man diese Rangfolge als **rational** bezeichnen kann, muss sie bestimmte Eigenschaften haben. Diese Eigenschaften werden in den folgenden Annahmen beschrieben. Die Rangfolge muss zunächst **vollständig** sein. Das bedeutet nicht, dass man zu jedem beliebigen Zeitpunkt die komplette Reihenfolge aller möglichen Optionen genau kennt. Es bedeutet nur, dass man im Prinzip für jegliche Wahlentscheidung, der man jemals ausgesetzt sein wird, eine Rangfolge aufstellen könnte. 100

Die zweite Annahme einer rationalen Präferenzordnung wird **Transitivität** genannt. Sie bedeutet, dass die Rangfolge der Optionen widerspruchsfrei ist. Zieht man also die Flasche Wein dem Stück Käse vor und den Käse einer Schale Kracker, kann man sagen, dass der Wein auch dem Salzgebäck vorgezogen wird.<sup>3</sup> 101

Angeht es um die ersten beiden Annahmen können Menschen mit einem begrenzten Budget ihren **Nutzen maximieren**. Im Rationalmodell wird angenommen, dass Menschen genau das tun, und zwar mit Perfektion: Sie 102

<sup>3</sup> Hierzu Rz. 65.



maximieren ihren eigenen Nutzen, indem sie aus den erreichbaren Optionen immer die am höchsten eingestufte auswählen.

- 103 Die dritte – eher technische – Annahme ist **Stetigkeit**. Wenn man eine Flasche Wein einem Stück Käse vorzieht, zieht man auch alles, was einer Flasche Wein ausreichend ähnlich ist (etwa ein wenig mehr oder weniger Wein) dem Käse vor. Diese Annahme besagt lediglich, dass kleine Änderungen an den Optionen nicht zu größeren Veränderungen in der Rangfolge führen. Zweck dieser Annahme ist es, Nutzen in mathematischen Funktionen auszudrücken und auf diese Funktionen mathematische Maximierungsmethoden anwenden zu können.
- 104 Jede dieser Rangfolgen sortiert die Optionen nach dem Nutzen für ein einziges Individuum. Neumann und Morgenstern gelang es zwar, unter Verwendung von Wahrscheinlichkeiten etwas darüber zu sagen, wie weit Optionen auf der Nutzenskala voneinander entfernt liegen. Diese Erkenntnis ändert aber nichts daran, dass jeder Vergleich von Nutzen zwischen Individuen bedeutungslos ist. Man mag sagen können, dass ich eine Flasche Wein doppelt so sehr wertschätze wie ein Stück Käse. Aber es ist sinnlos zu sagen, dass ich eine Flasche Wein doppelt so sehr wertschätze wie Sie.
- 105 Schließlich treffen Ökonomen üblicherweise ein paar vereinfachende Annahmen. Zum Beispiel nehmen Ökonomen typischerweise an, dass alle Dinge unverändert bleiben, außer dem einen Umstand, der ausdrücklich variiert wird (sogenannte *Ceteris-paribus*-Annahme). Sie mögen also eher Wein als Käse, doch Sie genießen das Schwätzchen mit der Person hinter der Käsetheke. Aus diesem Grund kaufen Sie letztlich lieber Käse. Ein Ökonom würde sich die Freiheit nehmen, den Nutzen der Lebensmittel separat zu bemessen. Er nähme daher an, dass sich Ihre Gefühle über das Schwätzchen mit dem Verkaufspersonal nicht verändern – egal, welche Lebensmittel Sie kaufen. Immer, wenn wir also sagen, dass wir Käse gegen Wein abwägen, können wir uns eine Wahl zwischen zwei Güterbündeln vorstellen: alles was wir haben plus eine Flasche Wein oder aber alles was wir haben plus ein Stück Käse. Ökonomen nehmen normalerweise auch an, dass mehr besser sei. **Güter** sind Dinge, die den Nutzen erhöhen. Diese Annahme ist natürlich fragwürdig, wenn Sie an Wein denken (Stichwort Kater!), aber sie vereinfacht die Dinge erheblich, solange man sich mit vernünftigen Mengen des Konsums beschäftigt.

### b. Indifferenzkurven: Messung des Werts eines Gutes an einem anderen Gut

Diese wenigen Annahmen reichen aus, um den Begriff der Wertschätzung ein wenig zu formalisieren. Angenommen, Sie haben eine gewisse Ausstattung an Wein und Käse. Und jetzt überlegen Sie: Wenn Sie eine halbe Flasche Wein aufgeben müssten, wie viel Käse würden Sie benötigen, um genauso glücklich zu sein, wie Sie es mit Ihrer ursprünglichen Ausstattung waren? Die Käsemenge, die Sie brauchen, drückt den Wert jener halben Flasche Wein aus, die Sie aufgegeben haben. Die Volkswirtschaftslehre hat diese Art des Denkens in Form einer Graphik formalisiert, die sie Indifferenzkurve (s. Abb. 3.1) nennt. Jede dieser Kurven stellt alle Kombinationen von Wein und Käse (oder anderen Gütern) dar, die einer Person denselben Nutzen bringen. 106

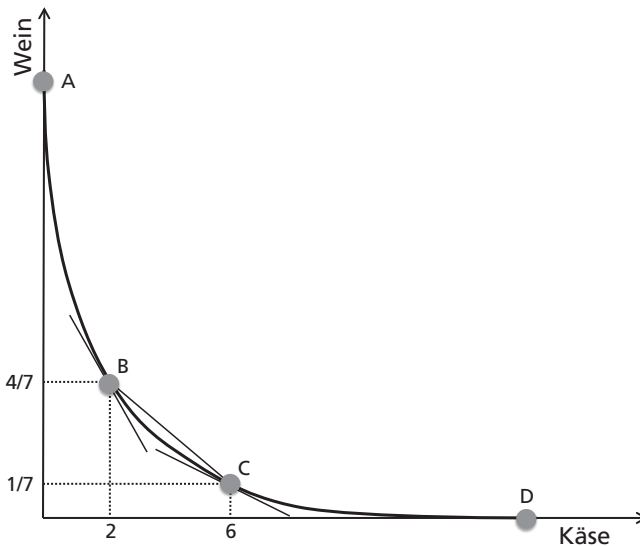


Abbildung 3.1: Indifferenzkurve

Eine Person mit den in der Indifferenzkurve abgebildeten Präferenzen wäre mit vier Gläsern Wein ( $4/7$  einer Flasche) und zwei Stücken Käse (Punkt B) ebenso zufrieden, wie mit einem Glas Wein ( $1/7$  einer Flasche) und 6 Stücken Käse (Punkt C). Es macht ihr also nichts aus, drei Gläser Wein für vier Stücke Käse aufzugeben. Die Steigung der Geraden, die durch diese beiden Punkte läuft, stellt den „Wechselkurs“ von Wein in Käse dar, den wir als **Substitutionsrate** bezeichnen. Wenn wir uns vorstel- 107

len, dass beide Güter in sehr kleine, tauschbare Einheiten umgewandelt werden können, sehen wir, dass die Steigung der Indifferenzkurve an einem bestimmten Punkt (vgl. etwa die Tangente in B) die Substitutionsrate für eine sehr kleine Gütereinheit darstellt. Diese Substitutionsrate für eine beliebig kleine Einheit wird **marginale Substitutionsrate** genannt.

108 Betrachtet man die Kurve, so sieht man, dass sich die marginale Substitutionsrate (d.h. die Steigung der Kurve) verändert (vgl. die Tangenten in B und in C). Die Indifferenzkurve beginnt steil abzufallen, wird dann aber flacher und flacher. Die Veränderung der Steigung bedeutet: Wenn die Person nur Wein hat, jedoch keinen Käse (Punkt A), ist sie bereit, viel Wein für nur sehr wenig Käse zu opfern. Am anderen Ende der Kurve (Punkt D) hat die Person nur Käse und ist deshalb bereit, sehr viel davon für lediglich etwas Wein zu opfern. Normalerweise haben wir lieber kleine Mengen von verschiedenen Dingen als große von einem einzigen. Das drückt die Kurve aus. Wein mit Käse ist viel besser als nur Käse oder nur Wein. In der Mitte der Kurve hat die Person eine vernünftige Kombination aus Wein und Käse zur Verfügung und wird daher Wein für Käse in relativ gleichen Mengen tauschen. Was für Wein und Käse gilt, trifft auch auf viele andere Güter zu. Niemand kann nur von Brot oder nur von Wasser leben. Und man zieht es sicher vor, in einer Stadt zu leben, in der es nicht nur ein Kino, sondern auch einen Konzertsaal gibt, denn wir mögen Abwechslung.

109 Wir können uns verschiedene Substitutionsbeziehungen zwischen zwei Gütern vorstellen. Güter können **Substitute** (Ersatzmittel) oder **Komplemente** (Ergänzungen) sein. Sind sie perfekte Substitute, bedeutet dies, dass es einem egal ist, ob man das eine oder das andere hat, da jedes Gut die betreffenden Bedürfnisse erfüllt. Da perfekte Substitute nun einmal perfekt sind, kommen sie sehr selten vor. Zucker und Süßstoff kommen perfekten Substituten sehr nahe. Viele würden auch verschiedene Weißweine als Substitute angesehen, etwa Weißburgunder und Chardonnay. Indifferenzkurven für perfekten Ersatz sind gerade. Die marginale Substitutionsrate ist im Verlauf der gesamten Indifferenzkurve konstant. Egal, wie viel Zucker oder Süßstoff man besitzt, man ist immer bereit, einen Löffel Zucker für die in Süßkraft entsprechende Menge Süßstoff abzugeben (s. Abb. 3.2).

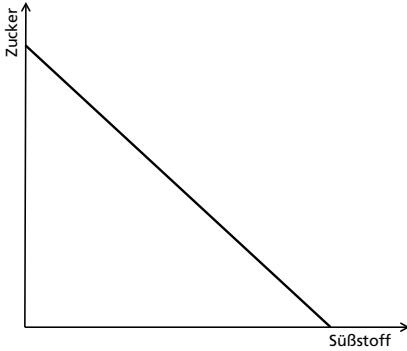


Abbildung 3.2: Indifferenzkurve für perfekte Substitute

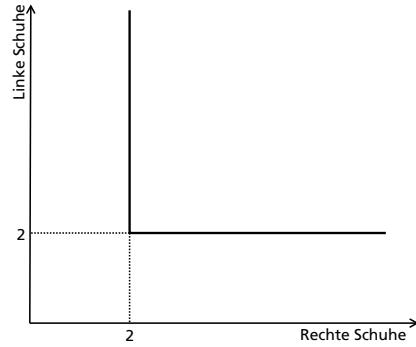


Abbildung 3.3: Indifferenzkurve für perfekte Komplemente (zwei Paar Schuhe)

Güter können auch perfekte Komplemente sein, d.h. man kann ein Gut nur nutzen, wenn vom anderen auch etwas vorhanden ist. Das Standardbeispiel hierfür sind Schuhe. Einen übriggebliebenen linken Schuh gibt man gerne auf, weil man ihn nicht nutzen kann, solange der passende rechte Schuh fehlt. Die Indifferenzkurven für perfekte Komplemente sehen aus wie jene in Abb. 3.3. Das Verhältnis zwischen den meisten Gütern liegt zwischen diesen beiden Extremen. Nehmen wir beispielsweise wieder Wein und Käse. Einerseits lieben Sie die Kombination der beiden, andererseits sind Sie jedoch stets bereit, den Verlust eines Nahrungsmittels mit dem Gewinn der richtigen Menge des anderen auszugleichen.

110

### c. Annahmen und Indifferenzkurven

Man kann die Indifferenzkurven mit den anfangs vorgestellten Annahmen rationaler Präferenzordnungen verknüpfen. Wir haben angenommen, mehr sei besser. Nun stellen wir uns die Ausstattung einer Person vor, dargestellt durch einen Punkt auf einer Indifferenzkurve dieser Person. Stellen Sie sich vor, die Person erhielte mehr, ohne dass sie gezwungen wäre, etwas dafür aufzugeben. Der Nutzen der Person würde zunehmen. Im Diagramm würde sich das Nutzenniveau nicht entlang der Kurve bewegen. Stattdessen würde es sich auf eine andere Indifferenzkurve bewegen, die nordöstlich der ursprünglichen Kurve läge, und ein höheres Nutzenniveau darstellt. Im Diagramm können wir uns eine unendliche Menge von Indifferenzkurven vorstellen, die sämtliche mögliche Nutzenniveaus einer Person darstellen (s. Abb. 3.4). Je weiter nordöstlich man voran-

111

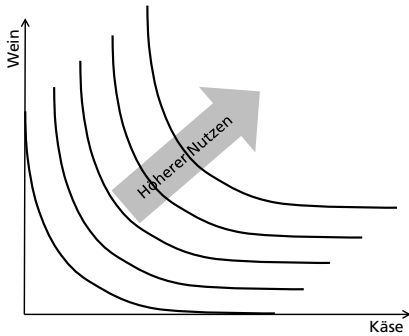


Abbildung 3.4: Vollständigkeit – es gibt unendlich viele Indifferenzkurven; „Mehr ist besser“ – das Nutzenniveau steigt mit größerem Konsum; Kontinuität – zwischen jedem Paar von Indifferenzkurven kann eine weitere Indifferenzkurve liegen.

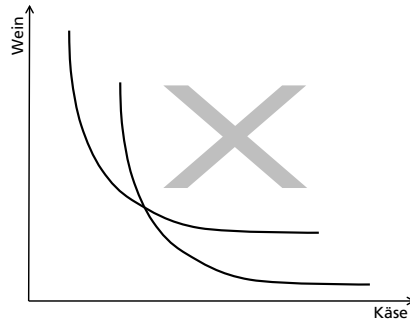


Abbildung 3.5: Transitivität – Indifferenzkurven einer Person schneiden sich nicht.

kommt, desto höher ist das Nutzenniveau. Somit wird durch die Indifferenzkurven dargestellt, dass mehr besser ist.

- 112 Wir haben auch die Vollständigkeit der Präferenzordnung angenommen. Im Prinzip können die Indifferenzkurven eine Rangfolge aller Optionen, mit denen eine Person konfrontiert ist, darstellen. Wie die Rangfolge verschiedener Bündel von Gut A oder Gut B dargestellt werden kann, haben wir gerade gesehen. Die Person zieht jedes Bündel von einer nordöstlicheren Indifferenzkurve dem Bündel vor, das sie in Händen hält. Sie ist indifferent zwischen dem Bündel, das sie hält, und jedem Bündel auf der Indifferenzkurve, auf der sie sich befindet. Und sie zieht das Bündel, das sie hält, jedem Bündel einer südwestlicheren Indifferenzkurve vor. Obwohl wir uns bisher auf Indifferenzkurven von lediglich zwei Gütern beschränkt haben, könnten wir die Anzahl der Güter erhöhen, indem wir Dimensionen hinzufügen. Wir könnten ein dreidimensionales Diagramm für drei Güter zeichnen, und die Mathematik gibt uns sogar die Mittel, bei Bedarf mit vieldimensionalen Indifferenzkurven eine Fülle von Gütern bearbeiten zu könnten. Indifferenzkurven sind also flexibel genug, um eine vollständige Präferenzordnung darzustellen. Um jedoch Tauschgeschäfte zwischen einem bestimmten Gut und allen anderen Konsummöglichkeiten darzustellen, können wir auch in zwei Dimensionen bleiben. Hierzu schauen wir uns das Gut an, das uns interessiert, sowie ein zweites, das alle anderen Konsummöglichkeiten darstellt. Dieses zweite Gut wird Geld genannt. Geben wir Geld für das Gut aus, das uns interessiert, repräsentiert die gezahlte Geldmenge (fast) alle Konsummöglichkeiten, die wir auslassen,

weil wir unsere Ressourcen in das ausgewählte Gut stecken. Wenn wir also Wein durch Geld ersetzen, sind unsere zweidimensionalen Diagramme eigentlich ziemlich allgemein.

Darüber hinaus haben wir auch Transitivität, d. h. eine widerspruchslöse Präferenzordnung, angenommen. Diese Annahme korrespondiert mit der Tatsache, dass sich Indifferenzkurven nicht kreuzen können. Wir sprachen davon, dass alle Punkte auf einer Indifferenzkurve dasselbe Nutzenniveau darstellen und alle Punkte, die weiter nordöstlich liegen, ein höheres. Gäbe es also einen Punkt, an dem sich zwei Indifferenzkurven kreuzen, so müssten beide Indifferenzkurven dasselbe Nutzenniveau darstellen. In Abb. 3.5 liegt aber links vom Schnittpunkt der Indifferenzkurven eine Kurve nordöstlicher als die andere und stellt demnach ein höheres Nutzenniveau als die zweite Kurve dar. Dies widerspricht jedoch der ersten Aussage, dass beide Indifferenzkurven dasselbe Nutzenniveau darstellen. Demnach wären sich kreuzende Indifferenzkurven widersprüchlich. Hieraus folgt, dass durch sich kreuzende Indifferenzkurven dargestellte Präferenzen nicht konsistent und daher durch die Transitivitätsannahme ausgeschlossen sind. Schließlich korrespondiert die Tatsache, dass es stets genügend Platz für eine weitere Indifferenzkurve zwischen zwei Indifferenzkurven gibt, mit der Annahme der **Kontinuität**. 113

## 2. Nutzenmaximierung

Das Modell der rationalen Wahl, das dem Nachfragekonzept zugrunde liegt, besagt, dass Agenten durch ihre Wahl ihren Nutzen maximieren. Stellen Sie sich also vor, dass Sie einen ruhigen Abend für sich planen, mit Wein und Käse. Zu diesem Zweck haben Sie zwei Flaschen Wein beiseitegestellt – aber Sie haben den Käse vergessen. Wie retten Sie Ihren Abend? Sie können von einem Ihrer Nachbarn etwas Käse „kaufen“ und mit Wein „bezahlen“. Sie wissen, dass Ihr Wein in Ihrer Feinschmeckernachbarschaft üblicherweise gegen vier Stücke Käse getauscht wird. 114

Wie viel Wein werden Sie also gegen Käse eintauschen? Ein Diagramm kann helfen. Die Tatsache, dass Sie zwei Flaschen Wein haben und die Tauschrate in Ihrer Nachbarschaft kennen, klärt, welche Bündel für Sie erreichbar sind und welche unerreichbar. Sie können entweder die beiden Flaschen behalten und ohne Käse auskommen. Oder Sie geben beide Flaschen für acht Stücke Käse weg. Sie können aber auch jeden Tausch zwischen diesen Extremen realisieren. Im Diagramm können sie diesen Gedanken durch eine gerade Linie zwischen den beiden Bündeln „zwei Flaschen Wein und kein Käse“ und „acht Stücke Käse und kein Wein“ 115

repräsentieren. Diese Linie wird Budgetlinie oder **Budgetgrenze** genannt. Sie können jedes Bündel südwestlich von oder direkt auf der Linie erreichen. Alle Bündel noröstlich sind für Sie unerreichbar.

- 116 Im Modell rationaler Wahl gilt die Annahme, dass Sie Ihren Nutzen maximieren. Sie werden kein Güterbündel unterhalb der Budgetlinie wählen, da solche Bündel bedeuten, dass Sie ohne Not auf Wein oder Käse verzichten. Stattdessen wählen Sie ein Bündel auf der Linie. Doch welches wird Ihren Nutzen maximieren? Dazu können wir die Indifferenzkurven konsultieren. Sie wollen das höchste Nutzenniveau erreichen. Daher möchten Sie im Diagramm auf eine Indifferenzkurve soweit nordöstlich wie nur möglich vorstoßen. Sie wählen also den Punkt aus, an dem die Budgetlinie die höchstmögliche Indifferenzkurve gerade noch tangiert (s. Abb. 3.6).

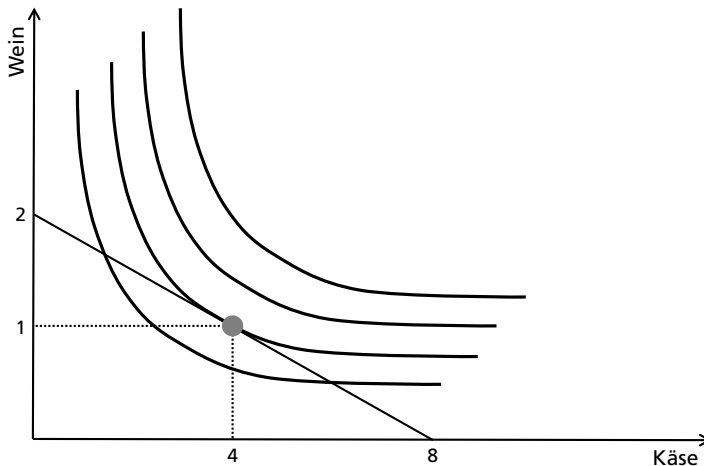


Abbildung 3.6: Nutzenmaximierung mit begrenztem Budget

- 117 Formaler ausgedrückt: Sie wählen den Punkt, an dem Ihre persönliche Substitutionsrate von Wein in Käse (die Steigung der Indifferenzkurve an einem gegebenen Punkt) genau dem in Ihrer Nachbarschaft vorherrschenden „Käsepreis“ in Wein entspricht (die Steigung der Budgetlinie). Sie kaufen so lange Wein für Käse, bis Ihnen die nächste gekaufte Käseeinheit gerade weniger Nutzen bringt als die nächste Menge Wein, die Sie dafür aufgeben.

### 3. Preisänderungen

Sie haben vielleicht gehört, dass die Nachfrage steigt, wenn der Preis sinkt. 118 Das ist das **Gesetz der Nachfrage**. Im Diagramm wird eine Preisänderung durch eine Änderung der Steigung der Budgetlinie (s. Abb. 3.7) dargestellt. Stellen Sie sich in unserem Wein-Käse-Beispiel vor, dass der „Käsepreis“ gefallen ist. Sie können nun für eine Flasche Wein sechs statt bisher nur vier Käsestücke bekommen. Dann wissen Sie auch, dass Sie zwei Flaschen Wein oder zwölf Stücke Käse oder jede beliebige Zuteilung dazwischen erreichen können. Wegen der Preisänderung wird die Budgetlinie flacher. Doch welche Auswirkung wird die Preisänderung auf Ihr Konsumverhalten haben? Sicher werden Sie sich nun mehr Käse besorgen, da er günstiger ist. Dieser Effekt nennt sich **Substitutionseffekt**. Darüber hinaus macht Sie der gefallene Käsepreis reicher. Er schiebt die Budgetlinie nach außen. Da Sie nun nicht mehr so viel bezahlen müssen, um eine adäquate Käsemenge zu kaufen, können Sie mehr Wein für sich behalten. Dieser Effekt nennt sich **Einkommenseffekt**. Graphisch können wir beide Effekte trennen. Wir nehmen das Ihnen bekannte Diagramm (Abb. 3.7) und schieben die Budgetlinie um vier Käsestücke auf der Horizontalachse nach außen. So stellen wir die Tatsache dar, dass die Preissenkung Ihnen nun erlaubt, zwei Flaschen Wein gegen bis zu zwölf Käsestücke einzutauschen. Dann gehen wir die ursprüngliche Indifferenzkurve vom Punkt, an dem wir begonnen haben, entlang bis zu dem Punkt, an dem die Kurve die Steigung der neuen Budgetlinie hat. Dies ist der Weg vom gestrichelten Kreis zum hohlen Kreis in Abb. 3.8, wobei die gestrichelte Linie der Steigung der neuen Budgetlinie entspricht. Die Zunahme des Käsekaufs (und die Abnahme des behaltenen Weins) ist eine Folge des Substitutionseffekts. Als nächstes gehen Sie weiter nordöstlich zu dem Punkt, an welchem die Budgetlinie die neue Indifferenzkurve tangiert. Dies ist der Weg vom hohlen Kreis zum gefüllten Kreis in Abb. 3.8. Die Zunahme des Wein- und Käsekonsums, die mit diesem Weg verbunden ist, stellt den Einkommenseffekt dar. Der Mechanismus einer Preissteigerung wird durch dieselbe Schrittfolge in umgekehrter Reihenfolge illustriert.



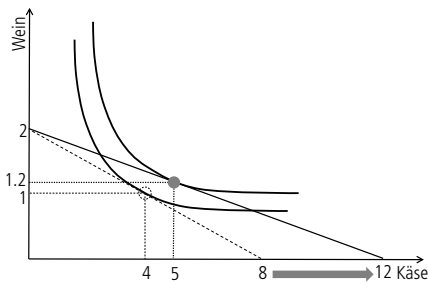


Abbildung 3.7: Reaktion des Verbrauchs auf Preisveränderung

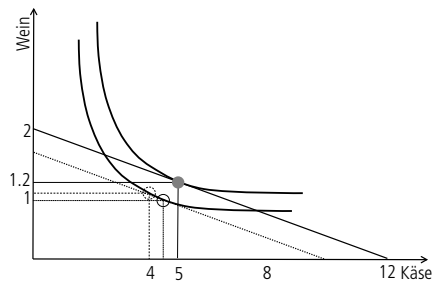


Abbildung 3.8: Substitutions- und Einkommenseffekt

#### 4. Nachfragefunktionen

- 119 Die Nachfragefunktion sagt aus, wie der Konsum eines Gutes auf Veränderungen seines Preises reagiert, wenn alle anderen Faktoren konstant bleiben.

##### a. Nachfragekurven von Individuen

- 120 Weil wir die Konsumveränderungen in Reaktion auf Preisveränderungen aus den Indifferenzkurven ableiten können (so wie wir das im letzten Abschnitt getan haben), können wir eine Funktion zeichnen, die uns eine Konsummenge für jeden möglichen Preis des Gutes anzeigt (Abb. 3.8.). Diese Funktion nennt sich **Nachfrage**. Die Konvention verlangt, dass wir die Preise auf der Vertikal- und die Mengen auf der Horizontalachse abtragen. Sind die Preise hoch, ist die Nachfrage gering. Sind die Preise niedrig, ist die Nachfrage hoch. Also fällt die Nachfragekurve und stellt so das Gesetz der Nachfrage dar. Normalerweise drücken wir Preise in Geld aus, weil Geld jede Art von möglichem Konsum und damit alle Abwägungen zwischen verschiedenen Konsumoptionen ausdrückt, denen sich ein Mensch gegenübersehen kann. Hier nehmen wir uns jedoch die Freiheit, den Käsepreis weiterhin in Weineinheiten auszudrücken, um daran zu erinnern, dass die Entscheidungstheorie, auf der wir das Konzept der Nachfrage aufbauen, von Geld ganz unabhängig ist.

##### b. Aggregierte Nachfragekurven

- 121 Nutzt man das Konzept der Nachfrage, betrachtet man aber üblicherweise Märkte, auf denen viele Individuen aktiv sind. Die für ein Individuum

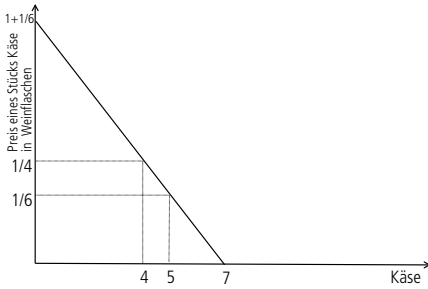


Abbildung 3.9: Nachfrage eines Individuums

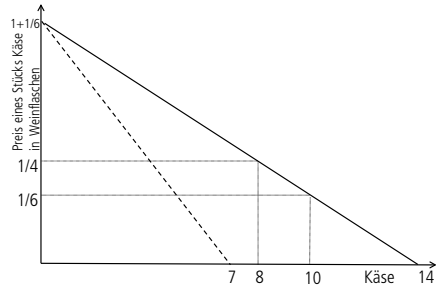


Abbildung 3.10: Aggregierte Nachfrage (zwei Individuen)

konstruierte Nachfragekurve kann leicht auf viele Individuen ausgedehnt werden: Wir addieren alle Mengen, die von Individuen zu einem bestimmten Preis gekauft werden, und konstruieren daraus eine neue Nachfragekurve, die aussagt, wie viel eine Gruppe von Personen zu einem bestimmten Preis konsumiert. Nehmen wir die individuelle Nachfrage in Abb. 3.9 und stellen uns vor, es gäbe zwei Personen, deren jeweilige Nachfrage durch diese Funktion dargestellt wird. Nun erstellen wir eine aggregierte Nachfragefunktion, die den Konsum beider darstellt. Bei einem Preis von einer viertel Flasche Wein pro Stück Käse konsumieren beide Individuen je vier Stücke Käse. Zusammen konsumieren sie also acht. Bei einem Preis von einem Sechstel konsumieren beide Individuen je fünf Einheiten, zusammen also zehn. Wenn der Preis auf 0 fällt, konsumieren beide je sieben Einheiten; zu diesem Preis ist die Nachfrage also vierzehn. Nun zeichnen wir eine neue Linie durch diese Punkte und erhalten die aggregierte Nachfragefunktion für beide Individuen in Abb. 3.10.

### c. Nachfrageelastizität

Die Steigung der Nachfrage sagt etwas darüber aus, wie stark die Nachfrage auf Preisänderungen reagiert. Fällt eine Nachfragefunktion steil ab, bedeutet das, dass sich die nachgefragte Menge nicht sehr verändert, wenn der Preis sich ändert. Um das zu illustrieren, kann man sich eine vertikale Nachfragefunktion vorstellen: Egal, wie hoch der Preis ist, der Konsument kauft stets dieselbe Menge. Wenn die Nachfrage nur wenig auf Preisänderungen reagiert, sprechen wir von **Inelastizität** (s. Abb. 3.11). Dementsprechend sieht eine Nachfrage, die sehr leicht auf Preisänderungen reagiert, flach aus und wird „**elastisch**“ genannt (s. Abb. 3.12).

122

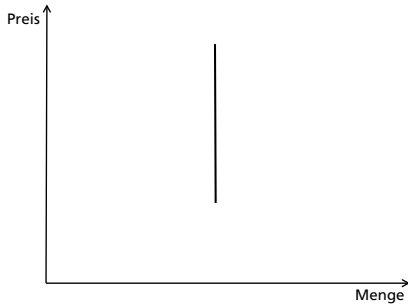


Abbildung 3.11: Vollständig inelastische Nachfrage – die nachgefragte Menge reagiert nicht auf Preisveränderungen.

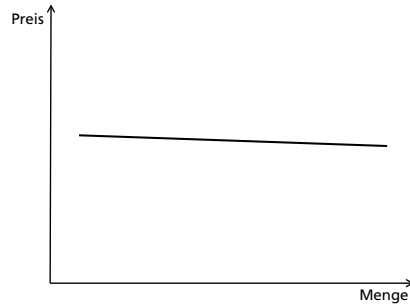


Abbildung 3.12: Elastische Nachfrage – eine kleine Preisänderung verändert die nachgefragte Menge stark.

#### d. Arbeiten mit der Nachfragekurve

- 123 Analysieren wir die Welt mit Hilfe von Nachfragekurven, übersetzen wir reale Phänomene entweder in eine Bewegung entlang der Nachfragekurve oder aber in eine Bewegung der Nachfragekurve selbst. Bewegungen entlang der Nachfragekurve stellen im Grunde Preisänderungen dar. Sie bestehen normalerweise aus einer Anwendung der Nachfragegesetze. Steigt der Preis, fragen Konsumenten weniger nach. Fällt der Preis, fragen Konsumenten mehr nach. Eine Bewegung entlang der Nachfragekurve wäre in Abb. 3.9 dargestellt, wenn man sie nutzte, um eine Preissenkung von  $1/4$  auf  $1/6$  zu analysieren. Wir könnten ablesen, dass bisher 4 Einheiten konsumiert würden und künftig 5 Einheiten.
- 124 Bewegungen der Nachfragekurve selbst drücken eine Änderung der Wertschätzung des Gutes aus. Ein typischer Grund für Verschiebungen der Nachfragekurve ist eine Änderung im Budget der Individuen. Wenn jemand mehr Wein zur Verfügung hat, um diesen gegen Käse einzutauschen, ist er oder sie auch bereit, mehr Wein für dieselbe Menge Käse auszugeben. Dies steigert die Nachfrage, d.h. die ganze Nachfragekurve verschiebe sich in nördliche Richtung. Ein weiteres Beispiel für die Bewegung der Nachfragekurve ist auch in Abb. 3.10 zu sehen. Hier ist die Bewegung der Kurve nicht das Resultat einer Budgetänderung, sondern der Tatsache geschuldet, dass eine zweite Person den Markt betritt. Wenn jetzt der Käsepreis ein Sechstel Weinflaschen beträgt, ist die Gesamtnachfrage nicht mehr fünf Käsestücke, sondern nunmehr zehn.

## e. Konsumentenrente

Aus den Nachfragekurven lässt sich ableiten, was die Konsumenten bei den Markttransaktionen gewinnen. (Nun werden wir endlich Geld einführen, um Preise auszudrücken.) Stellen wir uns zunächst die Nachfragefunktion einer Person vor. Bisher haben wir aus ihr abgelesen, wie viel eine Person zu welchem Preis kauft; ist der Preis sehr hoch, kauft die Person zum Beispiel lediglich ein Gut. Dies sagt uns aber auch, dass die Person die erste Einheit des Gutes genügend wertschätzte, um einen hohen Preis für sie zu zahlen. Wir können also an der Nachfragefunktion ablesen, wie viel eine Person maximal für eine Einheit zu zahlen bereit ist, das heißt, wie hoch die **Zahlungsbereitschaft** einer Person liegt.<sup>4</sup> Für die erste Einheit ist die Person viel zu zahlen bereit, für die zweite dann etwas weniger, noch weniger für die dritte und so fort. Wenn die Person nun eine Menge „q“ (für *quantity*) zu einem mittleren Preis „p“ erwirbt, kann man sagen, dass sie durch den Kauf um die Differenz zwischen ihrer Zahlungsbereitschaft und dem Preis reicher geworden ist. Die Fläche unterhalb der Nachfragekurve und oberhalb des Preises steht daher für das, was der Konsument mit der Transaktion verdient hat. Diese Fläche wird Konsumentenrente genannt (s. Abb. 3.13). Wird diese Fläche größer, können wir daraus schließen, dass der Nutzen des Konsumenten steigt. Wird sie kleiner, sinkt der Nutzen.

Wir können dieselbe Analyse mit Nachfragefunktionen vieler Personen durchführen, um Märkte zu betrachten. Wenn man die Marktnachfrage anstelle der Nachfrage eines Individuums verwendet, stellt die Fläche die Konsumentenrente für den Markt dar. Die Konsumentenrente eines Marktes ist also die Summe der in Geld gemessenen, durch Handel generierten Nutzenzuwächse der Konsumenten. Aber kann man Nutzenzuwächse verschiedener Individuen einfach addieren? Wir hatten festgestellt, dass der Nutzenvergleich zwischen Individuen sinnlos ist. Dann kann man aber auch nicht addieren. Man weiß ja nicht, ob ein Nutzenzuwachs von einem Euro mir genauso viel Nutzen bringt wie ein Nutzenzuwachs von einem Euro Ihnen. Wenn man aber nicht weiß, ob 1 immer gleich groß ist, weiß man auch nicht, ob 1+1 wirklich 2 ergibt. Das Problem verliert an Bedeutung, wenn man sicher ist, dass alle Konsumenten durch eine bestimmte Intervention gewinnen. Dann kann man der Ansicht sein, die Maximierung der Konsumentenrente würde nur sicherstellen, dass jeder Konsument so viel gewinne wie möglich. Aber es ist prinzipiell möglich, dass die Konsumentenrente insgesamt wächst, einige Verbraucher aber schlech-

<sup>4</sup> Hierzu Rz. 71.

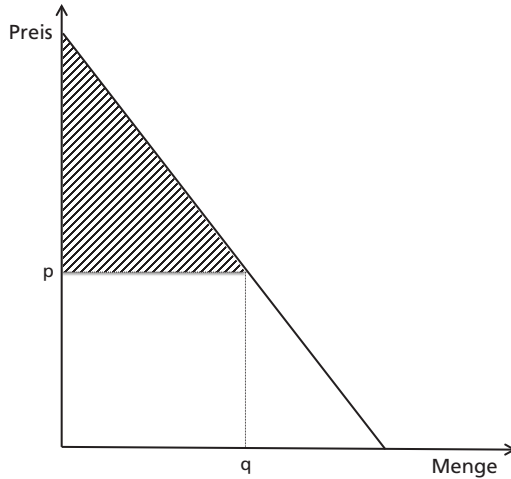


Abbildung 3.13: Konsumentenrente

ter gestellt sind. Es könnten etwa Produkte plötzlich nur noch als Bündel angeboten werden (etwa Wein und Käse – oder realistischer: Laptops und DVD-Laufwerke), so dass der Bündelpreis sänke, die Produkte einzeln aber nicht mehr zu erwerben wären. Dann ginge es den Konsumenten, die beide Güter bräuchten, besser. Denen, die nur eines bräuchten, ginge es aber unter Umständen schlechter. In der Wettbewerbspolitik wird der Einfachheit halber und als akzeptierter Näherungswert dennoch die Konsumentenrente als Maßstab herangezogen, um die Wohlfahrt zu bemessen, die von Märkten generiert wird. Nach dem **Konsumentenwohlfahrtsstandard** ist eine Zuteilung dann optimal, wenn die Konsumentenrente, maximiert wird.

#### f. Beispiel Flughafen (1)

- 127 Wir unterbrechen kurz, um einiges von dem, was wir gelernt haben, am Beispiel des Rechtsstreits zwischen Flughafenanwohnern und Flughafen anzuwenden. Aus ökonomischer Sicht wird eine ruhige Nacht so hoch geschätzt, wie die Anwohner bereit sind, dafür anderes aufzugeben. Die Möglichkeiten oder Optionen, die sie für eine ruhige Nacht aufgeben können, können grundsätzlich in Geld dargestellt werden, da das Geld, das einem Individuum zur Verfügung steht, die meisten Konsummöglichkeiten dieser Person repräsentiert. Nehmen wir an, die Anwohner können vor

Gericht gegen den Flughafen eine Unterlassungsverfügung erwirken und hiermit verhindern, dass der Flughafen nachts geöffnet bleibt. Nun stellen wir uns zwei Szenarien vor:

Im ersten entgehen dem Flughafen hohe Profite, wenn er nachts schließen muss. In diesem Fall könnte es sich der Flughafen leisten, durch Geld so viele Konsummöglichkeiten an die Anwohner zu transferieren, dass diese die transferierten Möglichkeiten höher schätzen als ruhige Nächte. Im Gegenzug würde der Flughafen verlangen, dass die Anwohner das Recht auf Erlass einer Unterlassungsverfügung an den Flughafen transferieren. In der Folge könnte der Flughafen auch nachts betrieben werden. Die Anwohner verlören ihre ruhigen Nächte, bekämen aber etwas, das sie höher schätzten. Der Flughafen stünde ebenfalls besser als bei einem Nachtflugverbot, weil er zumindest einen Teil seiner hohen Nachtflugprofite behalten könnte. Dies wäre eine *Pareto*-Verbesserung.<sup>5</sup>

128

Im zweiten Szenario entgehen dem Flughafen nur geringe Profite, wenn er nachts schließen muss. In diesem Fall würden die Profite, auf die er verzichten müsste, nicht ausreichen, um jene Konsummöglichkeiten an die Anwohner zu transferieren, die diese mehr schätzen als ruhige Nächte. Die Anwohner behielten daher ihr Recht auf Erlass einer Unterlassungsverfügung, würden dieses geltend machen und der Flughafen würde nachts schließen. Das ist *pareto*-effizient, weil es keine *Pareto*-Verbesserung darstellt, wenn dem Flughafen der Nachtbetrieb erlaubt würde, ohne dass die Anwohner entschädigt würden.

129

An diesem Beispiel wird auch deutlich, dass einerseits kaum in Erfahrung zu bringen sein wird, wie hoch die Wertschätzung ruhiger Nächte durch Anwohner ist, da „Wertschätzung“ in der Ökonomie eine fundamental relative und subjektive Angabe ist. Andererseits könnte ein funktionierender Markt gewährleisten, dass eine effiziente Allokation erreicht wird, wenn das Gesetz ein handelbares Recht schafft. Dann können Marktteilnehmer durch ihr Handeln ihre Wertschätzung offenbaren. Wir kommen am Ende dieses Kapitels auf diesen Punkt noch einmal zurück, denn der funktionierende Markt ist eine anspruchsvollere Einrichtung, als es in diesem Beispiel scheinen mag.

130

<sup>5</sup> Vgl. zu dem Begriff „*Pareto*-Verbesserung“ oder „*Pareto*-Superiorität“ Rz. 90.

### III. Angebot

- 131 Kommen wir zurück zu Ihrem ruhigen Abend mit Käse und Wein. Beide Produkte aus unseren Beispielen müssen irgendwoher kommen. Die Produkte werden von Unternehmen geliefert, deren Ziel es ist, Profit zu machen, indem sie ihre Produkte herstellen und verkaufen. Wir können uns diese Unternehmen als besondere Menschen vorstellen, die keine Präferenzen haben, jedoch auf ein anderes Ziel hinarbeiten: Profitmaximierung. Der zentrale Begriff, auf dem das Angebot ruht, ist der Kostenbegriff. Firmen können Güter nicht aus der hohlen Hand herstellen. Bei der Produktion haben sie Kosten. Insgesamt wollen sie wenig bezahlen und viel einnehmen. In diesem Abschnitt blicken wir jedoch zunächst allein auf die Kosten und verrechnen sie noch nicht mit den Einnahmen.

#### 1. Opportunitätskosten

- 132 Bei dem Kostenbegriff, den Ökonomen verwenden, handelt es sich um sogenannte **Opportunitätskosten**. Produktionskosten sind die Gewinnmöglichkeiten der besten alternativen Investitionsmöglichkeit (Opportunität), auf die ein Unternehmen verzichten muss, wenn es seine Mittel nun konkret in die Herstellung eines bestimmten Projekts investiert. Stellen Sie sich vor, Ihr Weinunternehmen erbt von Ihrer Großmutter das Château der Familie. Da Sie es umsonst erhalten haben (Steuern lassen wir hier außer Betracht), könnten Sie versucht sein, zu glauben, dass es Sie nichts kostet, das Château zur Weinherstellung zu nutzen. So denken Ökonomen aber nicht. Sie würden Ihnen stattdessen erklären, dass Sie, indem Sie das Château zur Weinherstellung verwenden, auf Einnahmen durch Tourismus verzichten oder auf die Möglichkeit, das Haus an einen anderen Weinproduzenten zu vermieten. Die Kosten der Nutzung des Châteaus sind ungefähr so hoch wie die Miete, die Sie daran verdienen könnten, wenn Sie es nicht selbst nutzten.
- 133 Einige deutsche Rechtswissenschaftler und Politiker hatten diese Logik nicht vollständig verinnerlicht, als der EU-Emissionshandel in Kraft trat. Dieser sieht vor, dass Unternehmen für jede Tonne ausgestoßenes CO<sub>2</sub> Zertifikate erwerben müssen, welche in einer von der EU begrenzten Anzahl am Markt gehandelt werden. Die Unternehmen haben somit die Wahl, Zertifikate zu kaufen oder ihre Emissionen zu reduzieren. Um zu verhindern, dass die Energielieferanten daraufhin ihre Preise für den Verbraucher erhöhen, wurden ihnen kostenlose Emissionszertifikate zugeteilt. Dennoch hoben sie in der Folge die Energiepreise an, weswegen manche

Juristen das Kartellamt aufforderten, den Preiserhöhungen Einhalt zu gebieten. Das Argument der Juristen: Marktmächtige Unternehmen – wie die großen Energieversorger – dürften ihre Preise gemäß Art. 102 AEUV nicht losgelöst von ihren Kosten anheben. Die kostenlos zugeteilten Zertifikate würden gerade keine Kostensteigerung darstellen, weswegen eine Preissteigerung nicht gerechtfertigt sei. Nun, da Ihnen das Konzept der Opportunitätskosten vertraut ist, erkennen Sie, dass die Verwendung der Zertifikate zur Energieproduktion gleichbedeutend mit dem Verzicht auf die Möglichkeit ihres Verkaufs ist. Die Verwendung eines Zertifikats, um CO<sub>2</sub> ausstoßen zu dürfen, würde also genau den Marktpreis des Zertifikats kosten, unabhängig davon, ob das Unternehmen für die Anschaffung des Zertifikates zahlen musste oder nicht. Demnach wurden die Preise tatsächlich in Reaktion auf eine Kostensteigerung erhöht. Das Kartellrecht wurde somit nicht verletzt.

## 2. Einige weitere wichtige Kostenbegriffe

Im Abschnitt über Nachfrage haben wir Geld als Medium interpretiert, 134  
welches (annähernd) jede mögliche Konsumform repräsentiert. Im Folgenden werden wir Geld als Medium interpretieren, das alle Investitionsmöglichkeiten repräsentiert. Eine weitere Vereinfachung, die wir anwenden, ist die Kostenminimierung. Eine spezifische, durch ein Unternehmen an den Markt gebrachte Menge eines Gutes (Output-Menge) kann zu sehr verschiedenen Kosten produziert werden. Denken wir an das Château zurück. Um Reben anzubauen, braucht es Arbeitskräfte und Maschinen. Falls Sie nur Arbeitskräfte einstellen, sieht das vielleicht sehr romantisch aus, aber Ihre Leute werden lange brauchen, bis sie die Ernte eingefahren haben. Stellen Sie nur eine Arbeitskraft mit vielen Maschinen ein, wird diese Person auch nicht viel schneller vorankommen. Sie kann nämlich zum Beispiel nicht den LKW nach Hause fahren und zeitgleich Wein lesen. Schließlich können Sie einige Maschinen verwenden und mehrere Arbeiter einstellen, was aller Wahrscheinlichkeit nach zu niedrigeren Kosten führen wird als die beiden vorherigen Kombinationen. Da Sie Ihr Château mit dem Ziel der Profitmaximierung betreiben und alle eingesparten Kosten den Profit erhöhen, liegt es in Ihrem Interesse, für eine gegebene Output-Menge die Kosten so gering wie möglich zu halten. Für die ökonomische Analyse bedeutet dies eine erhebliche Vereinfachung. Da wir wissen, dass Sie sich als Unternehmen bemühen, Ihre Kosten zu minimieren, müssen wir nur die Minimalkosten für jede mögliche Output-Menge in



Betracht ziehen. Dies resultiert in einer Kostenfunktion, bei der jede Output-Menge genau einem Kostenniveau entspricht.

- 135 Um zu sehen, welche Kostenarten mit der Herstellung eines Gutes verbunden sind, denken wir einmal mehr an das Château. Wenn Sie sich dafür entscheiden, nur ein Fass Wein pro Jahr herzustellen, beschränken Sie sich womöglich auf einen sehr kleinen Weinberg; Sie verwenden nur einen Teil des Châteaus zur Weinherstellung und können die anderen Teile, die Sie nicht verwenden, vermieten; auch können Sie den Betrieb allein führen, in nur ein paar Stunden pro Woche, und den Rest Ihrer Zeit investieren Sie in Freizeit oder Sie verdienen noch als Anwalt. Mit anderen Worten: Der Betrieb Ihres Unternehmens kostet Sie nicht viel. Wollen Sie hingegen die Weinproduktion ausbauen, steigen auch Ihre Kosten. Sie müssen die bestellte Parzelle ausweiten und deshalb auf die sich sonst daraus ergebenden Miet- bzw. Pachteinnahmen verzichten. Einen größeren Teil Ihres Châteaus verwenden Sie nun für die Weinherstellung und verzichten somit auf Einnahmen durch Besucher. Sie arbeiten mehr im Weinberg und geben ihre Tätigkeit als Anwalt auf. Vielleicht müssen Sie gar eine Hilfskraft einstellen.
- 136 Sie bemerken vielleicht schon, dass es verschiedene Arten von Kosten gibt, zwischen denen man unterscheiden kann. Manche Kosten ändern sich mit der Mengenausweitung, und manche nicht. Sie können den Keller, in dem die Fässer lagern, nicht als Restaurant benutzen – unabhängig davon, ob Sie nun drei Fässer herstellen oder kurzfristig auf 30 erhöhen. Diese entgangenen Einnahmen variieren nicht. Sie bleiben unabhängig von der produzierten Menge gleich. Auch brauchen Sie denselben Traktor – unabhängig davon, ob Sie drei oder 30 Fässer herstellen. Die Kosten für Maschinen variieren oft nicht mit der produzierten Menge. Diese Kosten, die von der produzierten Menge weitgehend unabhängig sind, heißen **Fixkosten**. Die Arbeit kann sich jedoch drastisch vermehren, wenn Sie statt der drei nun 30 Fässer herstellen. Kosten, die je nach produzierter Menge variieren, werden *variable Kosten* genannt. Arbeitskosten sind dafür das typische Beispiel. Die Unterscheidung zwischen fixen und variablen Kosten hat eine zeitliche Dimension. Auf kurze Sicht ist diese Unterscheidung vollkommen angemessen. Langfristig jedoch ändern sich alle Kosten mit der Herstellungsmenge, so dass alle Kosten variabel sind. Mit kleineren Output-Änderungen von einer Woche auf die nächste wird sich die Anzahl der Maschinen nicht verändern; hingegen werden Sie durchaus mehr Maschinen benötigen, wenn Ihre Produktion über die nächsten drei Jahre erheblich ansteigt. Wir werden uns zunächst auf die kurze Frist konzentrieren.

Es wird sich als nützlich erweisen, die Kosten zu betrachten, welche einer bestimmten Output-Einheit zugeschrieben werden können. Zu diesem Zweck kann man angeben, wie viel eine Output-Einheit bei einem gegebenen Output-Level durchschnittlich kostet (**durchschnittliche Kosten**). Wir können uns ebenfalls ansehen, welche dieser Kosten pro Einheit fix sind (**durchschnittliche Fixkosten**) und welche variabel sind (**durchschnittliche variable Kosten**). Nun gilt es zu überlegen, wie sich diese Kosten bei einem Anstieg des Outputs entwickeln würden. Die Fixkosten verändern sich nicht mit dem Output, weshalb sie einfach durch eine größere Anzahl von Output-Einheiten geteilt werden. Dies senkt die Fixkosten pro Einheit. Also gehen die durchschnittlichen Fixkosten zurück, wenn die Menge steigt. Denken Sie an den Traktor, den Sie gekauft haben, um den Wein anzubauen. Dieser hat Sie € 10.000 gekostet. Wenn er zehn Jahre hält und Sie mit seiner Hilfe pro Jahr drei Fässer herstellen, kostet Sie der Traktor im Schnitt etwa € 333 pro Fass. Stellen Sie hingegen 30 Fässer im Jahr her, kostet Sie der Traktor nur € 33 pro Fass.

Wie sich die durchschnittlichen variablen Kosten mit dem Output verändern, ist weniger klar. Letztlich nimmt man an, dass die Aufwendung variabler Kosten anfangs eine starke Output-Erhöhung nach sich zieht, während die Wirksamkeit der Aufwendungen in der Folge abnimmt, so dass spätere Aufwendungen nur noch zu kleineren Output-Erhöhrungen führen. Der Gedanke dahinter ist, dass die Effektivität dieser Investitionen irgendwann durch die Fixkosteninvestitionen beschränkt wird. Denken wir an den Keller. Es mag einfach sein, die Produktion von drei auf zehn Fässer hochzuschrauben. Ab einem bestimmten Zeitpunkt wird es jedoch eng im Keller werden: Die Arbeit an den Fässern würde durch die Enge erschwert. Daher nimmt man an, die durchschnittlichen variablen Kosten stiegen mit wachsendem Output. Die Kombination aus den sinkenden durchschnittlichen Fixkosten und den steigenden durchschnittlichen variablen Kosten ergibt eine U-förmige Funktion, die den Output zu den durchschnittlichen Kosten in Beziehung setzt (Abb. 3.14).

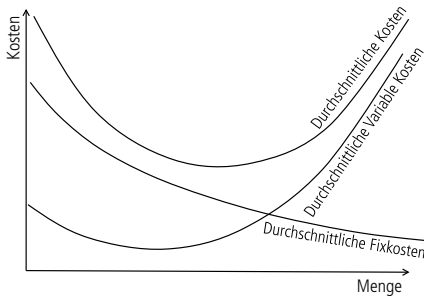


Abbildung 3.14: Durchschnittliche Kosten

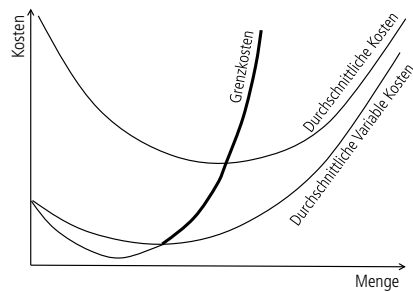


Abbildung 3.15: Grenzkosten und Angebot

- 139 Darüber hinaus ist auch die Unterscheidung zwischen **sozialen Kosten** und **privaten Kosten** wichtig. Alle in diesem Abschnitt über Nachfrage eingeführten Kosten sind private Kosten. Sie werden allein vom Unternehmen getragen. Soziale Kosten sind die Summe aller Kosten, die eine bestimmte Aktivität den verschiedenen Mitgliedern einer bestimmten ökonomischen Gemeinschaft auferlegt. Für die Effizienzanalyse kommt es auf die sozialen Kosten an.

### 3. Spezielle Kosten und die Angebotskurve

- 140 Der letzte – und wichtigste – Begriff aus dem Bereich der Kosten sind die **Grenzkosten**, die auch als **marginale Kosten** bezeichnet werden. Die Grenzkosten drücken die Steigung der Gesamtkostenfunktion an einem bestimmten Punkt aus. Die Steigung weist darauf hin, wie stark eine kleinere Änderung im Output die Gesamtkosten verändern kann. Wir können Grenzkosten also deuten als die Herstellungskosten der nächsten zusätzlichen Einheit. Mathematisch findet man die Steigung einer Funktion in ihrer ersten Ableitung. Daher ist die Grenzkostenfunktion die erste Ableitung der Gesamtkostenfunktion.
- 141 Wir erwähnten bereits, dass von variablen Kosten letztlich angenommen wird, sie stiegen mit steigender Herstellungsmenge. Das liegt daran, dass die Effektivität zusätzlicher Investitionen in variable Kosten durch die Investition in Fixkosten beschränkt wird. Daraus schließen wir, dass auch die Grenzkosten irgendwann steigen, weil die Interpretation von Grenzkosten derjenigen der variablen Kosten sehr ähnlich ist (Kosten der Produktion der nächsten Einheit). Dies ist deshalb so wichtig, weil jeder Hersteller nur dann eine zusätzliche Einheit herstellen wird, wenn die Kosten dieser Einheit geringer sind, als die durch diese generierten Einkünfte (Grenzer-

trag). Um die Output-Menge zu bestimmen, die ein Hersteller an den Markt bringen wird, muss man daher die Menge finden, an der die Grenzkosten dem Grenzertrag entsprechen.

Dies ist exakt dieselbe Technik, die wir bei der Analyse von Konsumentenentscheidungen als Nutzenmaximierung angewendet haben. Dort suchten wir nach der Menge, bei der der Grenznutzen gleich dem Preis war. Der Grenznutzen entspricht dem, was wir nun Grenzertrag nennen, und der Preis entspricht den Grenzkosten. Das Gleichsetzen von Grenzertrag und Grenzkosten ist eine Standardmethode der Maximierung und begegnet Ihnen in der Ökonomie, in der man häufig Kosten-Nutzen-Analysen betreibt, deshalb immer wieder.

142

Die Bedeutung der Grenzkostenfunktion geht noch weiter. Ihr Teil, welcher oberhalb des Schnittpunkts mit der Funktion der durchschnittlichen variablen Kosten liegt (der fettgedruckte Teil der Grenzkostenkurve in Abb. 3.15), stellt die **Angebotskurve** dar. Grundsätzlich wird das Unternehmen so lange Güter an den Markt bringen, bis die nächste hergestellte Einheit mehr kostet, als der Preis einbringt. Die Kosten der nächsten Einheit lassen sich an der Grenzkostenkurve ablesen. Zu Preisen unterhalb der Kurve, die die durchschnittlichen variablen Kosten darstellt, wird das Unternehmen aber nicht produzieren. Wenn die durchschnittlichen variablen Kosten über dem Preis liegen, kann das Unternehmen nicht einmal seine variablen Kosten einwerben – selbst wenn die Marge bei den letzten Einheiten positiv ist, also einen Ertrag ausweist (nicht fettgedruckter Teil der Grenzkostenkurve in Abb. 3.15). Wenn das Unternehmen aber nicht kostendeckend anbieten kann, wird es gar nicht aktiv werden. Daher repräsentiert die Grenzkostenkurve nur insoweit die Angebotskurve, als sie über den durchschnittlichen variablen Kosten liegt (fettgedruckter Teil der Grenzkostenkurve in Abb. 3.15). Für den Teil oberhalb der durchschnittlichen variablen Kosten gibt die Grenzkostenkurve an, wie viel ein Anbieter bereit ist, zu einem bestimmten Preis an den Markt zu bringen. Wir können also die Angebotsmenge aus der Grenzkostenkurve ablesen, sofern wir den Preis kennen.

143

Darüber hinaus ist es möglich, verschiedene Angebotskurven einzelner Anbieter in einer Marktangebotskurve zusammenzufassen, indem wir die Output-Mengen der Anbieter, die bei einem bestimmten Preisniveau produziert werden, addieren. Dies ist exakt dieselbe Operation, wie wir sie bereits weiter oben durchgeführt haben, als wir alle Mengen addierten, welche die Marktteilnehmer zu einem bestimmten Preis nachgefragt haben. Auch hier ist sie nützlich, weil man mit dem Angebotskonzept üblicherweise Märkte analysieren will, auf denen mehr als ein Anbieter aktiv ist.

144

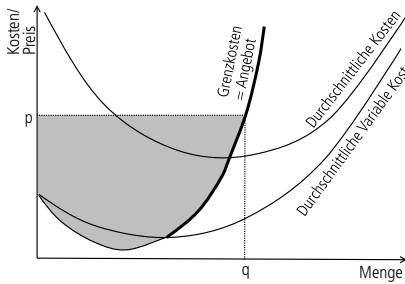


Abbildung 3.16: Produzentenrente

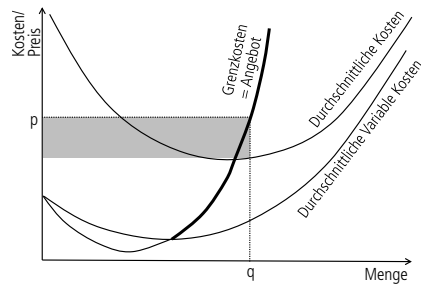


Abbildung 3.17: Profit

- 145 So, wie wir bezüglich der Nachfrage von Elastizität sprechen konnten, können wir auch im Kontext des Angebots von Elastizität sprechen. Das Angebot ist elastisch, wenn es stark auf Preisänderungen reagiert (also ist die Angebotsfunktion flach). Tut es dies nicht, ist das Angebot unelastisch (also ist die Angebotsfunktion steil).

#### 4. Produzentenrente

- 146 Die Produzentenrente berechnet sich aus dem Umsatz (Preis mal Menge) abzüglich der Grenzkosten. Die Logik ist ähnlich wie bei der Konsumentenrente. Die Grenzkostenkurve zeigt an, zu welchem Minimalpreis der Hersteller das Gut an den Markt gebracht hätte. Die Tatsache, dass er für das Gut mehr erhält, bedeutet, dass er am Verkauf verdient. Daher stellt die Fläche oberhalb der Angebotskurve bis zur Menge „q“ und unterhalb des Preises „p“ die Produzentenrente dar (s. Abb. 3.16). Diese Produzentenrente sollte jedoch nicht mit dem Profit verwechselt werden (s. Abb. 3.17). Obwohl der Hersteller eine Produzentenrente erwirtschaftet, kann er insgesamt Verluste machen, die ihn zwingen, den Markt zu verlassen. Die Fixkosten sind in der Produzentenrente nicht enthalten. Wenn also die Fixkosten zu hoch sind, können sie die Produzentenrente übersteigen, und der Hersteller fährt einen Verlust ein. Profite sind Umsatz minus durchschnittliche Kosten.

### IV. Der Markt

- 147 Nachdem wir Angebot und Nachfrage isoliert betrachtet haben, können wir diese in einem Markt zusammenführen und betrachten, wie sie zu-

sammenwirken. Zuvor müssen wir jedoch genauer beschreiben, wie wir uns diesen Markt vorstellen.

### 1. Perfekter Wettbewerb

Das Grundmodell eines Markts in der Ökonomie ist der perfekte Wettbewerb. Perfekter Wettbewerb bedeutet, dass kein am Markt auftretender Agent den Preis beeinflussen kann. Die Gesamtheit aller eigennützigen Aktionen machtloser Agenten generiert den Preis. Wir stellen uns eine unendliche Menge von Konsumenten am Markt vor, die jeweils einen winzigen Bruchteil der Nachfrage ausmachen. Sie alle maximieren ihren Nutzen. Ebenso stellen wir uns eine unendliche Zahl von Verkäufern vor, die jeweils einen winzigen Bruchteil des Angebots ausmachen und allesamt ihren Profit maximieren. Die Tatsache, dass wir uns jeden Agenten vorstellen als lediglich einen winzigen Bruchteil seiner Marktseite, bedeutet Folgendes: Jedes Mal, wenn ein Verkäufer höher anbietet als die anderen, findet er keinen Käufer mehr; umgekehrt verkaufen die Verkäufer an andere Käufer, wenn ein Käufer niedriger bietet als alle anderen. Dies repräsentiert die Tatsache, dass kein einzelner Akteur im perfekten Wettbewerb den Preis beeinflussen kann. Alle Agenten sind **Preisnehmer**. Dies passt zur Vorstellung eines Markts, der einen einzigen Preis ergibt. Implizit folgen wir dieser Annahme bisher schon im gesamten Kapitel. 148

Wo liegt also im perfekten Wettbewerb der Preis eines Gutes? Wir wissen, dass der Preis gleich sein wird mit den Grenzkosten der Herstellung (d. h. der Preis liegt auf der Angebotskurve), da jeder Anbieter seinen Output so lange erhöhen wird, bis die letzte produzierte Einheit genau so viel einbringt, wie ihre Bereitstellung kostet. Zugleich wissen wir, dass der Preis gleich der marginalen Zahlungsbereitschaft sein wird (d. h. er liegt auf der Nachfragekurve), denn jeder Konsument wird so lange weiter kaufen, bis er mit dem Preis, den er bezahlt, genau den Nutzen aufgibt, den er durch das gekaufte Gut erhält. Weil der Preis „p“ also zugleich den Grenzkosten der Bereitstellung und dem Grenznutzen des gehandelten Gutes entsprechen muss und eine eindeutig bestimmte Menge „q“ diesem Preis entsprechen muss, sind Preis und Menge an dem Punkt abzulesen, an dem sich die Nachfragekurve (marginale Zahlungsbereitschaft) und die Angebotskurve (Grenzkosten der Herstellung) schneiden. Dieser Punkt gibt uns ebenfalls die Menge „q“ an, zu der das Gut an den Markt gebracht und konsumiert wird (s. Abb. 3.18). 149

Ist dieses Ergebnis effizient? Als Maßstab für Effizienz am Markt nehmen wir typischerweise die Summe von Konsumentenrente und Produzenten- 150

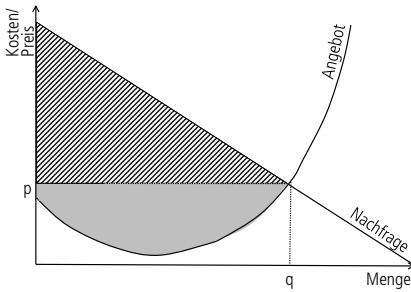


Abbildung 3.18: Preis unter perfektem Wettbewerb

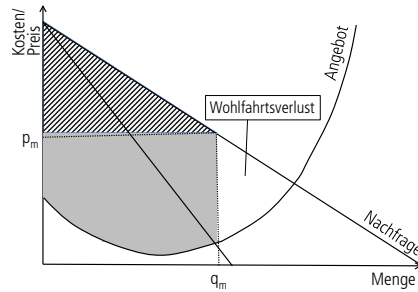


Abbildung 3.19: Preis unter Monopol

tenrente (Gesamtrente). Wird sie maximiert, ist die Allokation effizient im Sinne des *Kaldor-Hicks-Kriteriums*<sup>6</sup>. Will man prüfen, ob der Markt das Maximum, welches Konsumenten und Produzenten insgesamt hätten „verdienen“ können, realisiert hat, so hilft ein Blick auf das Diagramm. Dort sehen wir, dass ein Preis, der gleich den Grenzkosten und der marginalen Zahlungsbereitschaft ist, den gemeinsamen Überschuss maximiert. Läge der Preis höher, kauften die Nachfrager weniger: Wir müssten die diesem Preis entsprechende Menge am Markt an der Nachfragekurve ablesen und die umgesetzte Menge nähme ab. In dem Fall wäre ein kleiner Teil der Fläche, welche den Überschuss darstellt, nicht realisiert worden und daher verloren gegangen. Das ist der Grund, warum wir das verlorene Dreieck als *deadweight loss* bzw. **Wohlfahrtsverlust** bezeichnen (Abb. 3.19). Wenn wir den Preis künstlich unter den Wettbewerbspreis senkten, brächten die Produzenten eine geringere Menge an den Markt: Wir müssten die entsprechende Menge an der Angebotskurve ablesen, wieder nähme die umgesetzte Menge ab und ein Teil der Gesamtrente, die im Wettbewerb realisiert würde, fiel nicht an. Also ist die Wohlfahrt unter perfektem Wettbewerb maximal. Nebenbei merken wir an, dass eine Zuteilung, die von einem Markt generiert wird, auf dem perfekter Wettbewerb herrscht, auch *pareto-effizient* ist: Niemand kann bessergestellt werden, ohne dass ein anderer schlechtergestellt würde. Dies ist das sogenannte **erste Wohlfahrtstheorem**: Jedes Marktgleichgewicht unter perfektem Wettbewerb ist *pareto-effizient*. In der Literatur wurde gezeigt, dass theoretisch ein allgemeines Gleichgewicht möglich ist, in dem sich jeder Markt einer Volkswirtschaft unter perfektem Wettbewerb im Gleichgewicht befindet und dies auch ein Gleichgewicht zwischen allen Märkten darstellt. Weiter un-

<sup>6</sup> Vgl. Rz. 93 ff.

ten in diesem Kapitel werden wir einige der Gründe sehen, warum Theorie und Praxis hier auseinanderfallen. Wir verzeichnen auch, dass *Pareto-Effizienz* nichts aussagt über die Überschussverteilung zwischen Konsumenten und Produzenten. Ob Konsumenten- oder Produzentenrente größer ist, hängt von der jeweiligen Steigung von Angebot und Nachfrage ab.

Eigentlich hätte man von Beginn an wissen können, dass ein Preis, der Grenzkosten und Grenznutzen gleichsetzt, den Gesamtüberschuss maximiert. Warum? Weil die Gleichsetzung von Grenzkosten und Grenznutzen eine Maximierungsmethode ist, wie wir sie schon bei der Maximierung der Konsumentenrente und Produzentenrente angewendet haben. Die Gleichsetzung von Grenzkosten und Grenznutzen ergibt ein Maximum – in diesem Falle das Maximum des Gesamtüberschusses gemessen in Geld.

151

## 2. Güter als Bündel von Rechten

Wenn also der Markt „automatisch“ zu effizienten Ergebnissen führt, warum brauchen wir dann in der Marktwirtschaft rechtliche Interventionen? Im folgenden Abschnitt werden wir einige Gründe für Marktversagen kennenlernen, die rechtliche Eingriffe rechtfertigen können. Doch selbst wenn wir annehmen, dass alle Märkte in einer Marktwirtschaft gut funktionieren, brauchen wir das Recht, um zu definieren, mit welchen Gütern an Märkten gehandelt werden kann. Bisher haben wir mit leicht nachvollziehbaren Gütern gehandelt: Wein und Käse. Doch wir haben auch abstraktere Güter gesehen, mit denen gehandelt werden kann, etwa das Recht von Flughafenanwohnern auf Erlass einer Unterlassungsverfügung. Wir sehen, dass das Wort „Gut“ nicht als ein körperlicher Gegenstand oder ein „Produkt“ definiert werden kann. Tatsächlich wurde in unseren Beispielen mit Besitz und Eigentum an Wein bzw. Käse gehandelt. Was uns besonders interessierte, war das Recht auf Verbrauch der Sache, welches mit dem Eigentum einhergeht. Das ist nicht so trivial, wie es auf den ersten Blick erscheinen mag. Viele Märkte existieren gerade deshalb nicht, weil **Verfügungsrechte** entweder nicht definiert oder nicht durchsetzbar sind. Es gibt kein Eigentum an Menschen. Nicht einmal Sie selbst können sich als Sklave veräußern. Güter im Sinne dieses Verfügungsrechteansatzes sind Rechtebündel, die man übertragen kann. Diese Funktion des Rechts ist alles andere als neutral, wie wir etwa an geistigem Eigentum erkennen können. Welche Märkte existieren und wie sie funktionieren, wird von Juristen täglich aufs Neue definiert. Das Recht des geistigen Eigentums hat zur Aufgabe, neue Ideen handelbar zu machen. So kann man sich auf das Erfinden spezialisieren, ohne dass man seine Ideen selbst in die Praxis umset-

152



zen müsste. Das ermöglicht Arbeitsteilung. So landen die Ideen – so kann man hoffen – bei dem Unternehmen, das sie am besten realisieren kann.

## V. Marktversagen

- 153 Wir sagen, dass Märkte versagen, wenn die Zuteilung der Güter, die sie erreichen, ineffizient ist. Die Mittel gegen ein solches Marktversagen bestehen oft in Formen rechtlicher Intervention. Aber bevor wir uns der Frage zuwenden, wie man Marktversagen mit Hilfe des Rechts kompensiert oder gar verhindert, müssen wir zunächst begreifen, warum Märkte versagen. Daher soll das Marktversagen in den folgenden Kapiteln eingehender untersucht werden. Hierbei werden Sie vier unterschiedliche Ursachen kennenlernen: Marktmacht (1.), asymmetrische Information (2.), externe Effekte (3.) und nicht private Güter (4.). Externe Effekte werden uns in § 6 zu *Public Choice* wiederbegegnen.<sup>7</sup> Verschiedene Grade der Marktmacht<sup>8</sup> sowie verschiedene Formen nicht privater Güter<sup>9</sup> werden in § 4 zu Spieltheorie und Gemeinschaftsgütern ausführlicher behandelt, und asymmetrische Information ist ein Kernthema des Kapitels zur Vertragstheorie.<sup>10</sup>

### 1. Märkte ohne Wettbewerb

- 154 Der erste Grund, warum Märkte Güter nicht effizient verteilen, ist mangelnder Wettbewerb. Gibt es am Markt nur einen Verkäufer, kann dieser die Preise über das Wettbewerbsniveau hinaus erhöhen, ohne dass er fürchten muss, die Käufer an einen anderen Anbieter zu verlieren. Seine Preissetzungsmacht ist nicht durch Wettbewerber begrenzt, sondern lediglich durch die Möglichkeit des Käufers, nicht zu kaufen. Innerhalb dieses Spielraums kann er den Preis zu seinen Gunsten festsetzen. Der Monopolist ist daher kein Preisnehmer.
- 155 Da der Monopolist den Preis festsetzen kann, maximiert er seinen Profit anders als ein Preisnehmer im Wettbewerb. Ein Preisnehmer hat keinen Einfluss auf den Preis. Der Preis wird ihm vom Markt vorgegeben. Er setzt durch die Wahl der Menge seine Grenzkosten für die Herstellung mit seinem Grenzertrag (dem Preis pro Einheit) gleich. Der Monopolist setzt auch

<sup>7</sup> Vgl. dort insb. Rz. 329.

<sup>8</sup> Vgl. hierzu das Kartellilemma in Rz. 177 ff. und das Markteintrittsspiel in Rz. 248 ff. Vgl. auch § 5, Rz. 292 ff.

<sup>9</sup> Vgl. hierzu Rz. 221 ff.

<sup>10</sup> Vgl. Rz. 273 ff.

seine Grenzkosten der Herstellung mit seinem Grenzertrag gleich. Aber er weiß, dass er den Marktpreis beeinflusst: Stellt er eine weitere Einheit her, weitet er das Angebot aus, so dass der Preis sinkt. Sein Grenzertrag ist daher der Preis, den er für die nächste Einheit bekommt, abzüglich des Einkommens, das er verliert, weil der Preis sinkt, zu dem er die vorher produzierten Einheiten hätte verkaufen können. Sein Grenzertrag ist also nicht derselbe für jede Einheit. Stellen wir uns ein Zahlenbeispiel vor und nehmen an, die Nachfrage falle um einen Euro für jede weitere an den Markt gebrachte Einheit. Wenn der Monopolist eine Einheit herstellt, liegt der Preis am Markt bei, sagen wir, € 10. Stellt er eine weitere Einheit her, fällt der Preis um einen Euro auf € 9. Durch die Bereitstellung der zweiten Einheit verliert der Monopolist also € 1 auf die erste Einheit, nimmt aber € 9 für die zweite Einheit ein. Sein zusätzlicher Ertrag aus der Herstellung der zweiten Einheit ist daher € 9 (zusätzliche Einnahmen auf die zweite Einheit) minus € 1 (Preisreduktion auf die erste Einheit) – also insgesamt € 8. Die Produktion der dritten Einheit drückt den Preis für alle Einheiten auf € 8. Der Verkauf der dritten Einheit bringt € 8 ein, doch an der ersten und zweiten verliert der Monopolist insgesamt zwei weitere Euro an Einnahmen. Der zusätzliche Ertrag aus der Herstellung der dritten Einheit beträgt also € 6. Man sieht, dass der Grenzertrag fällt (10, 8, 6, ...), und zwar doppelt so schnell wie die Nachfragekurve (10, 9, 8, ...; s. Abb. 3.19). Das ist tatsächlich eine allgemeine Regel. Der Monopolist wird die Grenzkosten der Herstellung mit seinen Grenzeinnahmen gleichsetzen, um seinen Profit zu maximieren (s. Abb. 3.19). Daher wird die Quantität „ $q_m$ “, die im Monopol verkauft wird, immer kleiner sein, als wenn der Wettbewerb die Grenzkosten der Herstellung und den Grenznutzen des Konsums (die Nachfrage) ausgeglichen hätte (s. Abb. 3.18). Entsprechend wird der Monopolpreis „ $p_m$ “ höher sein als der Wettbewerbspreis.

Wenn wir die Gesamtrente des Monopolmarktes mit der des Wettbewerbsmarkts vergleichen (Abb. 3.18 und 3.19), sehen wir, dass sie unter einem Monopol kleiner ist. Wir sehen, dass die Monopolpreisbildung nicht nur Rente von den Konsumenten zum Hersteller verschiebt, sondern zusätzlich Wohlfahrt zerstört, die einfach verschwindet. Die Beseitigung des Wohlfahrtsverlusts durch eine Stärkung des Wettbewerbs, etwa indem die Wettbewerbspolitik höheren Output erzwingt, ist allerdings keine *Pareto*-Verbesserung, da der Monopolist durch diese Maßnahme verliert und nur die Konsumenten profitieren. Es ist vielmehr eine *Kaldor-Hicks*-Verbesserung, da die Konsumenten so viel gewinnen, dass sie den Monopolisten (hypothetisch) entschädigen könnten.

- 157 Monopole sind normalerweise instabil. Wenn Sie am heißen Strand einen einzigen Eisverkäufer sähen, der erfolgreich sein Eis zu hoffnungslos überkauerten Preisen verkauft, käme Ihnen dann nicht auch der Gedanke, ebenfalls am Strand Eis zu verkaufen? Die suprakompetitiven Profite der Monopole setzen einen Anreiz für Außenstehende, ebenfalls an diesem Markt tätig zu werden. Doch es gibt stabile Monopole, und viele von ihnen sind aus einem von zwei möglichen Gründen stabil. Entweder verhindert das Gesetz den Markteintritt, wie beispielsweise den Zugang zum Notarberuf in Deutschland. In einem Amtsbezirk kann nur der dort ernannte Notar tätig sein. Oder aber die Kostenstruktur einer bestimmten unternehmerischen Tätigkeit erlaubt es nur einer Firma, ihr Geschäft gewinnbringend auszuüben. Diese Kostenstruktur, die man **natürliches Monopol** nennt, wird durch große **Skaleneffekte** gekennzeichnet. Bei Skaleneffekten fallen die durchschnittlichen Kosten, während sich der Output vergrößert. Dies ist normalerweise der Fall, wenn es so hohe Fixkosten gibt, dass der Durchschnittskosten senkende Effekt durch die Verteilung der Fixkosten auf viele Output-Einheiten den Effekt zusätzlicher variabler Kosten durch weitere Output-Einheiten dominiert. Nehmen wir die Bahn als Beispiel. Es ist sehr teuer, ein Schienennetz zu bauen. Diese Kosten ändern sich mit Veränderung der beförderten Passagiere und Güter kaum. Ein weiterer Passagier oder ein weiterer transportierter Container produziert kaum zusätzliche Kosten. Hier fallen die variablen Kosten gegenüber dem Interesse, die hohen Fixkosten auf viele Kunden zu verteilen, praktisch nicht ins Gewicht. Im Extremfall des natürlichen Monopols kann ein Unternehmen seine durchschnittlichen Kosten nur dann unter den Maximalpreis drücken, den Kunden zu zahlen bereit wären, wenn es den gesamten Markt allein bedient.
- 158 Zwischen perfektem Wettbewerb und Monopol gibt es viele Marktformen mit einer begrenzten Wettbewerberzahl. Dabei sind die Wettbewerber weder Preisnehmer, noch können sie den Preis allein bestimmen. Der Preis entsteht aus dem Zusammenwirken aller Wettbewerber, bei dem alle Wettbewerber wissen, dass ihre jeweiligen Handlungen die der anderen beeinflussen. Diese wechselseitige Beeinflussung berücksichtigen die Wettbewerber und interagieren daher strategisch zur Beeinflussung der Preise. Diese strategische Interaktion wird mit Hilfe der Spieltheorie<sup>11</sup> analysiert.
- 159 Das Rechtsgebiet, das sich mit dem Erhalt des Wettbewerbs befasst, ist das Wettbewerbs- und Kartellrecht. In Europa sorgt Artikel 101 des Ver-

---

<sup>11</sup> Dazu § 4.

trags über die Arbeitsweise der Europäischen Union (AEUV) dafür, dass Unternehmen keine Vereinbarungen zu gemeinschaftlichem monopolistischen Handeln treffen.<sup>12</sup> In den Vereinigten Staaten erfüllt diese Aufgabe § 1 des Sherman Acts (ShA). Das Fusionskontrollregime verfolgt dasselbe Ziel mit anderen Mitteln. Artikel 102 AEUV und § 2 ShA verwehren es Unternehmen, die Monopolisten oder annähernd Monopolisten (marktmächtige Unternehmen) sind, ihre Marktmacht zu erhalten, auszuweiten oder auszunutzen.

## 2. Asymmetrische Information und verborgene Handlungen

Welche Informationen den Beteiligten einer Markttransaktion zur Verfügung stehen, kann wesentlichen Einfluss auf das Marktergebnis haben. Wenn man sich nicht sicher sein kann, welche Qualität ein Gebrauchtwagen hat, möchten man vielleicht keinen hohen Preis für den Wagen zahlen. Verkäufer guter Autos könnten auf diesen Umstand reagieren, indem sie ihre Wagen nicht verkaufen, so dass lediglich Wagen minderer Qualität zum Verkauf stehen.<sup>13</sup> Ein Arbeitgeber wird seinem Arbeitnehmer grundsätzlich gern ein hohes Gehalt zahlen, wenn der Arbeitnehmer hart arbeitet. Doch der Arbeitgeber wird oft nicht genau wissen, wie sehr sich der Arbeitnehmer ins Zeug legt. Er kann die Leistung des Arbeitnehmers nicht perfekt beobachten. Weil der Arbeitgeber das vorhersieht, wird er also keinen hohen Lohn zahlen und der Arbeitnehmer wird nicht hart arbeiten.<sup>14</sup> In beiden Beispielen können die Parteien nicht alle wertsteigernden Transaktionen verwirklichen (Verkauf guter Autos, hoher Lohn für gute Arbeit), weil einer Seite keine bzw. nur unzureichende Informationen zur Verfügung stehen. Um zu sehen, dass Informationen das Marktergebnis beeinflussen können, blicken wir auf das eben dargestellte Monopol zurück. Der Monopolist bestimmt einen Preis, der höher liegt als der Wettbewerbspreis. Dies generiert einen Wohlfahrtsverlust, indem manche Käufer am Kaufen gehindert werden, obwohl ihre Zahlungsbereitschaft die Herstellungskosten übersteigt. Information spielt eine Rolle bei der Entstehung dieses Wohlfahrtsverlusts. Wenn der Monopolist die Zahlungsbereitschaft jedes einzelnen Käufers kennen würde, könnte er viele verschiedene Preise setzen: einen für jeden Käufer. Er würde jedem Käufer genau das berechnen, was dieser höchstens zu zahlen bereit ist. So würde jeder

<sup>12</sup> Vgl. zu der geregelten Konstellation Rz. 177.

<sup>13</sup> Zu der Problematik *adverser Selektion* vgl. Rz. 274 ff.

<sup>14</sup> Zum Begriff *moral hazard* vgl. Rz. 302 f.

Käufer versorgt, dessen Zahlungsbereitschaft die Herstellungskosten übersteigt und der Monopolist verdiente gar noch mehr als unter einem einfachen Monopolpreis. Hierdurch würde Effizienz erreicht, allerdings bei einer sehr ungleichen Verteilung: Der Monopolist würde sich die gesamte entstehende Rente aneignen. Ein Grund, warum diese **perfekte Preisdifferenzierung** (auch Preisdiskriminierung ersten Grades) so schwer umzusetzen ist, liegt in der Tatsache, dass der Monopolist nicht die jeweilige Zahlungsbereitschaft jedes Konsumentens kennt. Der andere Grund ist, dass der Monopolist oftmals einen Wiederverkauf nicht verhindern kann. Er kann ja meist nicht beobachten und deshalb nicht kontrollieren, was der Käufer nach dem Kauf mit der Ware macht. Das verdirbt ihm die Preisdiskriminierung: Bildet sich unter seinen Käufern ein Weiterverkaufsmarkt mit perfektem Wettbewerb, würde sich auf diesem ein einziger Preis bilden. Jeder Käufer hätte dann Zahlungsbereitschaft in Höhe des auf dem Weiterverkaufsmarkt geltenden Preises – denn statt beim Monopolisten zu kaufen, könnte jeder Käufer ja auch am Weiterverkaufsmarkt kaufen. Kein Käufer wird daher dem Monopolisten mehr bezahlen, als er am Weiterverkaufsmarkt zahlen müsste.

### 3. Externe Effekte, Transaktionskosten und das Coase-Theorem

- 161 Die Theorie von Märkten geht davon aus, dass alle mit Herstellung und Konsum von Gütern verbundene Kosten und ebenso aller Nutzen des Güterkonsums bei den Parteien anfallen, die die Markttransaktion getätigt haben. Das ist oftmals eine vernünftige Annäherung an die Realität. Wenn Sie ein Glas Wein trinken, das Sie bei Ihrem Biowinzer gekauft haben, wird dieser hoffentlich die meisten Herstellungskosten getragen haben. Sie allein bezahlen mit dem Preis die Kosten und haben zugleich den Nutzen des Weintrinkens. Aber schon wenn wir das „Bio“ weglassen und an Probleme von Überdüngung und Pestizideinsatz in der Landwirtschaft denken, wird zweifelhaft, ob Sie und Ihr Winzer wirklich alle Kosten des Weinkonsums tragen. Klarer wird das noch im Flughafenbeispiel. Ein Flughafen verkauft Start- und Landemöglichkeiten an Fluggesellschaften. Doch der Lärm stellt Kosten für die Anwohner dar. Die Anwohner haben normalerweise nichts mit der Transaktion des Verkaufs von Start- und Landegenehmigungen zu tun. Wenn die Herstellungskosten zum Teil von einer Partei getragen werden, die abseits der Markttransaktion steht, dann werden diese Kosten Dritter nicht in der Angebotskurve dargestellt. Dies führt dazu, dass die Angebotskurve zu flach ist. Dadurch wird zum Marktpreis zu viel verkauft (s. Abb. 3.20). Kosten und Nutzen, die Dritten aufer-

legt werden, die nicht an der Markttransaktion beteiligt sind, nennen wir **externe Kosten** oder negative Externalitäten. Externe Kosten sind omnipräsent. Raucher stören vielleicht Menschen, die am Nachbartisch essen. Autofahrer nutzen einen Parkplatz, ohne ihn den Kindern abzukufen, die sonst dort spielen könnten. Und in den meisten Ländern tragen Kraftwerke zur Erderwärmung bei, indem sie CO<sub>2</sub> emittieren, ohne dass dies Bestandteil einer Transaktion wäre – etwas, das Märkte für Emissionszertifikate ändern sollen.

In einem berühmt gewordenen Artikel wies *Ronald Coase* darauf hin, 162 dass das Vorhandensein externer Effekte nicht bedeutet, dass Markttransaktionen ineffizient sind. Wäre das Recht, Externalitäten zu erzeugen oder zu unterbinden, handelbar, und entstünden beim Handel dieser Rechte keine Transaktionskosten, dann würden die Externalitäten vollständig internalisiert – mit dem Ergebnis, dass die Markttransaktionen effizient wären. Wenn das Gesetz den Flughafenanwohnern ein handelbares Recht auf Erlass einer Unterlassungsverfügung zugewiesen hätte, und wenn die nächtlich verdienten Profite des Flughafens ausgereicht hätten, um sowohl die Anwohner als auch den Flughafen reicher zu machen, hätte der Flughafen das Recht kaufen können. Davon hätten sowohl die Anwohner als auch der Flughafen profitiert und beide Seiten wären reicher gewesen.<sup>15</sup>

Zwei Vorbehalte gibt es jedoch, wobei der erste im *Coase*-Theorem explizit genannt wird: Das Theorem gilt nur „soweit es keine Transaktionskosten gibt“. Im vorgestellten Flughafenbeispiel haben wir die Transaktionskosten ignoriert. Ein Konflikt zwischen einem Flughafen und vielen Anwohnern ist aber eine typische Situation, in der die Transaktionskosten hoch sind. Transaktionskosten sind all jene Kosten, die Agenten aufwenden müssen, um eine Transaktion durchzuführen. Typische Beispiele sind die Kosten für die Suche nach Informationen oder nach dem richtigen Handelspartner, die Kosten für rechtliche Beratung oder die für das Aufsetzen eines Vertrags. Emotionale Hemmungen, einen Handel einzugehen, können auch als Transaktionskosten gelten, ebenso strategische Erwägungen, die Agenten daran hindern, zu einer Einigung zu kommen. Solche strategischen Hindernisse sind im Flughafenbeispiel sehr wahrscheinlich. Wenn wir viele Anwohner haben, von denen jeder ein Recht auf Erlass einer Unterlassungsverfügung hat, muss der Flughafen all diese Rechte der Anwohner erwerben, um nachts öffnen zu können. Die Rechte der Anwohner sind für den Flughafen wertlos, wenn auch nur eines der Rechte fehlt und dessen Inhaber eine Unterlassungsverfügung erwirkt. Daher 163

<sup>15</sup> Für Details zum *Coase*-Theorem vgl. Rz. 261 ff.

können wir uns die Situation so vorstellen, dass alle Rechte auf Erlass einer Unterlassungsverfügung für den Flughafen wertlos sind – bis auf das letzte, das der Flughafen erwirbt. Deshalb ist der Anwohner, der als Letztes verkauft, derjenige, der den höchsten Preis verlangen kann. Also will kein Anwohner früh verkaufen, sondern jeder Anwohner möchte der letzte sein, der verkauft. So beginnt niemand zu verkaufen. Dieses Problem bezeichnet man als **Holdout-Problem**. Im Ergebnis heißt das, es ist überhaupt nicht gewährleistet, dass in dem Flughafenbeispiel freie Transaktionen die Rechte auf Unterlassungsverfügungen dorthin lenken werden, wo sie am meisten geschätzt werden.

164 Der zweite Vorbehalt ist Verteilung. *Coase* spricht lediglich von Effizienz. Während unter den Annahmen von *Coase* die Effizienz nicht von der Zuteilung von Rechten betroffen ist, hat diese Zuteilung von Rechten durchaus einen Einfluss auf die Verteilung von Wohlstand. Stellen wir uns noch einmal vor, im Flughafenbeispiel gäbe es keine Transaktionskosten und diesmal läge das Recht zu entscheiden, ob Flugzeuge nachts starten dürfen, im Ausgangspunkt nicht bei den Anwohnern, sondern beim Flughafen. Wären die Profite hoch, würden die Anwohner ruhige Nächte nicht genügend wertschätzen, um den Flughafen für den Verzicht auf nächtliche Einnahmen zu entschädigen. Wären die Einnahmen niedrig, so wäre die Zahlungsbereitschaft der Anwohner ausreichend, um den Flughafen zu entschädigen, und das Recht ginge an die Anwohner. Aber wir sehen, dass die Anwohner ihre Ausstattung diesmal nicht erhöhen. Hier müssen die Anwohner entweder den Lärm ertragen – oder aber zahlen. Im vorherigen Beispiel hatten sie entweder ruhige Nächte oder eine Entschädigung. Also sind die Anwohner im zweiten Beispiel schlechter dran, obwohl in beiden Beispielen Effizienz erreicht wird. Einem an Effizienz interessierten Entscheidungsträger wird also eher egal sein, wem die Rechte im Ausgangspunkt zugeteilt werden, so dass er sich lieber auf die Minimierung von Transaktionskosten konzentrieren wird. Die betroffenen Agenten werden hingegen sehr daran interessiert sein, wo die Rechte zugeteilt werden, und – möglicherweise mit aller Kraft – versuchen, den politischen Entscheidungsträger zu beeinflussen. Dies eröffnet ein weites Feld an weiteren Fragen, mit denen wir uns im Kapitel über Public Choice<sup>16</sup> befassen.

165 Indem *Coase* erkannte, dass Rechte – soweit es keine Transaktionskosten gibt – effizient verteilt werden, schuf er das Fundament für die ökonomische Analyse des Rechts. Soweit das Recht nach Effizienz strebt, sollte es sich dort, wo Transaktionskosten gering sind, auf die Definition handel-

<sup>16</sup> Vgl. § 6, insb. Rz. 355 ff. zu *rent seeking*.

barer Eigentumsrechte beschränken. Dann, so die Hoffnung, können Agenten Ineffizienzen ganz einfach durch Transaktionen lösen. In der Realität ist die Welt jedoch voll von Transaktionskosten, wie wir am Flughafenbeispiel gesehen haben. Wo erhebliche Transaktionskosten bestehen und nicht beseitigt werden können, kann das Recht ökonomische Methoden nutzen, um effiziente Allokationen zu identifizieren und zu prüfen, welche Anreize sie herbeiführen können, um dann juristische Wege aufzufinden, um solche Allokationen zu realisieren.

#### 4. Nicht private Güter

Märkte funktionieren gut, wenn gehandelte Güter private Güter sind. Der Konsum eines privaten Guts ist ausschließbar, d.h. andere können daran gehindert werden, das Gut zu konsumieren. Zudem ist der Konsum von privaten Gütern rivalisierend, d.h. wenn Sie es konsumieren, kann kein anderer es konsumieren. Besitzen Sie ein Stück Käse, können Sie sicherstellen, dass Sie es für sich behalten (ausschließbar), und wenn Sie es essen, bleibt für andere nichts übrig (rivalisierend). Wenn den gehandelten Gütern diese Charakteristika fehlen, kann ein Markt sie nicht effizient zuteilen. Nehmen wir als Beispiel Allmendegüter. **Allmendegüter** sind im Konsum rivalisierend, aber nicht ausschließbar. Die Fischerei in den Meeren Europas ist ein Anwendungsgebiet. Was ich fische, können Sie nicht fischen (rivalisierend), aber ich kann Sie nicht realistisch am Fischen auf europäischen Meeren hindern (nicht ausschließbar). Wenn alle europäischen Fischer auf das Meer hinausfahren und fischen würden, wie sie es gerade wollen, gäbe es bald keine Fische mehr. Denn alle Fischer haben zwar das gemeinsame Interesse, eine gesunde Fischpopulation im Meer zu erhalten, damit die Menschen auf Jahre hinaus davon zehren können. Doch ein einzelner Fischer hat einen Vorteil, wenn er mehr fischt – egal, ob seine Kollegen sich an den ursprünglichen Plan halten oder nicht. Da dies für alle Fischer gilt, fischen sie zu viel. Der freie Markt für Fisch löst das Problem nicht, er verschärft es. Das europäische Recht versucht, dieses Problem durch das rechtliche Instrument von Fangquoten abzumildern. Wir werden dies eingehender beleuchten, wenn wir im nächsten Kapitel über das Recht und die Theorie strategischer Interaktion sprechen.<sup>17</sup>

<sup>17</sup> Hierzu Rz. 223, 226, 252.



## 5. Beispiel Flughafen (2)

- 167 Halten wir erneut kurz inne, um das zu Angebot und Märkten Gelernte auf das Flughafenbeispiel zu übertragen. Vergessen wir für einen Moment das Recht der Anwohner auf Erlass einer Unterlassungsverfügung. Nehmen wir an, der Flughafen operiere an einem Markt, wo die Flughäfen sogenannte Slots verkaufen, d.h. Zeiten, während derer Flugzeuge landen, beladen werden und wieder starten können. Fluglinien verwenden diese Slots als *Inputs* für die Transportdienstleistungen, die sie anbieten. Nehmen wir nun an, alle Flughäfen am Markt seien identisch und alle hätten Anwohner, die sich von Nachtflügen gestört fühlen. Und nehmen wir weiter an, unser Flughafen habe einen Marktanteil von zehn Prozent; auf dem Markt funktioniere der Wettbewerb passabel, so dass ein Modell perfekten Wettbewerbs eine hinreichende Näherung darstellt. Dann würde Abb. 3.20 ein System von Nachfrage für und Angebot von Slots darstellen. Nach diesem System würden 300 Slots an den Markt gebracht.
- 168 Würden die Kosten der Anwohner in der Angebotskurve abgebildet? Das wäre nur der Fall, wenn der Flughafen durch den Verkauf eines Slots die Möglichkeit verlöre, den Slot an die Anwohner zu verkaufen, die gern ihre Ruhe hätten (Opportunitätskosten, *Coase-Theorem*). Im Prinzip könnten die Anwohner Slots kaufen, wenn ihre Wertschätzung für Slots hoch genug wäre und es keine Transaktionskosten gäbe. Ein einzelner Anwohner hat wohl keine ausreichend hohe Zahlungsbereitschaft, um einen ganzen Flughafen nachts ruhigzustellen. Die Nachbarn müssten sich also notwendig zusammentun, um das Geld aufzubringen. Doch tragen sie genug bei? Ruhige Nächte sind im Konsum nicht ausschließbar. Wenn Frau Reich viel beiträgt, erhöht sie die Wahrscheinlichkeit, dass sie ruhig schläft. Aber sie kann Herrn Schmitt, der nicht beigetragen hat, nicht daran hindern, ebenfalls die Nachtruhe zu genießen. Zudem sind ruhige Nächte im Konsum nicht rivalisierend. Die Tatsache, dass Frau Reich ruhig schläft, macht die Nächte für Herrn Schmitt nicht lauter. Daher wird es zwischen den Flughafenanwohner zu einem sozialen Dilemma kommen. Jeder hat einen Anreiz, seine Wertschätzung ruhiger Nächte zunächst zu untertreiben, in der Hoffnung, dass andere zahlen und er am Ende doch ruhig schläft. So wird aufgrund des Dilemmas nicht genug Geld zusammenkommen und die Einwohner werden nicht in Slots investieren. Weil der Flughafenbetreiber an die Anwohner nichts verkaufen kann, verliert er nichts durch den Nachtflugbetrieb. Die Kosten der Anwohner werden daher nicht in der Angebotskurve abgebildet. Die Menschen, die nahe an unserem Flughafen wohnen, müssen also jede Nacht 30 Starts und Landungen er-

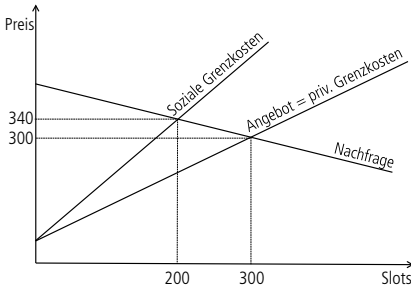


Abbildung 3.20: Angebot und soziale Grenzkosten

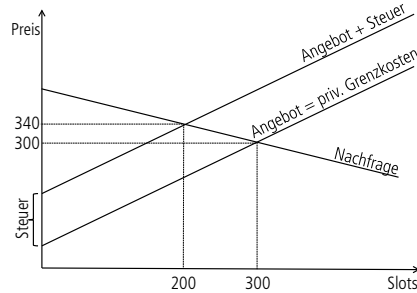


Abbildung 3.21: Pigou'sche Steuer

tragen (10 Prozent von 300). Könnten staatliche Eingriffe helfen, den Konflikt zwischen dem Flughafen und seinen Anwohnern beizulegen?

Das Problem ist, dass die Angebotskurve die privaten Grenzkosten des Flughafens darstellt. Die sozialen Grenzkosten des Nachtbetriebs sind aber höher: Allen Anwohnern entstehen Kosten durch jedes einzelne Flugzeug, das über ihre Dächer fliegt. Natürlich macht es einen Unterschied, ob drei Flugzeuge oder 30 die Nachtruhe stören. Daher stellt die private Grenzkostenkurve nicht die Gesamtgrenzkosten der Bereitstellung von Slots dar. So kreuzen sich Nachfrage und Angebot an einem Punkt, an dem die Menge gemessen an der effizienten Allokation – die das ganze soziale System betrachten sollte – zu hoch ist.

Die sozialen Grenzkosten würden schneller steigen als die privaten Grenzkosten des Flughafens.<sup>18</sup> Der Maßstab der Effizienz würde verlangen, dass das Angebot auf den sozialen Kosten einer Aktivität basiert. Nutzten wir die sozialen Grenzkosten als Angebotskurve, sähen wir, dass nachts weniger Flugzeuge fliegen würden (200 am Markt und 20 an unserem Flughafen). Das Gesetz kann den Flughafen zwingen, die soziale Grenzkostenkurve annähernd einzupreisen, indem jedes startende und landende Flugzeug besteuert wird (sogenannte *Pigou*-Steuer; s. Abb. 3.21). Doch es ist schwierig, die richtige Höhe dieser Steuer zu bestimmen. Ähnliches bewirkt die Schadenersatzhaftung. Stellen wir uns vor, der Flughafen müsste eine lebenslange Rente an jeden Anwohner zahlen, der durch den

<sup>18</sup> Je mehr Flugzeuge es gibt, unter deren Lärm Sie bereits leiden mussten, desto ärgerlicher ist jedes weitere Flugzeug. Das nehmen wir hier an. Vielleicht finden Sie es plausibler, dass alle Flugzeuge gleich ärgerlich sind; in dem Fall läge die Funktion der sozialen Grenzkosten zwar höher als die private Grenzkostenfunktion, verlief aber parallel.

Lärm krank geworden ist. Jedes Flugzeug würde die Wahrscheinlichkeit erhöhen, dass ein Anwohner krank wird. Mit dieser Wahrscheinlichkeit stiege dann auch die Wahrscheinlichkeit, dass der Flughafen haften müsste. Die Steigerung der erwarteten Kosten dieser Schadenersatzzahlungen würde in den Grenzkosten des Flughafens auftauchen. Auf diese Weise kann das Haftungsrecht Unternehmen motivieren, soziale Kosten zu berücksichtigen, die sonst nicht in der privaten Kostenfunktion erschienen.

## § 4 – Spieltheorie

**Literatur:** D. Baird/R. Gertner/R. Picker, *Game Theory and the Law*, 2000; S. Berninghaus/K.-M. Ehrhart/W. Güth, *Strategische Spiele*, 3. Aufl. 2010; K. Binmore, *Natural Justice*, 2005; A. Diekmann, *Spieltheorie*, 2009; E. Feess, *Mikroökonomie. Eine spieltheoretisch- und anwendungsorientierte Einführung*, 3. Aufl. 2004; D. Fudenberg/J. Tirole, *Game Theory*, 1996; J. Goldsmith/E. Posner, *The Limits of International Law*, 2005; A. Guzman, *How International Law Works*, 2008; M. Holler/G. Illing/S. Napel, *Einführung in die Spieltheorie*, 8. Aufl. 2016; R. McAdams/E. Rasmusen, *Norms and the Law*, in: A. Polinsky/S. Shavell (Hg.), *Handbook of Law and Economics II*, 2007, 1573–1618; R. Myerson, *Game Theory*, 1997; W. Poundstone, *Prisoner's Dilemma*, 1993; E. Rasmusen, *Games and Information*, 4. Aufl. 2006; F. Scharpf, *Interaktionsformen*, 2000; T. Schelling, *The Strategy of Conflict*, 1960; P. Straffin, *Game Theory and Strategy*, 5. Aufl. 2004.

### I. Spieltheorie und Recht

Die Spieltheorie ist eine allgemeine, mathematische Theorie des rationalen Entscheidens in strategischen Interaktionen. Sie hat sich zwar in ihren ersten Anfängen mit der Analyse von Gesellschaftsspielen befasst,<sup>1</sup> ist mittlerweile aber zu einem unentbehrlichen Analyseinstrument in der Ökonomie und in vielen anderen Sozialwissenschaften geworden. Diese Einführung möchte grundlegende Begriffe und Konzepte dieser komplexen und leistungsfähigen Theorie erläutern. Sie legt dabei die klassischen Annahmen rationaler und eigennütziger Individuen (*homo oeconomicus*) zugrunde.<sup>2</sup> Um die Zugangshürden für Juristen in Maßen zu halten, wird der

<sup>1</sup> von Neumann, Zur Theorie der Gesellschaftsspiele, *Mathematische Annalen* 100 (1928), 295 ff.; zur Entwicklung der Spieltheorie: Berninghaus et al., *Strategische Spiele*, 2010, S. 1 ff.

<sup>2</sup> Hierzu bereits Rz. 69 ff.; soziale Präferenzen und die Grenzen der Rationalität untersucht im Rahmen der Verhaltensökonomik die experimentelle Spieltheorie; dazu unten, § 8, sowie Englerth, *Behavioral Law and Economics – eine kritische Einführung*, in: Engel et al. (Hg.), *Recht und Verhalten*, 2007, 60 ff.; speziell zu Fairness-Präferenzen Magen, *Fairness, Eigennutz und die Rolle des Rechts*, ebd., 261 ff., sowie ausführlich ders., *Gerechtigkeit als Proprium des Rechts*, 2010, S. 93 ff.

für die Spieltheorie konstitutive mathematische Formalismus auf ein Minimum beschränkt und nicht formal eingeführt, sondern verbal erläutert.

### 1. Die Interdependenz von Interessen in juristischer und spieltheoretischer Perspektive

- 171 Rechtswissenschaften wie Spieltheorie haben mit dem gleichen Grundtatbestand menschlicher Sozialität zu tun: Regelmäßig hängt die Verwirklichung individueller Interessen vom Verhalten anderer Individuen ab und hat Auswirkungen auf deren Interessen. Beide Disziplinen betrachten die Probleme, die mit dieser wechselseitigen Abhängigkeit der Interessen verbunden sind, aber auf sehr unterschiedliche Weisen. Der Jurist befasst sich mit Interessen vor allem unter dem Aspekt des Konflikts und beurteilt diese vornehmlich im Lichte rechtlicher Normen, die zum Schutz und zum Ausgleich von Interessen in Geltung sind. Seine Aufgabe schließt dabei zwar auch eine Bewertung der betroffenen Interessen ein, aber nach den im Recht niedergelegten Gemeinwohlmaßstäben. Der Jurist sucht auch nur selten nach Wegen für eine optimale Interessenverwirklichung, sondern fragt eher nach dem angemessenen Interessenausgleich. Die Spieltheorie dagegen nimmt die individuellen Interessen als existent und gegeben hin und fragt, wie der Einzelne diese in Situationen wechselseitiger Interdependenz bestmöglich verfolgen kann, das heißt, welches Verhalten für den einzelnen Akteur **instrumentell rational** ist und wie ein **sozial effizienter** Zustand erreicht werden könnte.
- 172 Als Aussagen einer mathematischen Theorie sind diese Aussagen zunächst einmal nur analytischer Natur, sagen also weder etwas über die Wirklichkeit noch über normative Gebote. In den Wirtschafts- und Sozialwissenschaften werden die Spiele der Spieltheorie aber auch als (notwendig vereinfachende) **Modelle des tatsächlichen Verhaltens** von Menschen interpretiert und zur Grundlage von deskriptiven Verhaltenstheorien gemacht. Solche Theorien der **positiven Spieltheorie** können für Juristen in mehrfacher Hinsicht von Nutzen sein. Sie helfen nicht nur, genauer zu verstehen, wieso aus der Interdependenz von Interessen „Konflikte“ und andere soziale Probleme entstehen, sondern ermöglichen auch ein sehr viel differenzierteres Verständnis verschiedener Ursachen sozialer Probleme, die mit dem Begriff des „Konflikts“ alles andere als zureichend beschrieben sind. Die Spieltheorie wirft auch Licht auf Fragen wie die, wann und wieso Rechtsnormen notwendig sind, was die Aufgaben des Rechts sind, wie die Adressaten auf Normen reagieren werden, welche Probleme mit der Erreichung vorgegebener Regelungsziele verbunden sind und weshalb

der Appell an die Rechtstreue oder das Verantwortungsbewusstsein der Normadressaten mitunter nicht weiterhilft.

Die Spieltheorie findet aber auch im Bereich der **normativen Ökonomie** 173 Verwendung, wenn untersucht wird, ob das individuell rationale Ergebnis auch kollektiv rational oder effizient ist. Das sogenannte Mechanismus-Design als Zweig der Spieltheorie entwirft Regeln (etwa Vertragsklauseln oder das Design von Auktionen), um effizientere Ergebnisse zu erzielen. Natürlich sind Aussagen der normativen Ökonomie für das geltende Recht unverbindlich. Sie sind aber rechtspolitisch von Interesse und können auch das Verständnis des geltenden Rechts verbessern, wenn die gesetzgeberischen Ziele ökonomische Aspekte beinhalten (etwa beim Umweltschutz, der ökonomisch oft ein **Internalisierungsproblem** darstellt). Auch wenn dies nicht im Zentrum der ökonomischen Forschungen steht, kann die Spieltheorie schließlich ein wertvolles Instrument sein, um die Verteilungsfolgen strategischer Interaktionen sichtbar zu machen.

## 2. Individuelles Entscheiden und strategische Interdependenz

Die Spieltheorie befasst sich also mit Entscheidungssituationen, die durch **strategische Interdependenz** gekennzeichnet sind. Damit meint man Situationen, in denen das Resultat der Entscheidung eines Akteurs auch davon abhängt, wie sich der oder die anderen Akteure entscheiden. Strategische Interdependenz im eigentlichen Sinn liegt aber nur vor, wenn der Akteur diese Interdependenz berücksichtigen muss, um rational zu entscheiden. Letzteres ist aber nicht bei jeder Interesseninterdependenz der Fall. Zum Beispiel kann der einzelne Marktteilnehmer auf Märkten, auf denen vollständige Konkurrenz herrscht, das Verhalten der anderen Marktteilnehmer nicht strategisch beeinflussen, so dass er gleichsam gegen die Natur spielt. Aber auch ein Monopolist steht nicht vor einem strategischen Entscheidungsproblem, insoweit er keinen (potentiellen) Konkurrenten hat, dessen Verhalten er für seine Entscheidungen berücksichtigen müsste. Gegenstand der Spieltheorie in den Wirtschaftswissenschaften sind damit insbesondere Märkte zwischen den Extremen von Wettbewerbsmärkten und Monopolmärkten. Wir werden das sogleich an einem Beispiel erläutern. 174

## 3. Spieldefinition, Normalform und Extensivform

Spiele im Sinne der Spieltheorie sind formale Modelle, die den Reichtum 175 lebensweltlicher Situationen bewusst ausblenden, um bestimmte, für relevant erachtete Aspekte der jeweiligen Situation isolieren und mathema-

tisch analysieren zu können. Damit ein strategisches Problem mit den Mitteln der Spieltheorie analysiert werden kann, muss es aber formal als ein Spiel definiert werden. Dafür müssen bestimmte Elemente festgelegt werden: die **Spieler**, die **Regeln** (Aktionen, Zugfolge, Informationsstand), die möglichen Spielergebnisse und die **Auszahlungen**, das heißt die Bewertung der Ergebnisse nach den Präferenzen der Spieler (angegeben durch die Auszahlungsfunktion). Für die Definition eines Spiels gibt es zwei Darstellungsformen, die weitgehend, aber nicht vollständig äquivalent sind, nämlich die sogenannte **Normalform**<sup>3</sup> (oder strategische Form) und die **Extensivform**<sup>4</sup>. Beide unterscheiden sich vor allem darin, dass in der Extensivform die Reihenfolge der Züge und der Informationsstand der Spieler explizit modelliert werden, während man bei Normalformspielen, soweit nichts Anderes bestimmt ist, annimmt, die Spieler zögen simultan und ohne die Entscheidung der Mitspieler zu kennen.

- 176 An einem Beispiel werden im Folgenden zunächst das Normalformspiel eingeführt, grundlegende Lösungskonzepte erläutert und Möglichkeiten zur normativen Bewertung der Spielergebnisse dargestellt (II.). Sodann wird ein Überblick über verschiedene Konstellationen strategischer Interdependenz gegeben (III.). Im Anschluss werden Spiele in Extensivform behandelt (IV.) und die Modellierung von Recht und informalen Institutionen skizziert (V.).

## II. Spiele in Normalform

### 1. Das Kartelldilemma

- 177 Als Beispiel für ein strategisches Entscheidungsproblem wählen wir zwei Unternehmen, die auf einem von ihnen beherrschten Markt einen Monopolgewinn erzielen möchten, also den Gewinn, den man im Vergleich zum vollkommenen Wettbewerb zusätzlich erzielen kann, wenn man die Preise unabhängig vom Wettbewerb festlegen kann, etwa wenn man das Gesamtangebot verknappen und dadurch höhere Preise fordern kann.<sup>5</sup> Im Unterschied zu einem Monopolisten kann ein einzelner Duopolist einen solchen (anteiligen) Monopolgewinn allerdings nicht schon dadurch erzielen, dass er einseitig seine eigene Angebotsmenge reduziert, weil der Mo-

<sup>3</sup> Dazu unten in §4.II und 4.III.

<sup>4</sup> Siehe dazu § 4.IV.

<sup>5</sup> Beim sog. Cournot-Wettbewerb; vgl. Feess, Mikroökonomie, 2004, S. 369 ff.; Schwalbe/Zimmer, Kartellrecht und Ökonomie, 2. Aufl. 2011, S. 39 ff.

nopolgewinn vom Gesamtangebot abhängt, auf das auch der andere Duopolist Einfluss hat. Für diesen ist es aber möglicherweise vorteilhaft, auf eine Produktionsreduktion seines Konkurrenten mit einer Erhöhung der eigenen Produktion zu reagieren. Ein Duopolist, der seinen eigenen Gewinn maximieren möchte, muss deshalb das Verhalten des anderen antizipieren und berücksichtigen. Man kann dieses Entscheidungsproblem (sehr vereinfacht) wie folgt als ein Gefangenendilemma<sup>6</sup> modellieren:

### a. Spieler, Regeln und Ergebnisse

Zunächst müssen die Spieler, die Regeln (Aktionen, Zugfolge, Informationsstand) und die Spielausgänge festgelegt werden. Spieler sind die beiden Duopolisten  $D_1$  und  $D_2$ . Diese können verschiedene Produktionsmengen wählen. Für die Modellierung reduzieren wir sie auf zwei **Aktionen**, nämlich die Produktion einer kleinen Angebotsmenge (K) und die Produktion einer großen Angebotsmenge (G). Bezüglich der **Zugfolge** nehmen wir weiter an,  $D_1$  und  $D_2$  ziehen gleichzeitig und nur ein einziges Mal (simultanes Einmal-Spiel). Was den **Informationsstand** anbetrifft, unterstellen wir im Folgenden, soweit nichts anderes gesagt wird, dass alle Spieler die möglichen Aktionen aller Spieler, alle vorangegangenen Züge und die Auszahlungen aller Spieler kennen, also perfekte und vollständige Informationen haben. Ziehen die Spieler gleichzeitig, nimmt man an, dass der jeweilige Spielzug des anderen Spielers nicht bekannt ist. Hier kennen  $D_1$  und  $D_2$  die Aktionen, aus denen sie selbst und der Mitspieler auswählen können, und die jeweils resultierenden Auszahlungen, sie wissen aber bei ihrer Entscheidung nicht, wie der andere entscheidet. Aus den möglichen Kombinationen der  $D_1$  und  $D_2$  jeweils zur Verfügung stehenden Aktionen resultieren vier mögliche **Spielergebnisse**, nämlich a) beide wählen die kleine Angebotsmenge K, also (K, K); b) beide wählen die große Angebotsmenge G, also (G, G); sowie c) und d) jeweils einer wählt die kleine, und der andere die große Angebotsmenge, also (K, G) und (G, K).

### b. Präferenzen und Auszahlungen

Für jedes dieser Spielergebnisse müssen die Auszahlungen  $H_i$  der Spieler bestimmt werden. Sie geben für jedes mögliche Spielergebnis an, wie jeder Spieler dieses nach seinen Präferenzen bewertet. Diese Auszahlungen können, je nach Fragestellung, verschieden interpretiert werden, etwa als ordinale Rangordnung, als kardinaler (metrisch mess- und vergleichbarer) *von*

<sup>6</sup> Vgl. Rz. 219 f.



Neumann/Morgenstern-Nutzen, als Zahlungsbereitschaft oder als Gewinne in Geld.<sup>7</sup> Für die Zwecke dieser Einführung nehmen wir entsprechend einer verbreiteten Konvention im Folgenden an, die Spieler hätten **eigennützige Präferenzen**, bewerteten also nur die Folgen einer Handlung für sich selbst. Das lässt soziale Präferenzen wie Altruismus oder Präferenzen für Fairness unberücksichtigt, wie sie die Verhaltensökonomik untersucht.<sup>8</sup> Der Einfachheit halber nehmen wir weiter an, die Auszahlungen gäben bei Personen die individuelle Zahlungsbereitschaft wieder und bei Firmen den Gewinn.<sup>9</sup> Unsere Duopolisten achten also nur auf ihren eigenen Gewinn, während ihnen die Auswirkungen ihrer Entscheidungen auf den Gewinn des jeweils anderen gleichgültig sind.

180 Als Modell für reale Situationen sind Spiele immer starke Vereinfachung, die auf theoretischen Annahmen beruhen. Diese Annahmen sollten aber theoretisch begründet und empirisch angemessen sein.<sup>10</sup> Das gilt insbesondere für die Auszahlungen, die für die Natur des Spiels wesentlich sind. Unser folgendes Beispiel stützt sich auf ein vereinfachtes Modell duopolistischen Wettbewerbs, dem wir die **Marktergebnisse** entnehmen, die sich bei verschiedenen Angebotsmengen ergeben. Insoweit nehmen wir zunächst an, die kleinen Angebotsmengen  $K$  entsprechen zusammen der Menge, die ein gewinnmaximierender Monopolist anbieten würde. Wollen beide Duopolisten zu Lasten der Konsumenten kolludieren, dann wählen sie diese Menge – spielen also  $(K, K)$  – und realisieren dadurch anteilig den Monopolgewinn. Wir nehmen weiterhin an, dass die großen Angebotsmengen zusammen der Menge entsprechen, die sich als Marktgleichgewicht ergäbe, wenn  $D_1$  und  $D_2$  nicht kolludierten, sondern konkurrierten. Spielen die Duopolisten also  $(G, G)$ , erzielen sie einen geringeren Gewinn (der je nach Marktform – Mengen- oder Preiswettbewerb, homogene oder differenzierte Güter, etc. – auf das Niveau bei perfektem Wettbewerb reduziert sein wird oder zwischen diesem und dem Monopolgewinn liegen kann).<sup>11</sup> Insoweit  $D_1$  und  $D_2$  Kollusion gegenüber Konkurrenz vorziehen, sind ihre Interessen noch gleichläufig. Für die Auszahlung  $H$  der beiden Spieler  $i$  gilt insoweit  $H_i(K, K) > H_i(G, G)$ . Wählt allerdings nur ein Duopolist die ge-

<sup>7</sup> Hierzu Rz. 65 ff.

<sup>8</sup> Hierzu Rz. 502.

<sup>9</sup> Solche Auszahlungen beinhalten auch kardinale Präferenzordnungen und lassen Aussagen über den Erwartungsnutzen, Risikopräferenzen, gemischte Strategien und Wohlfahrtsgewinne zu.

<sup>10</sup> Hellwig, Neoliberales Sektierertum oder Wissenschaft?, Preprints des MPI für Gemeinschaftsgüter 2015/17, 9 ff.

<sup>11</sup> Vgl. Schwalbe/Zimmer, Kartellrecht und Ökonomie, 2. Aufl. 2011, S. 39 ff.

ringere, kollusive Angebotsmenge, während der andere die höhere, wettbewerbliche Menge anbietet, dann kann, so nehmen wir an, der konkurrierende Duopolist seinen Gewinn noch über die gemeinsame Kollusion steigern, während der Gewinn des sich einseitig kollusiv verhaltenden Duopolisten noch unter seinen Gewinn bei beiderseitigem Wettbewerb fällt. Was für  $D_1$  das beste Ergebnis ist, nämlich (K, G), ist für  $D_2$  also das schlechteste, und umgekehrt. Insoweit sind die Interessen von  $D_1$  und  $D_2$  gegenläufig.

### c. Spielmatrix

Normalformspiele mit zwei Personen und einer endlichen Strategiemenge kann man in Gestalt einer Matrix aufschreiben, die bei zwei Aktionsmöglichkeiten die simple Form einer 2x2-Matrix annimmt (sogenannte **2x2-Bimatrix-Spiele**). Die möglichen Aktionen von  $D_1$  sind dabei in den Reihen wiedergegeben,  $D_1$  kann also zwischen der oberen und der unteren Reihe wählen. Die möglichen Aktionen von  $D_2$  sind dagegen in den Spalten dargestellt,  $D_2$  kann also zwischen der rechten und der linken Spalte wählen. Durch die Entscheidungen von  $D_1$  und  $D_2$  wird eines der vier Felder ausgewählt, die mithin die vier möglichen Spielausgänge repräsentieren.<sup>12</sup> Die Auszahlungen für  $D_1$  werden jeweils links, die für  $D_2$  jeweils rechts angeführt:

|                      |              | D <sub>2</sub> wählt |             |
|----------------------|--------------|----------------------|-------------|
|                      |              | kleine Menge         | große Menge |
| D <sub>1</sub> wählt | kleine Menge | 3, 3                 | 1, 4        |
|                      | große Menge  | 4, 1                 | 2, 2        |

Tabelle 4.1: Spielmatrix des Kartell-Dilemmas

## 2. Lösungskonzepte für individuell rationales Verhalten

Spiele wie das genannte beschreiben eine strategische Situation, aber noch nicht, wie sich ein instrumentell rationaler Spieler verhalten sollte, was also die beste Strategiewahl ist. Mit dieser Frage befassen sich die sogenannten **Lösungs- oder Gleichgewichtskonzepte**. Ausgangspunkt der Lösungskonzepte sind die möglichen Strategien der Spieler. Unter Strategien versteht die Spieltheorie einen *ex ante* festgelegten Handlungsplan, der für jede Zugmöglichkeit festlegt, welche Aktion gewählt werden soll. Bei Spie-

<sup>12</sup> (K,K) links oben, (G,G) rechts unten, etc.

len wie dem vorstehenden, in denen nur einmalig und simultan gezogen wird, fallen die möglichen (reinen) Strategien allerdings mit den möglichen Aktionen zusammen. Würde  $D_2$  dagegen erst nach  $D_1$  und in Kenntnis von dessen Entscheidung ziehen, ergäben sich für ihn doppelt so viele Strategien, weil er auf zwei mögliche Aktionen des  $D_1$  mit zwei möglichen eigenen Aktionen reagieren kann.<sup>13</sup> Darauf wird zurückzukommen sein. Man sollte aber im Hinterkopf behalten, dass Strategien einen vollständigen *ex ante* Handlungsplan für jede Eventualität vorsehen müssen und deshalb komplexer sein können als Aktionen.

- 183 Welche Ergebnisse die von den Spielern gewählten Strategien produzieren, hängt natürlich von den Entscheidungen aller Spieler ab. Lösungskonzepte befassen sich deshalb nicht mit den Entscheidungen einzelner Spieler, sondern damit, wie die Entscheidungen der beteiligten Spieler zusammenwirken, nämlich mit **Strategieprofilen**. Unter einem Strategieprofil versteht man eine Kombination von Strategien der verschiedenen Spieler, die für jeden Spieler eine und nur eine mögliche Strategie enthält (im Beispiel etwa:  $D_1$  kolludiert,  $D_2$  kolludiert). Ein Strategieprofil induziert damit ein bestimmtes Spielergebnis ( $G, G$ ) und damit bestimmte Auszahlungen ( $2/2$ ). Als Lösungen oder Gleichgewichte eines Spiels bezeichnet man solche Strategieprofile, die für jeden Spieler eine beste Strategie enthalten. Die Lösungskonzepte versuchen nun eine Antwort auf die Frage zu geben, was eine beste Strategie ist. Grundlegend sind das Konzept der Dominanz und das sog. **Nash-Gleichgewicht**, die sogleich erläutert werden. Außerdem gibt es verschiedene Ansätze zur Verfeinerung des Konzepts eines Nash-Gleichgewichts, von denen die Teilspielperfektion an späterer Stelle bei den sequentiellen Spielen in Extensivform eingeführt wird.

#### a. Dominanz

- 184 Welche Strategie als beste gelten kann, muss bei strategischer Interdependenz im Verhältnis zu den möglichen Entscheidungen der anderen Spieler bestimmt werden. Es leuchtet aber ein, dass eine Strategie für einen Spieler jedenfalls dann die beste Wahl darstellt, wenn sie unabhängig davon, wie sich die anderen Spieler verhalten, immer zu den jeweils höchsten Auszahlungen führt. Eine solche Strategie dominiert die anderen Strategien dieses Spielers. Gibt es für jeden Spieler eine dominante Strategie, hat das Spiel ein Gleichgewicht in dominanten Strategien. Ob ein Spieler eine dominante Strategie hat, kann man an der Matrix-Darstellung eines Spiels ablesen.

<sup>13</sup> Eine der Strategien von  $D_2$  lautet z. B.: „Wenn  $D_1$  K spielt, spiele ich auch K; wenn  $D_1$  G spielt, spiele ich auch G.“

Man geht dafür alle möglichen Entscheidungen des Spielers durch und fragt jeweils, welche Entscheidung für den Spieler vorteilhafter ist, gegeben dass sich sein Mitspieler so oder so verhält. Im Beispiel wären also für  $D_1$  zwei Fragen zu stellen. Die erste ist: Welche Entscheidung ist für mich vorteilhaft, wenn  $D_2$  die kleine Menge wählt, also in der Matrix die linke (grau unterlegte) Spalte spielt?

|                         |              | D <sub>2</sub> wählt |             |
|-------------------------|--------------|----------------------|-------------|
|                         |              | kleine Menge         | große Menge |
| D <sub>1</sub><br>wählt | kleine Menge | ↓ 3, 3               | 1, 4        |
|                         | große Menge  | 4, 1                 | 2, 2        |

Tabelle 4.2: Vorzugsrelation für D1

In diesem Fall müsste sich  $D_1$  zwischen der oberen und der unteren Zelle in dieser Spalte entscheiden.  $D_1$  würde also die erste Spalte betrachten und die linken Auszahlungen in der oberen und in der unteren Zelle vergleichen, also 3 und 4. Weil diese seine Präferenzordnung wiedergeben, würde er von beiden Auszahlungen die höhere wählen. Welcher Zelle ein Spieler danach den Vorzug gibt, kann man zur Veranschaulichung mit sogenannten Abweichungspfeilen notieren, in diesem Fall mit vertikalen Pfeilen ( $\uparrow$  und  $\downarrow$ ) neben den Auszahlungen von  $D_1$ . Hier würde sich  $D_1$ , wenn  $D_2$  die linke Spalte wählt, für die untere Zeile entscheiden. Die Antwort wäre also: Wenn  $D_2$  kolludiert, würde  $D_1$  konkurrieren. Die andere Frage, die sich  $D_1$  stellt, ist: Welche Entscheidung ist für mich vorteilhaft, wenn  $D_2$  die große Menge wählt? Dafür kann man die zweite Spalte betrachten, und die linken Auszahlungen in der oberen und unteren Zelle vergleichen, also 1 und 2. Auch in diesem Fall ist es also für  $D_1$  rational zu konkurrieren. Das heißt, Wettbewerb ist für  $D_1$  in beiden Fällen vorzugswürdig, stellt für ihn also eine **dominante Strategie** dar.

185

|                         |              | D <sub>2</sub> wählt |             |
|-------------------------|--------------|----------------------|-------------|
|                         |              | kleine Menge         | große Menge |
| D <sub>1</sub><br>wählt | kleine Menge | 3, 3                 | ↓ 1, 4      |
|                         | große Menge  | 4, 1                 | 2, 2        |

Tabelle 4.3: Dominante Strategie für D1

$D_2$  stellt spiegelbildlich die gleichen Überlegungen an. Allerdings wählt  $D_2$  nicht zwischen den Zeilen (oben/unten), sondern zwischen den Spalten (links/rechts), und zwar jeweils unter der Annahme, dass  $D_1$  die eine oder die andere Zeile wählen würde. Welche Spalte  $D_2$  bevorzugt, gegeben

186

dass  $D_1$  eine bestimmte Zeile gewählt hat, ist mit horizontalen Pfeilen ( $\rightarrow$  und  $\leftarrow$ ) rechts neben den Auszahlungen des B ausgedrückt. Das Resultat ist, dass  $D_2$ , wenn  $D_1$  kolludiert, konkurrieren würde:

|                      |              | D <sub>2</sub> wählt |             |
|----------------------|--------------|----------------------|-------------|
|                      |              | kleine Menge         | große Menge |
| D <sub>1</sub> wählt | kleine Menge | 3, 3 $\rightarrow$   | 1, <u>4</u> |
|                      | große Menge  | 4, 1                 | 2, 2        |

Tabelle 4.4: Vorzugsrelation für D2

Wenn  $D_1$  konkurrieren würde, würde  $D_2$  ebenfalls konkurrieren wollen. Auch für  $D_2$  ist Wettbewerb mithin eine dominante Strategie.

|                      |              | D <sub>2</sub> wählt      |             |
|----------------------|--------------|---------------------------|-------------|
|                      |              | kleine Menge              | große Menge |
| D <sub>1</sub> wählt | kleine Menge | 3, 3                      | 1, 4        |
|                      | große Menge  | 4, <u>1</u> $\rightarrow$ | 2, <u>2</u> |

Tabelle 4.5: Dominante Strategie für D2

- 187 Die Abweichungspfeile, die notieren, ob ein Spieler einen Anreiz hat, von einer Entscheidung abzuweichen, ergeben zusammen betrachtet ein sogenanntes **Abweichungsdiagramm**:

|                      |              | D <sub>2</sub> wählt            |                     |
|----------------------|--------------|---------------------------------|---------------------|
|                      |              | kleine Menge                    | große Menge         |
| D <sub>1</sub> wählt | kleine Menge | $\downarrow$ 3, 3 $\rightarrow$ | $\downarrow$ 1, 4   |
|                      | große Menge  | 4, 1 $\rightarrow$              | <u>2</u> , <u>2</u> |

Tabelle 4.6: Abweichungsdiagramm

- 188 Das Diagramm zeigt, dass das Spiel ein **Gleichgewicht in dominanten Strategien** hat, bei dem beide Duopolisten die große Menge wählen, also nicht kolludieren. Man sieht auch, dass die Duopolisten das für sie vorteilhafte kollusive Ergebnis nicht erreichen können. Individuelle Rationalität und soziale Wohlfahrt treten hier auseinander, wie es für das **Gefangenen-Dilemma** im Speziellen und für Kooperationsprobleme im Allgemeinen kennzeichnend ist. Das **Kartell-Dilemma** ist allerdings insoweit ein Sonderfall, weil es ein sozial erwünschtes Dilemma ist. Denn wegen der (in der Spielmatrix nicht abgebildeten) negativen externen Effekte auf die Konsumenten und die Volkswirtschaft geht die Wohlfahrt der Duopolisten hier zu Lasten Dritter, so dass das Scheitern von **Kooperation** im Sinne des

Gemeinwohls liegen kann. Eine solche Kooperation, die für Dritte oder die Gesellschaft schädlich ist, wird häufig auch als Kollusion bezeichnet.<sup>14</sup>

Genauer gesagt hat man es im Gefangenen-Dilemma mit einem Fall der sogenannten **starken Dominanz** zu tun, die von der sogenannten **schwachen Dominanz** unterschieden wird. Stark dominante Strategien sind gegenüber jeder Strategie des anderen Spielers eine beste Wahl. Es kann aber auch sein, dass eine Strategie nur gegenüber manchen Strategien des anderen Spielers die beste Wahl ist, der Spieler gegenüber anderen Strategien aber lediglich indifferent ist. Wenn eine Strategie in diesem Sinn in mindestens einem Fall zu einem besseren, und in keinem anderem Fall zu einem schlechteren Resultat führt, spricht man von einer schwach dominanten Strategie. Ein Beispiel dafür gibt die nachfolgende Abwandlung des Kartell-Dilemmas. Der Unterschied betrifft hier den Fall, in dem  $D_1$  die große Menge wählt (die untere Reihe spielt). Anders als im Ausgangsfall macht es jetzt für  $D_2$  keinen Unterschied, ob er in Reaktion darauf die kleine oder die große Menge wählt. Denn seine Auszahlungen betragen jeweils 2. Nach dem Konzept der schwachen Dominanz ist es in einer solchen Konstellation für  $D_2$  dennoch rational, die große Menge zu wählen, weil er dadurch nur gewinnen und nichts verlieren kann.

189

|             |              | $D_2$ wählt  |             |
|-------------|--------------|--------------|-------------|
|             |              | kleine Menge | große Menge |
| $D_1$ wählt | kleine Menge | ↓ 3, 3 →     | ↓ 1, 4      |
|             | große Menge  | 4, 2         | <u>2, 2</u> |

Tabelle 4.7: Beispiel für schwache Dominanz

Das Konzept der Dominanz kann bei der Lösung von Spielen in zweierlei Weise helfen. Hat ein Spieler eine (stark oder schwach) **dominante Strategie**, ist die beste Strategie für ihn gefunden. Bei Spielen mit mehr als zwei Strategien kann es aber auch vorkommen, dass ein Spieler zwar keine dominante Strategie hat, aber eine Strategie im Vergleich zu anderen Strategien immer schlechter ist (oder zumindest in einem Fall schlechter und in keinem besser). Man spricht dann von (stark oder schwach) **dominierten Strategien**. Das Konzept der Dominanz fordert dann die Eliminierung dominierter Strategien, die für die Suche nach der besten Antwort ausscheiden, und bietet insoweit zumindest einen ersten Schritt für die weitere Lösung des Spiels.

190

<sup>14</sup> Vgl. zum Begriff der Kooperation unten Rz. 218.

### b. Nash-Gleichgewicht

- 191 Auch das sogenannte *Nash-Gleichgewicht* beruht auf dem Gedanken, solche Strategiekombinationen als Gleichgewichte anzusehen, bei denen keiner der Spieler einen Anreiz hat, von seiner Gleichgewichts-Strategie abzuweichen. Es ist aber insoweit ein schwächeres Lösungskonzept als das der Dominanz, als es nicht voraussetzt, dass die Gleichgewichtsstrategie unabhängig von dem Verhalten der Mitspieler eine beste Antwort sein muss, sondern nur unter der Voraussetzung, dass die anderen Mitspieler ebenfalls ihre Gleichgewichtsstrategien spielen. Genauer gesagt ist ein Strategieprofil ein *Nash-Gleichgewicht*, wenn die Gleichgewichtsstrategie jedes Spielers jeweils die beste Antwort auf die Gleichgewichtsstrategien aller anderen Spieler ist, so dass kein Spieler einen Anreiz hat, von seiner Strategie einseitig abzuweichen, vorausgesetzt die anderen spielen ihre Gleichgewichtsstrategien. Im *Nash-Gleichgewicht* stützen sich die Entscheidungen der Spieler also gegenseitig.
- 192 Das Konzept des *Nash-Gleichgewichts* lässt sich am Beispiel der Festlegung von Technologiestandards veranschaulichen.<sup>15</sup> Nehmen wir an, zwei Firmen (A und B) haben Nachfolgetechnologien für DVDs entwickelt, Firma A den Standard X, und Firma B den Standard Y. Beide Firmen können entweder ihren eigenen Standard am Markt einführen oder gegen eine Lizenzgebühr den Standard des Konkurrenten übernehmen. Werden beide Technologien nebeneinander am Markt eingeführt, ist wegen der Verunsicherung der Verbraucher zu erwarten, dass die Nachfrage nach den neuen Technologien kleiner ausfällt, so dass A wie B den geringsten Gewinn erhalten (im Sinne einer ordinalen Präferenzordnung beziffern wir ihn mit 1).<sup>16</sup> Kann eine Firma ihren eigenen Standard durchsetzen, erhält sie den größten Gewinn (wir beziffern ihm mit 3). Die andere Firma erhält in diesem Fall zwar auch einen größeren Gewinn als bei einem Nebeneinander der Standards, dieser fällt aber wegen der Lizenzgebühren geringer aus (wir beziffern ihn mit 2). Daraus ergibt sich folgendes Spiel:

<sup>15</sup> Vgl. *Besen/Farrell*, Choosing of how to compete: Strategies and tactics in standardization, *Journal of Economic Perspectives* 8 (1994), 117 ff.; *Farrell/Saloner*, Coordination Through Committees and Markets, *RAND Journal of Economics* 19 (1988), 235 (237 f.); *Baird et al.*, Game Theory and the Law, 1994, S. 208 ff.

<sup>16</sup>  $H_i(X, Y) = H_i(Y, X) = 1$ .

|               |            | Firma B wählt |             |
|---------------|------------|---------------|-------------|
|               |            | Standard X    | Standard Y  |
| Firma A wählt | Standard X | <u>3, 2</u>   | ↓ 1, ← 1    |
|               | Standard Y | ↑ 1, 1 →      | <u>2, 3</u> |

Tabelle 4.8: Spielmatrix des Standardisierungs-Spiels

Wieder kann man fragen, welche Entscheidung für den einen Spieler vor-  
 zugswürdig ist, gegeben dass sich der andere so oder so verhält, und die  
 Vorzugsrelationen mit Abweichungspfeilen kennzeichnen. Man sieht  
 dann, dass es in dem Spiel keine dominanten Strategien gibt: So ist für  
 Firma A der eigene Standard X vorzugswürdig, aber nur wenn auch Firma  
 B diesen Standard wählt (Zelle oben links). Wählt Firma B dagegen ihren  
 eigenen Standard Y, zieht auch Firma A diesen Standard vor (Zelle unten  
 rechts). Sowohl Standard X als auch Standard Y sind damit ein  
*Nash*-Gleichgewicht. Auch diese sind in der Matrix mit Hilfe der Abwei-  
 chungsdigramme zu erkennen. Zellen, von denen zumindest ein Abwei-  
 chungspfeil weg zeigt, sind keine Gleichgewichte, weil der Abweichungs-  
 Pfeil anzeigt, dass wenigstens ein Spieler einen Anreiz hat, eine andere  
 Strategie zu wählen. Zellen, auf die die Pfeile beider Spieler zeigen oder  
 von denen zumindest kein Pfeil weg zeigt, sind *Nash*-Gleichgewichte. Man  
 kann sie durch Unterstreichung der Auszahlungen markieren.

193

Unser Standardisierungs-Spiel hat mit seinen zwei Gleichgewichten kei-  
 ne eindeutige Lösung. Das Lösungskonzept des *Nash*-Gleichgewichts be-  
 sagt hier nur, dass es nicht rational wäre, kein Gleichgewicht zu spielen,  
 wenn man davon ausgeht, dass sich auch der andere Spieler rational ver-  
 hält und seinerseits ein Gleichgewicht spielt. In Spielen mit mehreren  
 Gleichgewichten wird damit die **Gleichgewichtsauswahl** zum zentralen  
 Problem. Man spricht deshalb von **Koordinierungsspielen**. Die beiden  
 Gleichgewichte unseres Standardisierungs-Spiels zeichnen sich weiter da-  
 durch aus, dass sie von den Firmen A und B unterschiedlich bewertet wer-  
 den: Firma A hätte lieber X als Standard, aber wenn Y als Standard fest-  
 steht, hat sie keinen Anreiz, einseitig an X als Standard festzuhalten. Ne-  
 ben dem Koordinierungsproblem ist das Spiel also mit einem Konflikt  
 belastet. Spiele dieser Art sind unter dem Namen „**Kampf der Geschlech-**  
**ter**“ bekannt.

194

### c. Gleichgewichte in gemischten Strategien

Im vorstehenden Beispiel haben wir angenommen, dass die Spieler im  
 Gleichgewicht aus den verfügbaren Strategien (Standard X, Standard Y)

195



eine auswählen (entweder Standard X oder Standard Y). Man spricht insoweit von einem Gleichgewicht in **reinen Strategien**. Die Spieltheorie lässt aber auch sogenannte **gemischte Strategien** zu, bei denen die Spieler aus mehreren reinen Strategien nach einem Zufallsmechanismus mit bestimmten Wahrscheinlichkeiten auswählen (z.B. Standard X mit einer Wahrscheinlichkeit von  $1/3$ , Standard Y mit einer Wahrscheinlichkeit von  $2/3$ ).<sup>17</sup> Auch solche gemischten Strategien können in einem *Nash*-Gleichgewichte stehen.

- 196 Man denke zum Beispiel an den Kampf der Polizei gegen Drogenumschlagsplätze. Nehmen wir im Sinne einer vereinfachenden Modellierung an, in einer Stadt kämen zwei Orte in Betracht, an denen sich Drogendealer und Drogenkonsumenten zur Abwicklung ihrer Geschäfte treffen können, nämlich A und B. Die Einsatzkräfte der Polizei reichen aber nur aus, um zu einem Zeitpunkt an einem Ort zu kontrollieren. Sofern die Polizei den Handel nicht stillschweigend duldet, um ihn unter Kontrolle zu behalten, „spielen“ Dealer und Polizei in dieser Situation ein **Diskoordinierungsspiel**:<sup>18</sup> Die Polizei will dort kontrollieren, wo der Drogenhandel stattfindet, aber die Dealer wollen genau dort handeln, wo die Polizei nicht kontrolliert. Die Matrix eines solchen Spiels könnte wie folgt aussehen:

|                         |   | Droghenhändler dealen in |         |
|-------------------------|---|--------------------------|---------|
|                         |   | A                        | B       |
| Polizei kontrolliert in | A | 2, -4 →                  | ↓ 0, 4  |
|                         | B | ↑ 0, 4                   | 2, ← -4 |

Tabelle 4.9: Spielmatrix des Diskoordinierungsspiels

- 197 Dieses Spiel hat kein strategisches Gleichgewicht in reinen Strategien. Denn für jede der möglichen Kombinationen aus reinen Strategien gilt, dass einer besser daran tut, seine Strategie zu ändern. Kontrolliert die Polizei in A, ist es besser, in B zu dealen. Wird in B gedealt, ist es besser, in B zu kontrollieren. Wird in B kontrolliert, ist es besser, in A zu dealen, usw. Das bedeutet aber nicht, dass es für die Beteiligten gleichgültig ist, wie sie sich verhalten. Für sie kommt es vielmehr darauf an, zu verhindern, dass das eigene Verhalten vom anderen Spieler vorhergesehen werden kann, also unberechenbar zu sein. Wüssten die Dealer zum Beispiel, dass die

<sup>17</sup> Die Summe der Wahrscheinlichkeiten der in einer gemischten Strategie gespielten reinen Strategien muss immer 1 betragen.

<sup>18</sup> Das klassische Beispiel für ein Diskoordinierungsspiel ist das *Matching-Pennies*-Spiel, vgl. *Fudenberg/Tirole*, Game Theory, 1996, S. 16 f.

Polizei immer in A kontrolliert, würden sie nach B ausweichen. In Situationen dieser Art ist es für die Kontrahenten am besten, nach einem Zufallsmuster mal die eine und mal die andere Alternative zu wählen, das heißt, ihre Strategien zu „mischen“. Dabei ist es nicht immer sinnvoll, alle Strategien mit der gleichen Wahrscheinlichkeit zu wählen. Wenn etwa Kontrollen für die Polizei in B in der Regel erfolgreicher sind, weil in B besser kontrolliert werden kann, können häufigere Kontrollen in B ratsam sein. Für die Drogendealer kann es dann wiederum besser sein, häufiger in A zu dealen. Passen die Akteure die Wahrscheinlichkeiten ihres Verhaltens so lange an, bis keiner der Beteiligten seine Situation durch Veränderung der Wahrscheinlichkeiten weiter verbessern kann, dann spielen sie ein **Gleichgewicht in gemischten Strategien**. Im Beispiel beinhaltet es, dass sowohl die Polizei als auch die Drogenhändler die Orte A und B mit der Wahrscheinlichkeit von  $\frac{1}{2}$  wählen. Gleichgewichte in gemischten Strategien kann man nicht mit Hilfe der Abweichungspfeile ermitteln, sondern nur ausrechnen.<sup>19</sup>

### 3. Soziale Wohlfahrt und politische Gemeinwohlziele

Lösungskonzepte wie das *Nash*-Gleichgewicht sagen etwas darüber, welche Lösung eines Spiels individuell rational ist. Diese Lösung muss aber weder den kollektiven Interessen der Spieler noch den Interessen Dritter oder einem anders bestimmten Gemeinwohl entsprechen. Die Definition eines Spiels als solche berücksichtigt insoweit nur die Präferenzen der Spieler. Die **Wohlfahrtskriterien** der *Pareto*-Effizienz und der *Kaldor-Hicks*-Effizienz ermöglichen es aber, auf Grundlage der individuellen Präferenzen bestimmte Spielergebnisse als sozial wünschenswert auszuzeichnen und zu fragen, ob diese wünschenswerten Ergebnisse erreicht werden. Die gleiche Frage lässt sich stellen, wenn sozial wünschenswerte Ergebnisse nach anderen, politischen Kriterien festgelegt werden.

198

#### a. Pareto-Optimalität

Nach dem ***Pareto-Kriterium***<sup>20</sup> stellt ein Spielausgang gegenüber einem anderen eine Verbesserung der sozialen Wohlfahrt dar, wenn wenigstens ein Spieler höhere Auszahlungen erhält (d. h., dieses Ergebnis gegenüber dem anderen vorzieht), und kein Spieler geringere Auszahlungen hat. Dieser

199

<sup>19</sup> Vgl. *Berninghaus et al.*, Strategische Spiele, 2010, S. 29 ff.; *Holler/Illing/Napel*, Einführung in die Spieltheorie, 2016, S. 65 ff.

<sup>20</sup> Hierzu Rz. 89 ff.

Spielausgang stellt dann gegenüber dem anderen eine *Pareto*-Verbesserung dar. Sind von einem Spielausgang keine *Pareto*-Verbesserungen mehr möglich, ist dieser Spielausgang *pareto-optimal*. Im Standardisierung-Spiel sind beide Gleichgewichte zugleich *pareto-optimal*, während die Spielausgänge, bei denen die Firmen unterschiedliche Standards verwenden, *pareto-inferior* sind. Der Übergang von einem gemeinsamen Standard zu dem anderen würde dagegen keine *Pareto*-Verbesserung bedeuten, sondern nur eine Umverteilung, durch die eine Firma besser und die andere schlechter stünde.

|               |            | Firma B wählt       |                     |
|---------------|------------|---------------------|---------------------|
|               |            | Standard X          | Standard Y          |
| Firma A wählt | Standard X | <u>3</u> , <u>2</u> | ↓ 1, ← 1            |
|               | Standard Y | ↑ 1, 1 →            | <u>2</u> , <u>3</u> |

Tabelle 4.10: Beispiel für *Pareto*-Optimalität

#### b. Kaldor-Hicks-Effizienz

- 200 Nach dem *Kaldor-Hicks-Kriterium*<sup>21</sup> können Nutzenverbesserungen für manche selbst dann noch als Wohlfahrtsverbesserung angesehen werden, wenn dadurch andere schlechter gestellt würden, nämlich wenn die Gewinner die Verbesserungen höher bewerten als die Verlierer die Nachteile. Das Problem dabei ist, wie dies ohne einen intersubjektiven Nutzenvergleich festgestellt werden kann, der in der Ökonomie weithin als unzulässig angesehen wird. Um einen solchen zu vermeiden, stellt das *Kaldor-Hicks-Kriterium* auf eine *hypothetische Pareto-Verbesserung* ab, nämlich darauf, ob die Gewinne der Gewinner ausreichen würden, um die Verluste der Verlierer in Geld zu kompensieren. Im Beispiel der Standardisierung wäre Standard X etwa gegenüber Standard Y eine Effizienz-Verbesserung, falls die Auszahlungen in der Matrix die Zahlungsbereitschaft der Firmen wiedergäben und wenn diese Zahlungsbereitschaften bei Standard X (5, 2) betrügen, im Gegensatz zu (2, 3) bei Standard Y. Firma A könnte dann Firma B hypothetischerweise mit 1 Zahlungseinheit kompensieren. Dadurch würden die Auszahlungen für Standard X (4, 3) betragen, was gegenüber den Auszahlungen von (2, 3) A besserstellt, aber B nicht schlechter. Nach dem *Kaldor-Hicks-Kriterium* hat dieses Spiel dann ein effizientes und ein ineffizientes Gleichgewicht.

<sup>21</sup> Vgl. Rz. 93 f.

|               |            | Firma B wählt       |                     |
|---------------|------------|---------------------|---------------------|
|               |            | Standard X          | Standard Y          |
| Firma A wählt | Standard X | <u>5</u> , <u>2</u> | ↓ 1, ← 1            |
|               | Standard Y | ↑ 1, 1 →            | <u>2</u> , <u>3</u> |

Tabelle 4.11: Beispiel für *Kaldor-Hicks*-Effizienz

### c. Politische Gemeinwohlziele

Außer die bekannten Wohlfahrtskriterien anzulegen, kann man im Grunde jeden beliebigen Spielausgang aufgrund politischer Gemeinwohlziele oder anderer Kriterien herausgreifen und mit den Mitteln der Spieltheorie fragen, ob dieses Ergebnis der individuellen Rationalität entspricht oder nicht. So kann es etwa sein, dass Standard Y zwar in Bezug auf den Gewinn der Firmen A und B ineffizient ist, aber für die Konsumenten Vorteile bringt, die die Nachteile der Firmen weit überragen (die aber in der Spiel-Matrix nicht abgebildet werden). Man kann dann fragen, welche Mechanismen es gibt, damit sich die Firmen auf diesem Standard koordinieren, und nicht auf dem anderen. 201

### d. Mechanismus-Design

Bisher wurde eine gegebene Interesseninterdependenz als Spiel modelliert, um dann zu fragen, welches Verhalten in dieser Situation individuell rational ist, und ob das individuell Rationale auch sozial wünschenswert ist. Das sogenannte Mechanismus-Design geht umgekehrt von dem sozialen Ziel aus und fragt, welche Regeln dafür sorgen, dass das soziale Ziel auch erreicht wird. Das Mechanismus-Design entwirft also Spiele, um bestimmte soziale Ziele zu implementieren.<sup>22</sup> 202

## III. Typen von Spielen

Mit dem Kartell-Dilemma, dem Standardisierungs-Spiel und dem Polizei- vs. Dealer-Spiel haben wir schon drei verschiedene Spiele kennen gelernt, nämlich ein „Gefangenen-Dilemma“, einen „Kampf der Geschlechter“ und eine Variante von „*Matching Pennies*“. Im Folgenden werden wir anhand von **2x2-Bimatrix-Spielen** eine Typologie wichtiger Spiele vorstellen. Diese Typologie soll einen ersten Überblick darüber geben, welche Arten 203

<sup>22</sup> Maskin, Mechanism Design: How to Implement Social Goals, American Economic Review 98 (2008), 567 ff.

von Problemen aus der Interdependenz von Interessen für die einzelnen Akteure und im Hinblick auf Wohlfahrts- oder Gemeinwohlziele erwachsen, und einen Eindruck davon vermitteln, wie die Spieltheorie die Funktionen des Rechts und anderer Institutionen behandelt. Wesentlich für die Natur eines Spiels ist zunächst, in welcher Rangfolge die Spieler die verschiedenen Spielergebnisse einander vorziehen. Wichtig ist darüber hinaus, ob und welche **Gleichgewichte** ein Spiel hat und ob diese Gleichgewichte der **sozialen Wohlfahrt** entsprechen.

### 1. Einfache Motive

- 204 Von einfachen Motiven kann man sprechen, wenn die Interessen der Akteure völlig gleichläufig oder völlig gegensätzlich sind. So liegt es bei Harmoniespielen und reinen Koordinationsspielen einerseits und reinen Konfliktspielen andererseits.

#### a. Harmonie

- 205 **Harmoniespiele** wie das folgende zeichnen sich durch die Abwesenheit jeglicher Probleme aus.

|   |            | B             |            |
|---|------------|---------------|------------|
|   |            | kooperiert    | defektiert |
| A | kooperiert | <u>4, 4</u> ← | 3, 3       |
|   | defektiert | ↑ 2, 2 ←      | ↑ 1, 1     |

Tabelle 4.12: Beispiel für ein Harmonie-Spiel

Beide Spieler bewerten hier die Spielausgänge in der gleichen Reihenfolge, es gibt also keinerlei Konflikt.<sup>23</sup> Das Spiel hat nur ein Gleichgewicht (in dominanten Strategien), so dass sich auch kein Koordinationsproblem

<sup>23</sup> Für die Natur des Spiels kommt es nur auf diese Reihenfolge und nicht auf die Höhe der Auszahlungen an. Legt man ordinale Auszahlungen zugrunde, sagen diese ohnehin nichts zu der Höhe des Nutzens aus. Legt man dagegen, wie hier, kardinale Auszahlungen zugrunde, etwa Zahlungsbereitschaften, können diese in der Höhe durchaus weit abweichen, ohne dass dies die Natur des Spiels ändert, soweit die Reihenfolge der Bewertungen, wie in dem folgenden Harmonie-Spiel, gleich bleibt.

|   |            | B                |            |
|---|------------|------------------|------------|
|   |            | kooperiert       | defektiert |
| A | kooperiert | <u>5.000, 20</u> | 1.000, ← 5 |
|   | defektiert | ↑ 2.000, 7       | ↑ 500, ← 1 |

stellt. Dieses Gleichgewicht ist zudem *pareto*-optimal, so dass individuelle Rationalität und soziale Wohlfahrt nicht in Konflikt geraten.

## b. Konflikt

Reine **Konfliktspiele** haben in verschiedener Hinsicht die gegenteiligen Eigenschaften von Harmoniespielen. Die Akteure bewerten in ihnen die Spielausgänge nicht in der gleichen, sondern in entgegengesetzter Reihenfolge, wie in folgendem Beispiel. 206

|   |            | B          |             |
|---|------------|------------|-------------|
|   |            | kooperiert | defektiert  |
| A | kooperiert | 4, 1 →     | <u>3, 2</u> |
|   | defektiert | ↑ 2, 3 →   | ↑ 1, 4      |

Tabelle 4.13: Beispiel für ein reines Konfliktspiel

Dieses Spiel ist ein Beispiel für ein Konstantsummenspiel, zu denen auch 207 das bekannte **Nullsummenspiel** gehört.<sup>24</sup> In diesen Spielen ist die Summe der Auszahlungen in jedem Spielausgang gleich, weshalb jede Verbesserung für einen Spieler zwangsläufig eine Verschlechterung für einen anderen Spieler bedeutet. So kann das Recht die Früchte eines Grenzbaumes nur entweder dem einen oder dem anderen Nachbarn zusprechen, es kann die Früchte aber nicht vermehren (§ 923 BGB verteilt sie zu gleichen Teilen). Da jede Veränderung in Konfliktspielen notwendig einen anderen schlechter stellt, gibt es keine Möglichkeit für *Pareto*-Verbesserungen, weshalb **jeder Spielausgang *pareto*-optimal** ist. Manche Konfliktspiele, wie das Polizei vs. Dealer-Spiel, haben kein Gleichgewicht (in reinen Strategien), sind also instabil, und manche, wie das obige Beispiel, haben eines. Aus juristischer Sicht ist das Interessanteste an reinen Konfliktspielen vielleicht, dass es sie so selten gibt, das heißt, dass kaum eine der vom Recht behandelten Situationen bei lebensnaher Betrachtung angemessen als ein reines Konfliktspiel modelliert werden kann. Spätestens ein Rechtsstreit über den Konflikt verursacht zusätzliche materielle und immaterielle Kosten, deren Vermeidung ein partiell gleichgerichtetes Interesse begründen kann (dazu unten das sogenannte Falke-Taube-Spiel).

<sup>24</sup> *Straffin*, Game Theory and Strategy, 2004, S. 5. Bei Nullsummenspielen ist die Summe der Auszahlungen der Spieler bei allen Spielergebnissen gleich Null, bei Konstantsummenspielen wie in obigem Beispiel (die zu Nullsummenspielen äquivalent sind) ist die Summe der Auszahlungen der Spieler bei allen Spielergebnissen konstant. Vorausgesetzt ist jeweils, dass die Auszahlungen kardinal und vergleichbar sind.

## c. Koordination

- 208 Bei reinen Koordinationsspielen haben die Spieler keinen Konflikt, müssen aber zwischen mehreren Gleichgewichten wählen. Ihre Interessen sind also gleichgerichtet, aber nicht eindeutig. Ein bekanntes Beispiel ist die Frage, auf welcher Seite im Straßenverkehr gefahren wird. In der Sache ist es unerheblich, ob man sich für Rechts- oder Linksverkehr entscheidet, aber die Entscheidung sollte einheitlich sein, um Unfälle zu vermeiden. Man kann das Problem wie folgt modellieren:

|   |        | B           |             |
|---|--------|-------------|-------------|
|   |        | links       | rechts      |
| A | links  | <u>2, 2</u> | ↓ 1, ← 1    |
|   | rechts | ↑ 1, 1 →    | <u>2, 2</u> |

Tabelle 4.14: Beispiel für ein reines Koordinationsspiel

- 209 Das Spiel hat zwei Gleichgewichte, zwischen denen weder ein Konflikt noch ein Wohlfahrtsgefälle besteht. Man spricht dann von **reinen Koordinationsspielen**. Ein Problem werfen reine Koordinationsspiele insofern auf, als es den Handelnden gelingen muss, sich auf einem der Gleichgewichte zu koordinieren. Wie einfach oder schwierig das ist, wird von verschiedenen Faktoren beeinflusst. Können die Beteiligten miteinander kommunizieren, können sie leicht ein Gleichgewicht festlegen (da es keinen Konflikt über die Auswahl gibt). Handeln die Beteiligten nacheinander und können sie ihre Handlungen beobachten, dann kann der, der zuerst handelt, das Gleichgewicht vorgeben, und es gibt keinen Grund, weshalb die anderen ihm nicht folgen sollten. In Interaktionen, die sich häufig mit den gleichen Beteiligten wiederholen, wird sich mit der Zeit ein Gleichgewicht etablieren, das dann relativ stabil sein wird. Interagieren die Handelnden nur einmal und können sie vorher nicht kommunizieren, kann Koordination dagegen zum Problem werden. Das gilt umso mehr, je mehr Gleichgewichte in Betracht kommen, etwa wenn man sich in einer fremden Stadt treffen will, aber keinen Treffpunkt abgemacht hat. Jeder mögliche Ort ist dann ein Gleichgewicht. Bei der Auswahl kann es dann helfen, ein Gleichgewicht zu wählen, das aus irgendeinem Grund die Aufmerksamkeit der Akteure auf sich zieht (sogenannte **Brennpunkte** oder *focal points*). In der Hoffnung, dass der andere auf die gleiche Idee kommt, könnte man zum Beispiel am Hauptbahnhof warten. Gerade in modernen Massengesellschaften haben nicht wenige Institutionen die Aufgabe, solche Brennpunkte zu setzen, um eine Koordination zu erleichtern, wenn persönliche Übereinkünfte fehlen. An Bahnhöfen und Flughäfen gibt es zum Beispiel

in der Regel ausgeschilderte Treffpunkte, die als „Brennpunkte“ die wechselseitige Koordination bei fehlenden Absprachen erleichtern. Eine vergleichbare Funktion können z.B. Straßenverkehrsregeln wahrnehmen, soweit die Straßenverkehrsteilnehmer ein gemeinsames Interesse daran haben, Unfälle zu vermeiden.<sup>25</sup>

## 2. Gemischte Motive

Spiele zwischen den Extremen völlig gleichlaufender und völlig gegensätzlicher Interessen nennt man Spiele mit gemischten Motiven. Diese, meist unter blumigen Namen firmierenden, Spiele modellieren verschiedene Situationen, in denen Interessen teils gleich und teils gegeneinander laufen. Zu ihnen gehören der „Kampf der Geschlechter“, das „Falke-Taube-Spiel“, die „Hirschjagd“ und das „Rambo-Spiel“, aber auch Kooperationsspiele wie das „Gefangenen-Dilemma“.

### a. Kampf der Geschlechter

Den Kampf der Geschlechter haben wir bereits in Gestalt eines Standardisierungs-Spiels kennen gelernt: Zwei Firmen möchten sich auf einen gemeinsamen Technologiestandard einigen, bevorzugen aber jeweils den eigenen Standard.<sup>26</sup> Spiele dieser Art sind gemischte Koordinations- und Konfliktspiele. Sie haben zwei Gleichgewichte, die von beiden Spielern gegenüber den anderen Spielergebnissen vorgezogen werden und die auch *pareto*-optimal sind. Insoweit sich die Spieler auf einem dieser Gleichgewichte treffen wollen, handelt es sich um ein Koordinationsspiel. Hinsichtlich der Auswahl des Gleichgewichts besteht aber ein Konflikt, der nun die Koordination gefährdet, weil jeder versucht, das ihm günstige Gleichgewicht zu etablieren. Das Problem derartiger Situationen liegt in den sozialen Kosten des Koordinationsversagens, nämlich in der Gefahr, dass eine Koordination an den divergierenden Interessen scheitert und die Interak-

<sup>25</sup> Ausführlich zur Konzeption von Recht als Brennpunkt *McAdams*, A Focal Point Theory of Expressive Law, Virginia Law Review 86 (2000), 1649 ff.; *ders.*, The Expressive Power of Adjudication, Illinois Law Review 5 (2005), 1043 ff.

<sup>26</sup> Mit folgender Spielmatrix:

|               |            | Firma B wählt |             |
|---------------|------------|---------------|-------------|
|               |            | Standard X    | Standard Y  |
| Firma A wählt | Standard X | <u>3, 2</u>   | ↓ 1, ← 1    |
|               | Standard Y | ↑ 1, 1 →      | <u>2, 3</u> |



tion deshalb einen individuell wie kollektiv unerwünschten Ausgang nimmt.

### b. Falke-Taube-Spiel

- 212 Nicht wenige Konflikte, die sich bei isolierter Betrachtung als Nullsummenspiele darstellen, erscheinen in einem anderen Licht, wenn man mitberücksichtigt, dass die Austragung des Konflikts selbst Schäden anrichtet, sei es, weil die Austragung des Konflikts materielle oder immaterielle Kosten verursacht, sei es, dass der Konflikt die Vertrauensbasis zwischen den Parteien zerstört und dadurch die Fortsetzung einer Kooperation verhindert. Dieser auch für viele Rechtsprobleme typische Aspekt von Konflikten lässt sich mit dem Falke-Taube-Spiel abbilden. Das **Falke-Taube-Spiel** ist, wie der Kampf der Geschlechter, ein gemischtes Koordinations- und Konfliktspiel. In ordinalen Präferenzen sieht es wie folgt aus:

|   |       | B                   |                     |
|---|-------|---------------------|---------------------|
|   |       | Taube               | Falke               |
| A | Taube | ↓ 3, 3 →            | <u>2</u> , <u>4</u> |
|   | Falke | <u>4</u> , <u>2</u> | ↑ 1, ← 1            |

Tabelle 4.15: Beispiel für ein Falke-Taube-Spiel

- 213 Nehmen wir an, zwei Parteien befinden sich im Streit. Sind beide unnachgiebig, spielen also beide „Falke“, eskaliert der Konflikt und beide erleiden erheblichen Schaden. Das ist der ungünstige Ausgang des Spiels, den beide am liebsten vermeiden würden. Für beide wäre es besser, wenn sie den Streit durch einen Kompromiss friedlich beilegen, also beide „Taube“ spielten. Das ist freilich kein Gleichgewicht, weil die Kompromissbereitschaft durch den anderen Spieler ausgenutzt werden kann. Denn gegenüber einer „Taube“ kann man gefahrlos „Falke“ spielen und den Konflikt zum eigenen Vorteil entscheiden. Spielt der andere aber „Falke“, kann man die eigene Position nicht mehr dadurch verbessern, dass man selbst auch zu kompromisslosem Verhalten übergeht, denn das würde zur Eskalation führen. Das perfide an derartigen Konstellationen ist, dass beiderseitiges Nachgeben kein Gleichgewicht, also nicht individuell rational ist. Nur die asymmetrischen Konstellationen, in denen einer unnachgiebig und der andere nachgiebig ist, sind *Nash*-Gleichgewichte. Da aber jeder versuchen wird, das für sich günstigere Gleichgewicht durchzusetzen, enden beide leicht im schlechtesten der möglichen Zustände. Für zwei Falken nimmt das Spiel keinen glücklichen Verlauf. Auch beim Falke-Taube-Spiel liegt das soziale Problem mithin in den **Kosten des Koordinationsversagens**.

Dieses Problem wird umso dringlicher, umso höher der Schaden im Fall einer Eskalation ist, wie etwa beim atomaren Wettrüsten. Für solche Anwendungen wird das Spiel häufig mit hohen negativen Auszahlungen für den unkooperativen Ausgang modelliert und als **Feigling-Spiel** bezeichnet (*chicken game*).<sup>27</sup>

In Situationen wie dem Falke-Taube-Spiel kann es sinnvoll sein, wenn 214 das eigene Verhalten für den anderen Spieler nicht vorhersagbar ist. Das erreicht man, wenn man die Strategiewahl einem Zufallsmechanismus überlässt, also eine **gemischte Strategie** spielt. Statt sich *ex ante* für eine Strategie zu entscheiden („Taube“ oder „Falke“), könnte man sich zum Beispiel darauf festlegen, vor der Aktion zu würfeln und bei geraden Zahlen „Taube“ zu spielen, und bei ungeraden „Falke“. Interessant ist, welchen Einfluss die Höhe des drohenden Schadens darauf hat, mit welcher Wahrscheinlichkeit die unnachgiebige und die nachgiebige Strategie gespielt werden sollte: Umso höher der drohende Schaden ist, umso häufiger sollten beide Spieler „Taube“ spielen, und umso seltener „Falke“. Mit ansteigendem Schaden konvergieren deshalb die gemischten Strategien in Richtung beiderseitigen Nachgebens. Das Verhalten nähert sich dann der sogenannte **Maximin-Lösung** an, bei der die Spieler einen denkbaren Verlust möglichst gering halten.

### c. Hirschjagd

Die Hirschjagd ist nach einer Parabel von Rousseau benannt. Zwei Jäger 215 können entweder gemeinsam einen Hirsch jagen, oder jeweils alleine Hasen. Von ihrem Anteil an einem Hirsch haben beide mehr als von einem Hasen. Es kann sich also keiner durch die Hasenjagd besser stellen (das ist ein wichtiger Unterschied zum Gefangenen-Dilemma). Lässt allerdings ein Jäger den anderen im Stich, um doch Hasen zu jagen, geht der andere leer aus, während er selbst den Hasen gewiss hat. Die Hirschjagd verspricht also einen höheren, aber unsicheren Ertrag, während die Hasenjagd einen sicheren, aber geringeren Ertrag bedeutet.

<sup>27</sup> Die Matrix eines Feigling-Spiels könnte wie folgt aussehen:

|   |           | B               |                               |
|---|-----------|-----------------|-------------------------------|
|   |           | Feigling        | Aggressor                     |
| A | Feigling  | ↓ <u>3, 3</u> → | 2, 4                          |
|   | Aggressor | 4, 2            | ↑ <u>-100</u> , ← <u>-100</u> |

|        |        | B jagt      |             |
|--------|--------|-------------|-------------|
|        |        | Hirsch      | Hasen       |
| A jagt | Hirsch | <u>6, 6</u> | ↓ 0, ← 2    |
|        | Hasen  | ↑ 2, 0 →    | <u>2, 2</u> |

Tabelle 4.16: Beispiel für ein Hirschjagd-Spiel

- 216 In einem solchen Spiel gibt es zwei Gleichgewichte, nämlich wenn beide gemeinsam einen Hirsch jagen oder wenn jeder für sich einen Hasen jagt. Im Unterschied zum Gefangenen-Dilemma ist hier also auch wechselseitige Kooperation ein Gleichgewicht. Dieses ist auch für beide Spieler mit den höchsten Auszahlungen verbunden, ist also *pareto*-optimal. Die Hirschjagd ist also ein Koordinationsspiel ohne Konflikt und ohne Widerspruch zwischen individueller und kollektiver Rationalität.
- 217 Das Problem des Spiels liegt aber darin, dass die Gleichgewichtsstrategien Hirschjagd und Hasenjagd unterschiedlich riskant sind, falls eine Koordination beider Spieler auf einem Gleichgewicht nicht zustande kommt. Denn wer Hasen jagt, erhält unabhängig vom Verhalten des anderen Jägers eine Auszahlung von 2, trägt also kein Risiko. Wer dagegen Hirsch jagt, geht leer aus, wenn er alleine bleibt. Das Problem liegt also in dem Risiko, beim Scheitern der Koordination auf dem kooperativen Gleichgewicht einen Verlust zu erleiden. Nach welchen Grundsätzen zwischen in diesem Sinn unterschiedlich riskanten Gleichgewichten ausgewählt werden soll, ist nicht eindeutig. Dafür bedarf es weiterer Kriterien: Nach dem Kriterium der sogenannte Auszahlungs- oder *Pareto*-Dominanz sollte man nach der Höhe der Auszahlungen auswählen. In der Hirschjagd sollte man sich danach in jedem Fall für das *pareto*-überlegene Hirsch-Gleichgewicht entscheiden. Legt man dagegen Wert darauf, einen Verlust so gering wie möglich zu halten, und entscheidet nach dem *Maximin*-Prinzip,<sup>28</sup> würde man immer das sicherere Hasen-Gleichgewicht bevorzugen. Nach dem Konzept der Risiko-Dominanz sollte man dagegen das Gleichgewicht mit dem höheren Erwartungsnutzen wählen, wobei der Erwartungsnutzen von der Wahrscheinlichkeit abhängt, mit welcher der andere Spieler sich für die eine oder die andere Strategie entscheidet. Welches Gleichgewicht nach dem Kriterium der Risiko-Dominanz gewählt wird, hängt dann von den **Risikopräferenzen** der Spieler ab und deren Annahmen darüber, mit welcher Wahrscheinlichkeit ihr Gegenüber kooperieren wird. Sind die Spieler risikoneutral und halten sie mangels näherer Anhaltspunkte beide Ent-

<sup>28</sup> Zum *Maximin*-Prinzip vgl. z.B. Holler/Illing/Napel, Einführung in die Spieltheorie, 2016, S. 55 f.

scheidungen ihres Gegenübers für gleich wahrscheinlich, dann liegt der Erwartungsnutzen im obigen Beispiel für die Hasen-Strategie bei 2, für die Hirsch-Strategie aber bei 3.<sup>29</sup> Hirsch ist dann gegenüber Hase risiko-dominant. Sind die Spieler allerdings hinreichend risikoavers, kann die Hase-Strategie risiko-dominant werden.

### 3. Kooperation

Ebenfalls zu den Spielen mit gemischten Motiven gehören die Kooperations- 218  
spiele, denen wir wegen ihrer praktischen Bedeutung einen eigenen Abschnitt widmen. Sie modellieren Situationen, in denen individuelle Rationalität und soziale Wohlfahrt im Widerspruch stehen, weil die individuell rationalen Gleichgewichte des Spiels *pareto*-nachteilig sind, während der oder die sozial erwünschten Spielausgänge keine Gleichgewichte sind.<sup>30</sup> Wenigstens ein Akteur hat dann einen individuellen Anreiz, vom sozial wünschenswerten Verhalten abzuweichen und ein *pareto*-unterlegenes Resultat herbeizuführen.<sup>31</sup> Der spieltheoretische Begriff der Kooperation ist also enger als das umgangssprachliche Verständnis von Kooperation, das alle möglichen Formen der gemeinschaftlichen Verfolgung von Zielen umfasst, die sich in spieltheoretischer Sicht aber nur zum Teil als Kooperations- 219  
spiele, zum Teil aber auch als Harmonie-, Konflikt- oder Koordinations-  
spiele darstellen.

#### a. Gefangenen-Dilemma

Das bekannteste Modell für ein Kooperationsproblem ist das Gefangenen- 219  
Dilemma, das uns bereits in Gestalt des Kartell-Dilemmas<sup>32</sup> begegnet ist. Im Gefangenen-Dilemma und allgemein in Kooperationsspielen bezeichnet man die Aktionen, die das sozial wünschenswerte Ergebnis herbeiführen (würden) als Kooperation, und die Aktionen, die zu dem individuell

<sup>29</sup> Wenn A Hirsch spielt und B ebenfalls Hirsch, erhält A eine Auszahlung von 6; wenn A Hirsch spielt und B Hase, erhält A eine Auszahlung von 0. Sind beide Möglichkeiten gleich wahrscheinlich, ist der Erwartungsnutzen des A für die Strategie Hirsch:  $(\frac{1}{2} \cdot 6) + (\frac{1}{2} \cdot 0) = 3$ .

<sup>30</sup> In der sozialwissenschaftlichen Literatur wird dies auch als Konflikt zwischen individueller und kollektiver Rationalität interpretiert.

<sup>31</sup> *Beckenkamp*, A game-theoretic taxonomy of social dilemmas, Central European Journal of Operations Research 3 (2006), 337 (338); ausführlich *ders.*, A Game Theoretic Taxonomy of Social Dilemmas; Preprint 2002/11, MPI for Collective Goods.

<sup>32</sup> Hierzu Rz. 177 ff.

rationalen Ergebnis führen, als Defektion oder opportunistisches Verhalten. **Defektion** ist im Gefangen-Dilemma eine **dominante Strategie**.

|   |             | B           |             |
|---|-------------|-------------|-------------|
|   |             | Kooperation | Defektion   |
| A | Kooperation | ↓ 3, 3 →    | ↓ 1, 4      |
|   | Defektion   | 4, 1 →      | <u>2, 2</u> |

Tabelle 4.17: Beispiel für ein Gefangen-Dilemma

- 220 Gefangen-Dilemmata sind aber nicht auf 2-Personen-Spiele beschränkt, sondern können auch als Mehr- oder n-Personen-Spiele auftreten. Man spricht dann auch von **Problemen kollektiven Handelns**. Auch bei vielen Personen ändert sich die grundsätzliche Anreizstruktur des Gefangen-Dilemmas nicht, weshalb man oft das 2-Personen-Spiel als vereinfachte Modellierung für Probleme kollektiven Handelns verwendet. Eine nicht unwichtige Ursache von Kooperationsproblemen sind negative oder positive **Externalitäten**. Damit ist gemeint, dass die Handlung eines Akteurs positive oder negative Folgen für den Nutzen eines anderen hat, die der Handelnde selbst (wie mit der Eigennutzannahme vermutet) bei seiner Entscheidung nicht berücksichtigt. Im Kartell-Dilemma zum Beispiel verursacht ein Duopolist, der seine Produktion steigert, beim anderen Duopolisten negative Externalitäten im Sinne von Verlusten<sup>33</sup> Zu dem für Kooperationsprobleme typischen Konflikt zwischen individueller und kollektiver Rationalität kommt es dann, weil der Einzelne dem anderen stärker schadet, als er sich selbst nutzt, im Beispiel-Spiel des Kartell-Dilemmas also, weil die Verluste des anderen Duopolisten höher sind als die Gewinne aus der Produktionssteigerung. Das Kartell-Dilemma ist allerdings in verschiedener Hinsicht ein spezieller Fall negativer Externalitäten. Erstens werden

<sup>33</sup> Dessen Auszahlungen entsprechen denen des Gefangen-Dilemmas hier im Text. Im Kartell-Dilemma entsprechen die summierten Angebotsmengen von  $D_1$  und  $D_2$ , wenn diese jeweils die kleine Menge wählen, der Angebotsmenge eines Monopolisten. Bei ihr ist der gemeinsame Gewinn von  $D_1$  und  $D_2$  am höchsten (im Beispiel bei 6). Jede Steigerung der Produktion führt demgegenüber zu einer Verringerung des gemeinsamen Gewinns (auf 5, wenn einer mehr produziert, auf 4, wenn beide mehr produzieren). Diese Einbuße fällt aber auch bei der anderen Firma an, auch wenn diese ihre Produktion nicht erhöht. Wenn z. B.  $D_1$  von der kleinen auf die große Menge übergeht, steigt sein Gewinn von 3 auf 4, weil er die Einbuße durch eine Ausweitung seines Marktanteils überkompensieren kann, während der Gewinn von  $D_2$  auf 1 sinkt. Reagiert darauf  $D_2$  seinerseits mit einer Produktionssteigerung, kann er trotz weiterer Verringerung des Gesamtgewinns (von 5 auf 4) seinen Gewinn durch Ausweitung seines Marktanteils steigern (von 1 auf 2), schmälert dadurch aber auch den Gewinn von  $D_2$  (von 4 auf 2).

die negativen Externalitäten zwischen den Duopolisten nicht, wie beim Umweltschutz und anderen Gemeingütern, über natürliche oder andere Kausalprozesse vermittelt (sogenannte **technologische Externalitäten**), sondern über den Marktpreis (sogenannte **pekuniäre Externalitäten**). Und zweitens ist das Scheitern von Kollusion zwar für die Duopolisten nachteilig, für die gesellschaftliche Wohlfahrt aber vorteilhaft.<sup>34</sup> Das Kartell-Dilemma ist deshalb ein sozial erwünschtes Dilemma.<sup>35</sup>

### b. Kollektivgüter

Ein Spezialfall des Externalitätenproblems und eine Form des **Marktversagens** sind Kollektivgüter, d. h. Güter, die keine privaten Güter sind. **Private Güter** zeichnen sich durch Ausschließbarkeit und Rivalität des Gebrauchs aus. Dabei meint **Ausschließbarkeit**, dass Dritte vom Gebrauch des Gutes ausgeschlossen werden können, diese das Gut also nicht ohne die Mitwirkung des Eigentümers oder sonst Berechtigten nutzen können (ohne Wohnungsschlüssel kann man keine Mietwohnung bewohnen). Ausschließbarkeit gibt dem Berechtigten die Möglichkeit, für die Nutzung des Gutes durch Dritte ein Entgelt zu verlangen, und damit einen Anreiz zur Bereitstellung oder Produktion des Gutes. (**Konsum-**)**Rivalität** meint, dass der Gebrauch eines Gutes durch einen Nutzer den Gebrauch des gleichen Gutes durch andere Nutzer beeinträchtigt (man kann ein Brot nur einmal essen). Weist ein Gut beide Eigenschaften auf, kommt es weder zu positiven noch zu negativen Externalitäten: Der Produzent erhöht zwar den Nutzen der Konsumenten, erhält dafür aber ein Entgelt und berücksichtigt dadurch die positiven Wirkungen seiner Produktion. Der Konsument verringert zwar durch den Verbrauch des Gutes den Nutzen des Produzenten, muss dafür aber ein Entgelt bezahlen und internalisiert damit die negativen Wirkungen seines Konsums.

**Kollektive Güter** zeichnen sich nun dadurch aus, dass ihnen eine oder beide dieser Eigenschaften privater Güter fehlen (und Internalisierung dadurch zum Problem wird). Man unterscheidet insoweit zwischen (reinen) öffentlichen Gütern, Klubgütern (bzw. Mautgütern oder ausschließbaren öffentlichen Gütern) und Gemeingütern (oder Almendegütern). Bei ihnen fehlen die Eigenschaften der Ausschließbarkeit und/oder Rivalität in folgender Kombination:

<sup>34</sup> Vgl. hierzu bereits Rz. 154 ff.

<sup>35</sup> Vgl. Engel, Wettbewerb als sozial erwünschtes Dilemma, in: ders./Möschel (Hg.), Recht und spontane Ordnung, Festschrift für Ernst-Joachim Mestmäcker zum 80. Geburtstag, 2006, 155 ff.

|                   |      | Rivalität     |                   |
|-------------------|------|---------------|-------------------|
|                   |      | ja            | nein              |
| Ausschließbarkeit | ja   | private Güter | Klubgüter         |
|                   | nein | Gemeingüter   | öffentliche Güter |

Tabelle 4.18: Private und kollektive Güter

- 223 **Fehlende Ausschließbarkeit** hat insoweit zur Folge, dass der Nutzer des Gutes die Kosten seiner Nutzung nicht berücksichtigen muss. Bei Gemeingütern, deren Gebrauch rivalisiert, bewirkt dies einen Anreiz zur Übernutzung. Dieses Problem stellt sich häufig bei der gemeinschaftlichen Nutzung natürlicher Ressourcen. Die Ausbeutung von Fischbeständen in internationalen Gewässern ist dafür ein Beispiel. Werden diese überfischt, kann sich der Fischbestand nicht mehr ausreichend erneuern, so dass der Ertrag insgesamt sinkt. Das ist mehr als ein Verteilungsproblem, weil durch die Überfischung die soziale Wohlfahrt insgesamt geschmälert wird.
- 224 Bei (reinen) **öffentlichen Gütern** und **Klubgütern** (ausschließbaren öffentlichen Gütern) ist Übernutzung kein Problem, weil der Gebrauch dieser Güter mangels Rivalität keine Kosten hat. Das Problem liegt vielmehr in der Produktion. Unter dem Gesichtspunkt der Effizienz sollten solche Güter bereitgestellt werden, soweit die Kosten der Bereitstellung den Nutzen ihres Gebrauchs nicht übersteigen. Bei (reinen) öffentlichen Gütern besteht dazu aber kein Anreiz, weil die Kosten der Bereitstellung infolge fehlender Ausschließbarkeit nicht auf die Nutzer überwälzt werden können. Öffentliche Güter werden deshalb häufig vom Staat bereitgestellt werden, der sie über Steuern finanzieren kann. Die Gewährleistung äußerer und innerer Sicherheit durch Bundeswehr und Polizei sind Beispiele für öffentliche Güter in diesem Sinne.<sup>36</sup>
- 225 Bei Klubgütern ist durch die Ausschließbarkeit zwar eine Finanzierung der Bereitstellung über Entgelte oder Gebühren möglich. Diese kann aber ineffizient sein, weil sie auch Nutzungen abschreckt, die mangels Rivalität keine Kosten verursachen. Mögliche positive Externalitäten, die keine Kosten verursachen, werden dann nicht realisiert. So kann man öffentliche Straßen zwar über eine Maut finanzieren, sie werden dann aber weniger genutzt.

<sup>36</sup> Das gilt allerdings dort nicht mehr, wo Sicherheitsdienstleistungen auf einzelne Personen oder Objekte beschränkt werden können, so dass Ausschließbarkeit gegeben ist. Ob es in diesen Konstellationen allerdings ökonomisch oder politisch sinnvoll ist, Sicherheitsleistungen privat anzubieten, ist damit noch nicht gesagt.

Öffentliche Güter und Gemeingüter werden häufig als Gefangenen-Dilemmata modelliert. Bei ihnen ist es aber nicht immer so, dass spieltheoretisch formuliert Defektion eine dominante Strategie ist, also jeder immer einen Anreiz hat, sich gegen das gemeinsame Wohl zu verhalten. Bei der Ausbeutung von Fischbeständen zum Beispiel gibt es häufig einen Punkt, an dem es sich auch für das einzelne Fischereiunternehmen nicht mehr lohnt, das Gemeingut noch stärker auszubeuten, weil der geringe zusätzliche Ertrag die höheren Ausbeutungskosten nicht mehr deckt.<sup>37</sup> Ab diesem Punkt können dann individuelle und kollektive Interessen wieder parallel laufen, wie es das folgende Kooperations-Spiel modelliert (das kein Gefangenen-Dilemma ist).

226

|                |          | Staat B fischt |             |          |
|----------------|----------|----------------|-------------|----------|
|                |          | wenig          | mittel      | intensiv |
| Staat A fischt | wenig    | ↓ 3, 3 →       | ↓ 1, 4      | ↓ 0, ← 3 |
|                | mittel   | 4, 1 →         | <u>2, 2</u> | 1, ← 1   |
|                | intensiv | ↑ 3, 0 →       | ↑ 1, 1      | ↑ 0, ← 0 |

Tabelle 4.19: Beispiel für Nicht-Dominanz von Defektion

Die Auszahlungen in den vier Zellen oben links (mit den Aktionen „wenig“ und „mittel“) entsprechen hier den Auszahlungen eines normalen Gefangenen-Dilemmas. Trotz der Erweiterung um eine weitere Aktion (intensiv) ist das Strategieprofil (mittel, mittel) auch in diesem Spiel ein Gleichgewicht in dominanten Strategien. Denn die Spieler haben nach den Auszahlungen für die Aktion „intensiv“ keinen Anreiz, die Bestände noch weiter auszubeuten. Entsprechend ist das sozial noch nachteiligere Strategieprofil (intensiv, intensiv) kein Gleichgewicht. Allgemein gesprochen kann es Kooperationsprobleme auch ohne dominante Strategien und auch bei mehreren Gleichgewichten geben, wenn nämlich der sozial wünschenswerte Zustand selbst kein Gleichgewicht ist.<sup>38</sup> Das Gefangenen-Dilemma modelliert insoweit nur den **Extremfall** eines Kooperationsproblems.

227

### c. Kooperation mit Konflikt

Während das Gefangenen-Dilemma Kooperationsprobleme überzeichnen kann, weil in ihm Defektion eine dominante Strategie ist, bildet es einen

228

<sup>37</sup> Für ein ausführliches Beispiel vgl. *Beckenkamp*, Preprint 2002/11, MPI for Collective Goods, 2 ff.

<sup>38</sup> Die Wertung als sozial wünschenswert bezieht sich dabei nur auf die Interessen der Spieler. Für Dritte kann Kooperation durchaus negative Folgen haben, wie im Beispiel des Kartell-Dilemmas.



anderen Aspekt von Kooperationsproblemen gar nicht ab, der aber in der Realität eine große Rolle spielt. Gemeint ist der Umstand, dass es oft mehrere Möglichkeiten der Kooperation gibt, hinsichtlich derer die Akteure unterschiedliche Präferenzen haben. Dieses Problem tritt zum Beispiel auf, wenn Kooperationsgewinne oder Kooperationslasten unterschiedlich verteilt werden können. Zusätzlich zu dem Kooperationsproblem müssen die Akteure dann auch noch einen Verteilungskonflikt bewältigen. Dieses Problem modellieren gemischte Kooperations- und Konfliktspiele. Nehmen wir an, zwei Firmen leiten Abwässer in ein Gewässer ein, welche die eigene Produktion und die der anderen Firma beeinträchtigen und dadurch Kosten verursachen. Zur Reinigung der Abwässer könnten die Firmen einen von zwei Filtern einbauen oder untätig bleiben. Aufgrund der unterschiedlichen Einbaukosten und den differenzierten Folgen für die Firmen A und B ergäben sich folgende Auszahlungen:

|         |             | Firma B         |                  |                     |
|---------|-------------|-----------------|------------------|---------------------|
|         |             | Filter 1        | Filter 2         | kein Filter         |
| Firma A | Filter X    | ↓ 5, 5 →        | ↓ <u>4</u> , 8 → | ↓ 0, 9              |
|         | Filter Y    | ↓ 8, <u>4</u> → | ↓ 5, 5 →         | ↓ 1, 6              |
|         | kein Filter | 9, 0 →          | 6, 1 →           | <u>2</u> , <u>2</u> |
|         |             |                 |                  |                     |

Tabelle 4.20: Beispiel für ein gemischtes Kooperations- und Konfliktspiel

- 229 Dieses Kooperationsproblem hat ein Gleichgewicht in dominanten Strategien, bei dem keine der Firmen (freiwillig) einen Filter einbaut (durch Unterstreichungen gekennzeichnet). Dieses Gleichgewicht ist *pareto*-inferior gegenüber sechs *pareto*-optimalen Spielausgängen, von denen vier als Gegenstand möglicher Kooperationsbemühungen in Betracht kommen (kursiv gekennzeichnet).<sup>39</sup> Von diesen sind zwei mit Auszahlungen von (5, 5) symmetrisch. Gegenüber diesen sind die anderen beiden *pareto*-optimalen Gleichgewichte mit Auszahlungen von (8, 4) und (4, 8) effizienter im Sinne des *Kaldor-Hicks*-Kriteriums (Gesamtauszahlungen von 12 gegenüber 10), aber stark ungleich verteilt. Wie dieses, vielen Kooperationssituationen immanente, Verteilungsproblem zu lösen ist, sagt die Spieltheorie nicht. Seine Existenz stellt aber für das praktische Gelingen von Kooperation ein zusätzliches und nicht selten zentrales Hindernis dar, und eine der wesentlichen Funktionen des Rechts könnte darin bestehen, dieses Kooperationshindernis zu überwinden.<sup>40</sup> Der Klimaschutz gibt dafür ein aktu-

<sup>39</sup> Für die Spielausgänge mit den Auszahlungen 9, 0 und 0, 9 gilt dies nicht, weil sie gegenüber dem Gleichgewicht eine *Pareto*-Verschlechterung darstellen.

<sup>40</sup> So Magen, Gerechtigkeit als Proprium des Rechts, 2010.

elles Beispiel von globalem Ausmaß.<sup>41</sup> Auch wenn weitgehend Einigkeit besteht, dass eine Begrenzung der CO<sub>2</sub>-Emissionen in bestimmtem Umfang im Wohl jedenfalls sehr vieler Staaten liegt, kommen die Verhandlungen für ein Nachfolgeabkommen zum Kyoto-Protokoll zur Zeit nicht voran, weil sich die Staaten nicht über die Verteilung der Reduktionsverpflichtungen und damit der Klimaschutzkosten einigen können. Auch bei der Einführung eines europäischen Emissionshandelssystems war und ist die Verteilung der Reduktionslasten einer der politischen Hauptstreitpunkte, der z.B. auf nationaler Ebene durch die Zuteilungsgesetze (ZuG 2007 und ZuG 2012) verbindlich entschieden werden musste.

#### d. Kooperation im weiteren Sinn

Kooperation im eigentlichen Sinn liegt nur vor, wenn Kooperation eine *Pareto*-Verbesserung bewirkt, also durch Kooperation niemand im Vergleich zu dem unkooperativen Gleichgewicht schlechter gestellt wird. Eine lange Diskussion über die ökonomische Analyse des Rechts hat aber gezeigt, dass diese Voraussetzung selten erfüllt ist.<sup>42</sup> Sobald man auf konkrete staatliche Vorhaben sieht, gibt es allzu oft einige Betroffene, die nicht gewinnen, sondern unter dem Strich verlieren. Wer etwa in der Nähe der Flugschneise einer neuen Startbahn wohnt, wird ungeachtet der wirtschaftlichen Vorteile für die Region auf deren Bau möglicherweise lieber ganz verzichten wollen.<sup>43</sup> In solchen Konstellationen, in denen Kooperation auch Netto-Verlierer hinterlässt, kann man dennoch von Kooperation im weiteren Sinn sprechen, wenn das Kooperationsziel, wie in folgendem Beispiel, zumindest im Sinne des *Kaldor-Hicks*-Kriteriums noch als sozial erwünscht beschrieben werden kann oder wenn das Kooperationsziel aufgrund einer politisch legitimen Entscheidung als Gemeinwohlziel festgelegt worden ist.

230

<sup>41</sup> Ebd., S. 456 ff.

<sup>42</sup> *Hellwig*, Effizienz oder Wettbewerbsfreiheit? Zur normativen Grundlegung der Wettbewerbspolitik, in: Engel/Möschel (Hg.), Recht und spontane Ordnung, Festschrift für Ernst-Joachim Mestmäcker zum 80. Geburtstag, 2006, 231 (238); *Posner*, Economic Analysis of Law, 8. Aufl. 2011, S. 18; vgl. aber andererseits *Kirchgässner*, Gemeinwohl in der Spannung von Wirtschaft und politischer Organisation: Bemerkungen aus ökonomischer Perspektive, in: Brugger *et al.* (Hg.), Gemeinwohl in Deutschland, Europa und der Welt, 2002, 289 (317); das *Pareto*-Kriterium decke vielleicht den wichtigsten Teil staatlicher Aktivitäten ab.

<sup>43</sup> Vgl. etwa die Klage gegen die Flugplatzenerweiterung beim Airbus-Werk in Hamburg; OVG Hamburg, NVwZ-RR 2006, 97 ff.

|   |            | B          |             |
|---|------------|------------|-------------|
|   |            | kooperiert | defektiert  |
| A | kooperiert | ↓ 6, 6 →   | ↓ 1, 8      |
|   | defektiert | 8, 1 →     | <u>7, 2</u> |

Tabelle 4.21: Beispiel für ein Kooperationsspiel im weiteren Sinn

#### 4. Wiederholte Spiele

- 231 Bislang war die stillschweigende Annahme, dass es sich bei den in den Matrizen dargestellten Spielen um Einmal-Spiele handelt (*one shot games*). Manche Interaktionen, etwa zwischen Nachbarn, in laufenden Geschäftsbeziehungen oder in Arbeitsverhältnissen, zwischen Parteien und Fraktionen oder zwischen Staaten dauern aber über einen längeren Zeitraum an. Die Spieltheorie modelliert solche Situationen als wiederholte Spiele (*repeated games*). Dabei wird angenommen, dass die Spieler das in der Matrix dargestellte und in diesem Zusammenhang Basisspiel (*stage game*) genannte Spiel über mehrere Runden spielen, wodurch ein neues Spiel, das Gesamtspiel oder *super game*, entsteht. Nun sind zwar alle Gleichgewichte des Basisspiels als wiederholte Aktionen auch Gleichgewichte des Gesamtspiels. Das wiederholte Spiel kann aber eine Vielzahl zusätzlicher Gleichgewichte aufweisen. Ob das der Fall ist, hängt wesentlich davon ab, ob das Spiel unendlich oder unbestimmt oft wiederholt wird oder ob es für eine im Voraus bestimmte Zahl von Runden gespielt wird.

##### a. Unbestimmt oft wiederholte Spiele und Folk-Theorem

- 232 Betrachten wir zunächst den Fall, dass die Spieler nicht wissen, in welcher Runde das Spiel endet, weil dieses unbestimmt oft wiederholt wird. Von den Bewohnern zweier benachbarter Reihenhäuser hört zum Beispiel vielleicht der eine gerne laute Musik, was seinen Nachbar belästigt, während der Nachbar gerne geräuschintensive handwerkliche Arbeiten verrichtet, was wiederum ersteren stört.<sup>44</sup> Beide ziehen es vor, dass der Nachbar Rücksicht nimmt, sie aber ihren Vorlieben nachgehen können.<sup>45</sup> Eine solche Situation kann ein Gefangenens-Dilemma sein, so dass der Konflikt eigentlich unvermeidbar ausbrechen sollte.<sup>46</sup> Allerdings müssen die Nach-

<sup>44</sup> Zu zivilrechtlichen Abwehransprüchen in derartigen Konstellationen vgl. OLG München, NJW-RR 1991, 1492 ff.

<sup>45</sup> Wenn man soziale Präferenzen zunächst außer Betracht lässt.

<sup>46</sup> Vorausgesetzt, die Nachbarn finden beiderseitigen Lärm immer erträglicher als

barn auf Dauer miteinander auskommen, so dass sie ihre Entscheidungen im **Schatten der Zukunft** treffen.<sup>47</sup> Spieltheoretisch betrachtet kommt es nicht mehr nur auf die Auszahlungen im Basisspiel, sondern auf die addierten (und diskontierten) Auszahlungen des Gesamtspiels an. Deshalb kann der Zeithorizont die Natur des Spiels ändern, vorausgesetzt, die Wahrscheinlichkeit der Wiederholung des Basisspiels ist groß genug und die Beteiligten werten zukünftige Vorteile nicht zu stark ab, wenn also die sogenannten Diskontraten nicht zu hoch sind. In einem wiederholten Gefangenen-Dilemma ist dann zwar immer noch beiderseitige Defektion ein Gleichgewicht (wenn der eine nie Rücksicht nimmt, ist Rücksichtnahme auch für den anderen nicht ratsam, und *vice versa*). Im langen Zeithorizont verschwindet aber die Unausweichlichkeit der Dilemmastruktur. Kooperiert nämlich der eine, aber nur solange, wie auch der andere kooperiert, dann liegt es im eigenen Interesse des anderen, die Kooperation fortzusetzen, weil er andernfalls die langfristigen Vorteile aus einer Fortsetzung der Kooperation verlieren würde, während der Gewinn aus einer einseitigen Defektion nur kurzfristig ist, weil sein Gegenüber dann ebenfalls zu unkooperativem Verhalten übergehen würde. Anders gesagt haben die Spieler im langen Zeithorizont eine glaubhafte Sanktionsdrohung. In einem unbestimmt oft wiederholten Gefangenen-Dilemma ist deshalb auch gegenseitige Kooperation ein strategisches Gleichgewicht. Aus einem Spiel mit einem dominanten Gleichgewicht wird dann ein Koordinierungsproblem, bei dem es zwischen verschiedenen Gleichgewichten zu wählen gilt. Das Gesamtspiel eines wiederholten Gefangen-Dilemmas hat dann die Struktur einer Hirschjagd,<sup>48</sup> bei dem man (u.a.) zwischen einem Gleichgewicht mit geringen, aber sicheren Auszahlungen, und einem Gleichgewicht mit hohen, aber riskanten Auszahlungen wählen muss.<sup>49</sup>

Allgemeiner gesagt lassen sich in unbestimmt oft wiederholte Spielen nach dem sogenannten **Folk-Theorem** (bei hinreichend hoher Wiederholungswahrscheinlichkeit und hinreichend niedrigen Diskontraten) alle Auszahlungen realisieren, in denen jeder Spieler im Durchschnitt wenig-

233

---

eigenes einseitiges Nachgeben. Ist dies nicht der Fall, etwa weil sie besonders lärmempfindlich sind, handelte es sich um ein Falke-Taube-Spiel.

<sup>47</sup> Bild von *Axelrod*, *The Evolution of Cooperation*, 1984, S. 174.

<sup>48</sup> Hierzu oben in Rz. 215 f.

<sup>49</sup> *Skyrms*, *The Stag Hunt and the Evolution of Social Structure*, 2004, S. 5 ff. Etwas genauer gesagt, ist das nicht-kooperative Gleichgewicht gegenüber dem kooperativen risiko-dominant, wenn die Wahrscheinlichkeit, dass das Spiel wiederholt wird, in einem bestimmten mittleren Bereich liegt. Liegt sie darunter, ist auch das wiederholte Spiel ein Gefangenen-Dilemma, liegt sie darüber, ist das kooperative Gleichgewicht risiko-dominant.

stens das erhält, was ihm die anderen Spieler im Basisspiel nicht einseitig nehmen könnten (z.B. die Auszahlungen für beiderseitige Defektion im Gefangenenden-Dilemma). Der Grund ist, dass die Spieler ihre Aktionen auch als Sanktionen gegen den Mitspieler nutzen können, indem sie ihre Aktionen so wählen, dass die Auszahlungen des Mitspielers möglichst niedrig sind (sogenannten *Maximin*-Auszahlung).<sup>50</sup> Die Spieler können deshalb implizit oder explizit vereinbaren, bestimmte Aktionen des Basisspiels zu spielen und Abweichungen von diesem Plan durch die Maximin-Auszahlungen zu bestrafen. Ein solcher sanktionsbewehrter Plan ist für alle möglichen Strategien oberhalb der Maximin-Auszahlungen möglich, so dass wiederholte Spiele unendlich viele Gleichgewichte haben.<sup>51</sup> Sie werfen damit vor allem ein Koordinationsproblem auf.

### b. Endlich wiederholte Spiele und Rückwärts-Induktion

- 234 Wird das Basisspiel endlich oft wiederholt, haben jedenfalls Spiele, die als Einmal-Spiel nur ein strategisches Gleichgewicht haben, auch als wiederholte Spiele nur dieses eine (teilspielperfekte)<sup>52</sup> Gleichgewicht. Das ergibt sich aus der Methode der sogenannten **Rückwärtsinduktion** (*backward induction*), bei der das Gesamtspiel vom letzten Basisspiel aus gelöst wird: In der letzten Runde sind die Anreize die gleichen wie im Einmalspiel, weshalb das Gleichgewicht des Einmalspiels gespielt werden wird. Das gleiche gilt dann auch für die vorletzte Runde, die vorvorletzte, usw. Deshalb ist Kooperation im endlich oft wiederholten Gefangenenden-Dilemma kein teilspielperfektes Gleichgewicht.<sup>53</sup>

<sup>50</sup> Vgl. *Holler/illing/Napel*, Einführung in die Spieltheorie, 2016, S.148. In dem Gefangenenden-Dilemma im Text ist die Maximin-Auszahlung 2, weil sich die Spieler unabhängig davon, wie sich der andere verhält, zumindest eine Auszahlung von 2 sichern können. Nähert sich die Diskontrate 0 und die Wiederholungswahrscheinlichkeit 1, dann können in dem Gesamtspiel in strategischen Gleichgewichten alle Auszahlungen realisiert werden, in denen die Spieler im Durchschnitt eine beliebige Auszahlung zwischen 2 und 4 bekommen.

<sup>51</sup> Vorausgesetzt ist freilich, dass die Spielzüge aller Spieler beobachtbar sind.

<sup>52</sup> Zur Teilspielperfektion noch sogleich, Rz. 239f.

<sup>53</sup> Die Rückwärtsinduktion ist zwar logisch zwingend, widerspricht aber häufig der Intuition und dem Verhalten der Menschen. Diesen Widerspruch bringt das sog. „Handelsketten-Paradox“ von *Reinhard Selten* zum Ausdruck, vgl. *Selten*, The chain store paradox, *Theory and Decision* 9 (1978), 127 (133).

## IV. Spiele in Extensivform

## 1. Definition eines Spiels in Extensivform

Neben der oben eingeführten Normalform gibt es noch eine weitere Form 235  
 der Darstellung von Spielen, nämlich die **Extensivform**. In der Extensiv-  
 form wird anhand eines Spielbaums explizit angegeben, in welcher Rei-  
 henfolge die Spieler ziehen (Zugfolge) und welche Informationen sie bei  
 jeder Entscheidung über den bisherigen Spielverlauf haben (Informations-  
 stand). Ein Spielbaum besteht aus **Knoten** und aus **Kanten**. Knoten be-  
 zeichnen Entscheidungssituationen (**Entscheidungsknoten**) und Endsitu-  
 ationen (**Endknoten**). Kanten bezeichnen die **Aktionen**, zwischen denen ein  
 Spieler an einem Entscheidungsknoten wählen kann. Sie führen zu Ent-  
 scheidungsknoten, wenn noch eine Entscheidung getroffen werden kann,  
 oder zu Endknoten, wenn das Spiel beendet ist. Die Endknoten bezeichnen  
 mögliche Spielausgänge und die mit ihnen verbundenen Auszahlungen.

Zur Veranschaulichung können wir wieder die Frage der Standardset- 236  
 zung heranziehen, nur dass wir jetzt annehmen, die Firmen A und B  
 könnten nicht gleichzeitig entscheiden, sondern zuerst Firma A und dann,  
 in Kenntnis der Entscheidung von A, Firma B. Das Spiel ist jetzt ein **se-**  
**quentielles Spiel mit perfekter Information**. In Extensivform stellt es sich  
 wie folgt dar:

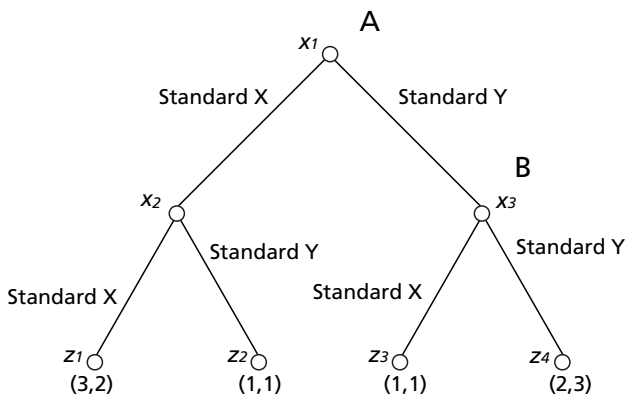


Abbildung 4.1: Spielbaum des sequentiellen Standardisierungs-Spiels

Der erste Entscheidungsknoten ( $x_1$ ) markiert die Zugmöglichkeit für A, 237  
 die von ihm wegführenden Kanten seine Handlungsmöglichkeiten (Stan-  
 dard X und Standard Y). Je nachdem, wie sich A entscheidet, befindet sich

B in der zweiten Ebene am linken ( $x_2$ ) oder am rechten ( $x_3$ ) Entscheidungsknoten. Auch er hat die gleichen zwei Handlungsmöglichkeiten, und je nachdem, wie er sich entscheidet, endet das Spiel an einem der vier Endknoten ( $z_1, z_2, z_3, z_4$ ).

- 238 Die Lösung dieses Spiels ist intuitiv relativ einleuchtend, wenn man sich zunächst die Situation von Firma B verdeutlicht, wenn sie am Zug ist. Firma A hat sich zu diesem Zeitpunkt bereits entschieden. B weiß also, ob sie sich an  $x_2$  oder an  $x_3$  befindet, und wird die für sie vorteilhaftere Handlung wählen. Wenn A sich für Standard X entschieden hat, wird sie ebenfalls Standard X wählen, und wenn A Standard Y gewählt hat, wird B ebenfalls diesen Standard wählen. Firma A antizipiert dies und kann sich folglich aussuchen, welchen Standard sie wählt. Sie wird sich für den für sie günstigeren Standard X entscheiden. Sie hat hier den für **sequentielle Koordinationsspiele** typischen sogenannten *first mover advantage*.

## 2. Teilspielperfektion

- 239 Spieltheoretisch liegen die Dinge allerdings nicht so einfach. Um die *Nash-Gleichgewichte* zu ermitteln, kann man den Spielbaum in eine Matrix umformulieren. Während Firma A insoweit die gleichen Aktionen/Strategien hat wie im simultanen Spiel, wird bei Firma B der Unterschied von Strategien und Aktionen relevant. Da **Strategien** nämlich *ex ante* formulierte **vollständige Handlungspläne** sind, muss B in jeder Strategie für jeden seiner Entscheidungsknoten ( $x_2, x_3$ ) eine Aktion auswählen.<sup>54</sup> Seine Strategien enthalten deshalb jeweils zwei bedingte Aktionen, etwa: „Wenn A Standard X wählt, wähle ich auch Standard X; wenn A Standard Y wählt, wähle ich auch Standard Y“, oder kürzer: ( $x_2$ : X/ $x_3$ : Y) bzw. noch kürzer (X/Y). Abhängig vom Verhalten des A wird aber von diesen beiden konditionell spezifizierten Aktionen nur eine tatsächlich gespielt. In dem sequentiellen Koordinierungsspiel ergeben sich für Firma B damit die folgenden vier möglichen Strategien: immer Standard X (X/X); immer Standard Y (Y/Y); den gleichen Standard wie Firma A (X/Y); den entgegengesetzten Standard wie Firma A (Y/X). Das Spiel lässt sich dann in folgender Matrix darstellen:<sup>55</sup>

<sup>54</sup> Genauer: für jede seiner Informationsmengen.

<sup>55</sup> Die Auszahlungen ergeben sich daraus, welche Aktion B ausführt. Wenn A links wählt, also die obere Zeile spielt, wählt B immer die linke der in seiner Strategie spezifizierten Aktionen; wenn A die untere Zeile spielt, also rechts wählt, die rechte Aktion.

|         |   | B wählt                       |                 |                                 |                                |
|---------|---|-------------------------------|-----------------|---------------------------------|--------------------------------|
|         |   | X/X                           | X/Y             | Y/X                             | Y/Y                            |
| A wählt | X | <u>3, 2</u>                   | <u>3, 2</u>     | 1, $\leftarrow$ 1               | $\downarrow$ 1, $\leftarrow$ 1 |
|         | Y | $\uparrow$ 1, 1 $\rightarrow$ | $\uparrow$ 2, 3 | 1, $\leftarrow$ 1 $\rightarrow$ | <u>2, 3</u>                    |

Tabelle 4.22: Spielmatrix des sequentiellen Standardisierungs-Spiels

Man sieht nun, dass das Spiel drei *Nash*-Gleichgewichte hat. Das fett markierte entspricht dabei unserem intuitiven Ergebnis: B wählt die Strategie, der Entscheidung des A zu folgen, A kann damit das Gleichgewicht festlegen und entscheidet sich für das ihm vorteilhaftere Gleichgewicht mit seinem Standard X. Daneben gibt es aber noch zwei weitere *Nash*-Gleichgewichte, die hier kursiv markiert sind. Sie verdanken ihre Gleichgewichtseigenschaft einer impliziten Drohung des B. Interessant ist insoweit insbesondere das Gleichgewicht in der Zelle ganz rechts unten (Y, Y/Y). Die Strategie von B lautet hier, den eigenen Standard Y zu wählen, unabhängig davon, wie sich A verhält.

240

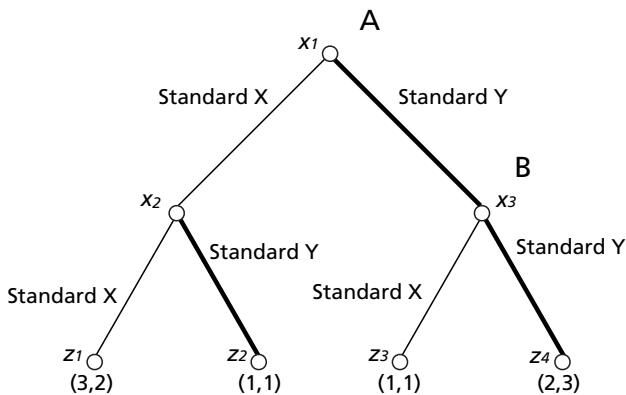


Abbildung 4.2: Gleichgewichtsstrategien des Gleichgewichts (Y, Y/Y)

In diesem Gleichgewicht gibt es zwar für Firma A keinen Anreiz, den für sie günstigen Standard X zu wählen, aber nur weil sich Firma B auch für diesen Fall vorher darauf festgelegt hat, ihren Standard Y zu wählen (an  $x_2$ ). Deshalb erhielte A, wenn sie sich für ihren Standard entschied, nur die schlechte Auszahlung von (1,1) am Endknoten  $z_2$ . Das Problem ist allerdings, dass es für B selbst nachteilig wäre, seine Drohung auch auszuführen, denn auch für Firma B sind die Auszahlung am Endknoten  $z_2$  schlechter als an  $z_1$ . Die Drohung des B ist mithin nicht glaubwürdig.

241



242 Die Frage ist, ob Gleichgewichte, die auf unglaublichen Drohungen beruhen, eine überzeugende Lösung des Spiels darstellen. Eine Antwort darauf gibt das von *Reinhard Selten* entwickelte Konzept der **Teilspielperfektion**, das eine Verfeinerung des *Nash*-Gleichgewichts darstellt.<sup>56</sup> Teilspiele sind Spiele, die nicht am ersten Entscheidungsknoten beginnen, sondern an einem nachfolgenden. In diesem Beispiel gibt es also zwei Teilspiele (beginnend mit  $x_2$  und  $x_3$ ), von denen aber, je nach der vorangehenden Entscheidung des A, nur eins gespielt werden kann. In dem problematischen *Nash*-Gleichgewicht (Y; Y/Y) würden nur die Aktionen  $x_1$ : Standard Y und  $x_3$ : Standard Y gespielt. Auf diesem sogenannten **Gleichgewichtspfad** liegt nur das Teilspiel  $x_3$ . Die Gleichgewichtsstrategien spezifizieren zwar auch eine Entscheidung der B für das Teilspiel  $x_2$ . Eben diese Entscheidung enthält die Drohung, die dafür sorgt, dass A lieber Standard Y wählt. Aber diese Entscheidung wird im Gleichgewicht nicht gespielt, sie liegt jenseits des Gleichgewichtspfads. Das Konzept der Teilspielperfektion verlangt nun, dass die Gleichgewichtsstrategien in jedem Teilspiel ein *Nash*-Gleichgewicht darstellen müssen, also auch in solchen, die jenseits des Gleichgewichtspfades liegen. Das ist aber für das *Nash*-Gleichgewicht (Y; Y/Y) nicht der Fall, und Gleiches gilt für das *Nash*-Gleichgewicht (X; X/X). Nur das Gleichgewicht, in dem A den für ihn vorteilhafteren Standard X wählt und B dieser Entscheidung folgt, ist teilspielperfekt.

243 Anhand des Spielbaums lässt sich diese Lösung auch mit der Methode der sogenannten Rückwärtsinduktion (*backward induction*) ermitteln. Bei der Rückwärtsinduktion löst man das Spiel für jedes Teilspiel rückwärts von den letzten Entscheidungsknoten ( $x_2$  und  $x_3$ ) bis zum Anfangsknoten ( $x_1$ ). An  $x_2$  würde B Standard X wählen, an  $x_3$  Standard Y. Weil A dies antizipiert, kann er für seine Entscheidung die Kanten zu den Endknoten abschneiden, die B nicht wählen würde (sogenannten *tree pruning*). Der Spielbaum verkürzt sich für ihn damit wie folgt:

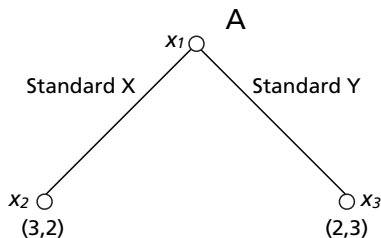


Abbildung 4.3: Beispiel für *tree pruning*

<sup>56</sup> *Selten*, Spieltheoretische Behandlung eines Oligopolmodells mit Nachfrageträgheit, Zeitschrift für die gesamte Staatswissenschaft 121 (1965), 301 (306 ff.).

### 3. Imperfekte Informationen und Informationsmengen

Bislang haben wir angenommen, Firma B kenne die Entscheidung, die Firma A zuvor getroffen hat. Anders gesagt, B weiß, ob sie sich am Entscheidungsknoten  $x_2$  oder am Knoten  $x_3$  befindet. Kennt ein Spieler alle vorangegangenen Züge seiner Mitspieler, spricht man von perfekter (oder vollkommener) Information. Man kann sich aber auch Situationen vorstellen, in denen der eine Spieler zwar früher zieht, der andere diese Entscheidung bei seiner Entscheidung aber (noch) nicht kennt. Man spricht dann von **imperfekter** (oder unvollkommener) **Information**. Bezogen auf den Spielbaum hieße das in unserem Standardsetzungs-Spiel, dass B bei seiner Entscheidung nicht weiß, ob er sich am Knoten  $x_2$  oder an  $x_3$  befindet. Die **Entscheidungsknoten**, zwischen denen ein Spieler nicht unterscheiden kann, gehören in der Terminologie der Spieltheorie zu einer Informationsmenge. Entgegen der Intuition weiß ein Spieler also umso weniger über die Züge seines Mitspielers, je größer seine Informationsmenge ist. Kennt also B die Entscheidung von A nicht, gehören  $x_2$  und  $x_3$  zur gleichen Informationsmenge. Im Spielbaum wird dies dadurch gekennzeichnet, dass man Entscheidungsknoten, zwischen denen ein Spieler nicht unterscheiden kann, durch eine gestrichelte Linie verbindet.

244

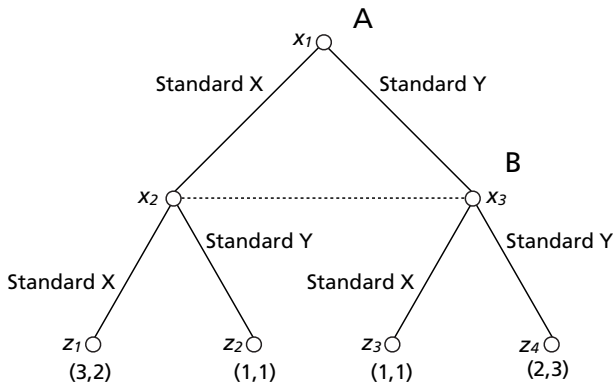


Abbildung 4.4: Beispiel für ein Spiel mit imperfekter Information

Dieses Spiel lässt sich mittels Rückwärtsinduktion nicht mehr lösen, weil Firma B ja nicht weiß, wo sie sich befindet, und deshalb ihre Aktionen nicht mehr von der Entscheidung von Firma A abhängig machen kann.<sup>57</sup> Ihre Strategiemenge verkürzt sich deshalb wiederum auf die Entscheidung

245

<sup>57</sup> Die mit den Entscheidungsknoten  $x_2$  und  $x_3$  beginnenden Teile des Spielbaums

für einen der beiden Standards, also auf ihre beiden Aktionen. Das sequentielle Standardsetzungs-Spiel mit imperfekter Information ist damit äquivalent mit dem simultanen Standardsetzungs-Spiel, das wir zunächst analysiert hatten.<sup>58</sup>

#### 4. Unvollständige Informationen

- 246 Als gemeinsames Wissen (*common knowledge*) bezeichnet man Charakteristika des Spiels oder der Spieler, die jeder Spieler kennt und von denen jeder Spieler weiß, dass sie jeder Spieler kennt, und von denen jeder Spieler weiß, dass jeder Spieler weiß, dass sie jeder Spieler kennt, usw. Solches gemeinsames Wissen wird von den spieltheoretischen Lösungskonzepten in verschiedenen Hinsichten vorausgesetzt. Zum Beispiel setzt das Konzept des *Nash*-Gleichgewichts voraus, dass sich alle **Spieler rational** verhalten und dies auch gemeinsames Wissen ist.
- 247 Von **vollständigen Informationen** spricht man, wenn die Charakteristika aller Spieler gemeinsames Wissen sind, das heißt, wenn alle Spieler die Auszahlungen aller Mitspieler kennen und auch wissen, welche Strategien und Informationen ihnen zur Verfügung stehen (während sich der Begriff der perfekten Informationen nicht auf die Kenntnis der Charakteristika der Spieler, sondern der vorangegangenen Spielzüge bezieht). Viele unter strategischen Gesichtspunkten interessante Situationen zeichnen sich durch unvollständige Informationen in diesem Sinne aus, zum Beispiel, wenn ein Akteur nicht sicher ist, welchen Nutzen ein Kooperationspartner tatsächlich aus der gemeinsamen Kooperation zieht. Man kann solche Probleme **unvollständiger Information** modellieren, indem man annimmt, die „Natur“ wähle die fraglichen Charakteristika der Spieler in einem Zufallszug aus, über den aber Unkenntnis bestehe. So übersetzt man Spiele mit unvollständiger Information in Spiele mit imperfekter Information und kann sie mit Hilfe von Informationsmengen in der oben vorgestellten Weise analysieren.
- 248 Das sei am Beispiel eines **Markteintrittsspiels** dargestellt:<sup>59</sup> Ein Markt wird von einem Monopolisten M beherrscht, der auf dem Markt einen Monopolgewinn erzielt ( $G_M = 100$ ). Ein Konkurrent A überlegt, in den

sind dann auch keine Teilspiele mehr, weil sie über eine Informationsmenge miteinander verbunden sind.

<sup>58</sup> Hierzu bereits Rz. 192 ff.

<sup>59</sup> Das Spiel geht zurück auf *Selten*, The chain store paradox, Theory and Decision 9 (1978), 127 ff.; das Beispiel hier orientiert sich an *Holler/Illing/Napel*, Einführung in die Spieltheorie, 2016, S. 15 f., S. 47 ff.

Markt einzutreten. Tut er dies, kann der Monopolist mit einer aggressiven Abwehrstrategie oder mit Marktteilung reagieren. Welche Auswirkungen ein Markteintritt für den Gewinn von M und A haben würde, hängt nun von der Kostenstruktur des M ab, aufgrund derer sich der Monopolist gegenüber einem Angreifer entweder in einer schwachen oder starken Position befindet. Ist der Monopolist schwach, dann erzielen M und A bei einem Markteintritt des A einen anteiligen Duopolgewinn ( $G_D = 40$ ), der in der Summe allerdings geringer ausfällt als der Monopolgewinn ( $80 < 100$ ). Ein Abwehrkampf würde dagegen für M und A zu Verlusten führen ( $G_A = -10$ ). Ist der Monopolist dagegen in einer starken Position, erzielt er (z.B. wegen hoher Skalenerträge) auch bei einem Angriff noch Gewinne ( $G_{DM} = 30$ ), während der Angreifer bei einem Markteintritt Verluste erleidet ( $G_{DA} = -10$ ). Bei einer Marktteilung dagegen fällt der anteilige Duopolgewinn für M (z.B. wegen sinkender Skalenerträge) geringer aus als sein Gewinn bei einem Abwehrkampf ( $G_{TM} = 20$ ). Ein schwacher Monopolist würde also auf den Angriff mit Marktteilung reagieren, ein starker Monopolist dagegen mit einem Abwehrkampf.

Allerdings kennt A die Kostenstruktur des M nicht. Der Angreifer hat es also mit zwei möglichen Spielen zu tun, von denen er nicht weiß, welches

249

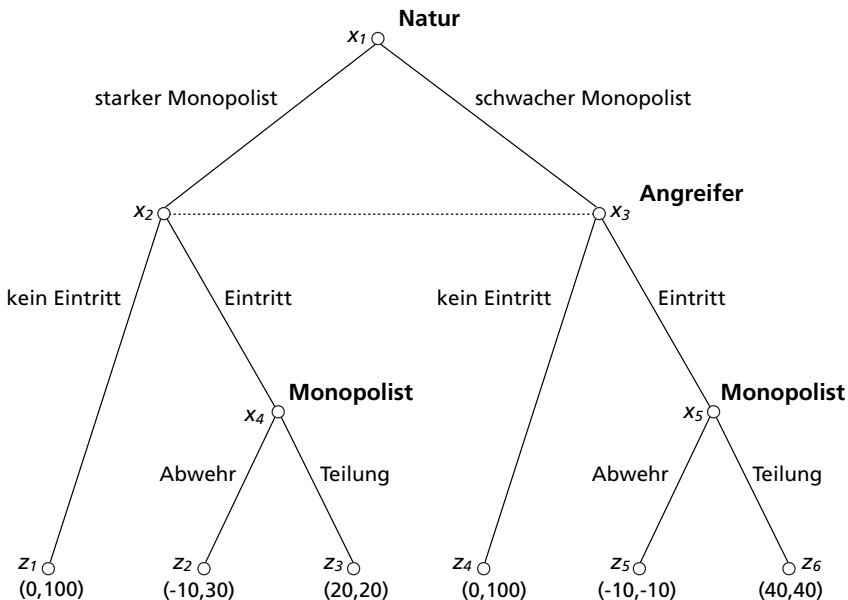


Abbildung 4.5: Beispiel für ein Markteintrittsspiel

gespielt wird. Man kann dies modellieren, indem man zu Beginn des Spielbaums einen Zug „Natur“ vorsieht, bei dem diese mit bestimmten Wahrscheinlichkeiten zwischen beiden möglichen Spielen wählt. Es ergibt sich dann folgendes Spiel, wobei das Teilspiel  $x_2$  den Fall des starken Monopolisten, und Teilspiel  $x_3$  den Fall des schwachen Monopolisten darstellt. Da A nicht weiß, welches Teilspiel er spielt, weil er den Zug der Natur nicht kennt, befinden sich die Knoten  $x_2$  und  $x_3$  in einer Informationsmenge. Der Monopolist dagegen kennt seine Kostenstruktur und weiß deshalb, ob er sich in Spiel  $x_4$  oder in Spiel  $x_5$  befindet, weshalb diese beiden Knoten nicht einer Informationsmenge angehören.

250 Löst man zunächst das Teilspiel  $x_2$  mit der Methode der Rückwärtsinduktion, sieht man, dass sich ein starker Monopolist bei einem Markteintritt (an  $x_4$ ) für den Abwehrkampf mit den Auszahlungen  $(-10/30)$  entscheiden wird. Dies antizipierend würde sich der Angreifer (an  $x_2$ ) gegen einen Eintritt entscheiden.<sup>60</sup> Im Teilspiel  $x_3$  dagegen würde sich der schwache Monopolist bei einem Angriff (an  $x_5$ ) für die Marktteilung entscheiden  $(40/40)$ , so dass der Konkurrent in diesem Fall (an  $x_3$ ) den Markteintritt wählen würde.<sup>61</sup> Allerdings weiß der Angreifer nicht, ob er sich an  $x_2$  oder an  $x_3$  befindet (die derselben Informationsmenge angehören). Wir können das Spiel aber um die Züge kürzen, die M nicht spielen würde:

251 Jetzt sehen wir, dass die Entscheidung des A nicht mehr von strategischen Erwägungen abhängt, sondern nur noch vom Erwartungsnutzen für die beiden Entscheidungsmöglichkeiten des A (Auszahlungen x Ein-

<sup>60</sup> Die Normalform für Teilspiel  $x_2$  sieht wie nachstehend aus. Wegen der sequentiellen Struktur des Spiels hat der Monopolist eigentlich vier bedingte Strategien zur Auswahl, von denen aber jeweils zwei identische Auszahlungen haben und auf eine reduziert werden können (sog. **reduzierte Form**). Das Teilspiel hat nur ein *Nash*-Gleichgewichte, nämlich (kein Eintritt/Abwehr).

|           |               | starker Monopolist |          |
|-----------|---------------|--------------------|----------|
|           |               | Abwehr             | Teilung  |
| Angreifer | kein Eintritt | <u>0, 100</u>      | ↓ 0, 100 |
|           | Eintritt      | ↑ -10, 30          | 20, ← 20 |

<sup>61</sup> Die Normalform für Teilspiel  $x_3$  sieht wie nachstehend aus. Dieses Teilspiel hat ein weiteres *Nash*-Gleichgewicht, nämlich (Eintritt/Teilung), während das Gleichgewicht (kein Eintritt/Abwehr) nicht mehr teilspielperfekt ist.

|           |               | schwacher Monopolist |               |
|-----------|---------------|----------------------|---------------|
|           |               | Abwehr               | Teilung       |
| Angreifer | kein Eintritt | <u>0, 100</u>        | ↓ 0, 100      |
|           | Eintritt      | ↑ -10, -10 →         | <u>40, 40</u> |

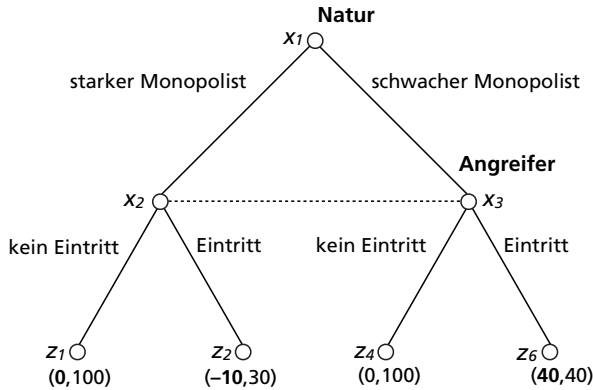


Abbildung 4.6: Markteintrittsspiel nach Rückwärtsinduktion

trittswahrscheinlichkeit).<sup>62</sup> Ist es z. B. gleich wahrscheinlich, dass der Monopolist stark oder schwach ist, beträgt der Erwartungsnutzen bei Markteintritt<sup>63</sup> 15 und ohne Markteintritt 0. Der Konkurrent würde also in den Markt eintreten und je nachdem, ob der Monopolist tatsächlich stark oder schwach ist, eine Auszahlung von -10 ( $z_2$ ) oder von 40 ( $z_6$ ) erhalten.

## V. Recht und informale Institutionen

### 1. Recht als Preis oder Brennpunkt

Von der Rolle des Rechts war bisher nicht die Rede. Das bedeutet aber nicht, dass die Spieltheorie nur rechtsfreie Räume beschreibt. Das Recht bildet vielmehr ebenso wie die natürlichen Handlungsbedingungen eine Restriktion, die bei der Modellierung eines Entscheidungsproblems Berücksichtigung finden kann. Allerdings hat die Spieltheorie für Normativität keinen systematischen Platz, so dass die im geltenden Recht ausgesprochene Verpflichtung als solche irrelevant ist. Zunächst kann Recht über die Sanktionen Berücksichtigung finden, mit denen eine Rechtspflicht bewehrt ist, indem der Erwartungswert der Sanktion (Sanktionshöhe  $\times$  Eintrittswahrscheinlichkeit) quasi als Preis einer bestimmten Aktion in den Aus-

252

<sup>62</sup> Sowie von den Risikopräferenzen, die wir hier der Einfachheit halber als neutral annehmen.

<sup>63</sup>  $-10 \times 0,5 + 40 \times 0,5 = -5 + 20 = 15$ .

zahlungen berücksichtigt wird.<sup>64</sup> Nehmen wir das Problem der Überfischung von Fischbeständen als Beispiel, bei dem die sozial wünschenswerte Zurückhaltung kein Gleichgewicht ist. Wenn nun mittlere und intensive Befischung verboten und mit einem Bußgeld von 8 belegt werden, welches mit einer Wahrscheinlichkeit von  $\frac{1}{4}$  auch tatsächlich verhängt wird, verringern sich die Auszahlungen dieser Strategien um 2. Es ergibt sich dann das folgende Spiel, in dem Kooperation eine dominante Strategie ist:

|                |          | Staat B fischt |            |            |
|----------------|----------|----------------|------------|------------|
|                |          | wenig          | mittel     | intensiv   |
| Staat A fischt | wenig    | 3, 3           | 1, ← 2     | 0, ← 1     |
|                | mittel   | ↑ 2, 1 →       | ↑ 0, ← 0   | ↑ -1, ← -1 |
|                | intensiv | ↑ 1, 0         | ↑ -1, ← -1 | ↑ -2, ← -2 |

Tabelle 4.23: Beispiel für geänderte Auszahlungen durch Sanktionen

- 253 Auch hier gibt es in der Spielmatrix keine explizite Modellierung für Recht, aber die geänderten Auszahlungen berücksichtigen den institutionellen Rahmen. Und indem Recht die Auszahlungen verändert, ändert es die Natur des Spiels, das nun kein Kooperationsspiel mehr ist, sondern ein Harmoniespiel. Zum anderen kann das Recht Informationen über das zu erwartende Verhalten der Mitspieler enthalten und als Brennpunkt für die Koordination dienen. Das Recht lässt dann die Natur des Spiels unberührt, beeinflusst aber, welches von mehreren Gleichgewichten gespielt wird. Die Entscheidung zwischen Links- oder Rechtsverkehrs als Koordinierungsproblem kann das verdeutlichen.<sup>65</sup> Reist man zum Beispiel nach Namibia und möchte dort am Straßenverkehr teilnehmen, reicht in der Regel die Kenntnis der Rechtslage (Linksverkehr), um zu einem entsprechenden Verhalten motiviert zu werden. Bereits das Eigeninteresse an der Vermeidung von Verkehrsunfällen ist Anreiz genug, in diesem wiederholten Spiel das Gleichgewicht zu wählen, das die anderen Verkehrsteilnehmer bereits spielen, so dass Sanktionen nicht notwendig sind.

## 2. Recht und soziale Normen

### a. Informale Institutionen als Gleichgewichte in wiederholten Interaktionen

- 254 Für Juristen ist es naheliegend, dass ohne Recht Konflikte und Chaos herrschten. Aus spieltheoretischer Sicht würde dies von der Natur der

<sup>64</sup> Hierzu ebenso Rz. 65 ff. sowie Rz. 494 f.

<sup>65</sup> Hierzu Rz. 208.

Spiele abhängen, die die Menschen dann spielten. Agieren die Menschen in dichten Gruppen (*close-knit communities*), deren Interesseninterdependenzen sich als unbestimmt oft wiederholte Spiele mit vielen Personen modellieren lassen, dann stehen ihnen nach dem **Folk-Theorem** unabhängig von der Struktur der jeweiligen Einzelinteraktion (der Basisspiele) viele Gleichgewichte offen, darunter auch kooperative. Natürlich ist die Interaktionsdichte in modernen Gesellschaften weniger dicht als in Dorfgemeinschaften. Aber auch in der Moderne nehmen die Menschen an einer Vielzahl einander überschneidender Gruppen teil, die vergleichsweise häufig und dauerhaft miteinander interagieren (z. B. in Firmen und Behörden oder deren Abteilungen, in Berufsgruppen von z. B. Rechtsanwälten oder Richtern) und deshalb einer Analyse als wiederholtes Spiele zugänglich sein können. Es können dann soziale Normen und andere Institutionen entstehen, deren Stabilität sich aus spieltheoretischer Sicht aus der Stabilität des ausgewählten Gleichgewichts ergibt. Wir nennen solche Institutionen informale Institutionen. Sie bestehen u. a. deshalb fort, weil es für den Einzelnen nachteilig wäre, einseitig von dem etablierten Gleichgewicht abzuweichen. Das heißt nicht, dass andere Gleichgewichte bzw. Institutionen nicht zu erreichen wären, sondern dass dafür das Verhalten vieler geändert werden müsste.

### b. Konventionen, soziale Normen und benachteiligende Normen

Es ist instruktiv nach der Natur des Basisspiels verschiedene Arten informaler Institutionen zu unterscheiden. Alle informale Institutionen im hiesigen Verständnis sind als Gleichgewichte per definitionem selbstdurchsetzend. Es macht aber einen Unterschied, ob das Verhalten, das die Institution vorsieht, im Basisspiel ein Gleichgewicht ist oder nicht. Ist das Basisspiel z. B. ein Koordinationsspiel (wie beim Rechts- oder Linksverkehr), dann sind auch im wiederholten Spiel keine Sanktionen notwendig, um die Stabilität der Institution zu gewährleisten. Man kann insoweit von **Konventionen** reden.<sup>66</sup> Ist das vorgesehene Verhalten dagegen, wie etwa bei Kooperationsnormen in einem Gemeingutsdilemma, kein Gleichgewicht im Basisspiel, dann folgt die Stabilität der Institution erst aus den glaubhaften Sanktionen, die sich die Mitglieder der Gruppe (als Teil der Gleich-

255

<sup>66</sup> Der Sprachgebrauch ist insoweit nicht einheitlich und rekurriert auf unterschiedliche Differenzierungskriterien, vgl. *McAdams/Rasmussen*, Norms and the Law, in: *Polinski/Shavell* (Hrsg.), *Handbook of Law and Economics*, 1573 ff.; *Sugden*, Spontaneous order, *Journal of Economic Perspectives* 3 (1989), 85 ff.; *Ullmann-Margalit*, The Emergence of Norms, 1977; *Young*, The Evolution of Conventions, *Econometrica* 61 (1993), 57 ff.



gewichtsstrategie) gegenseitig androhen. Solche Institutionen kann man in weitgehender Übereinstimmung mit dem rechtssoziologischen Sprachgebrauch als **soziale Normen** bezeichnen, weil sie mit der Erwartung einer Sanktion im Fall der Normverletzung verbunden sind. Sanktionsdrohungen sind nach dem Folk-Theorem in ihrer einfachen Variante dezentral in dem Sinn, dass alle anderen Spieler mit der Einstellung von Kooperation bedroht werden. Mit zunehmender Komplexität der Interaktionen werden Sanktionen häufig spezialisierten Instanzen übertragen, die die Norm Einhaltung überwachen und selektiv nur noch den Normbrecher sanktionieren. Spätestens mit diesem Schritt wird der Übergang von sozialen Normen zum Recht vollzogen (was nicht heißt, dass es nicht Rechtsformen gibt, die ohne zentrale Sanktionsinstanzen auskommen, wie etwa das Völkerrecht). Dass Konventionen und soziale Normen selbstvollziehend sind, bedeutet allerdings keineswegs, dass sie auch die Interessen aller Betroffenen gleichermaßen berücksichtigen. Vielmehr trifft man auf vielerlei informale **Institutionen, die soziale Ungleichheiten oder Diskriminierungen aufrechterhalten**.<sup>67</sup> Dann liegt es nahe, dass das Basisspiel kein reines Kooperations- oder Koordinierungsspiel ist, sondern z. B. ein Kampf der Geschlechter,<sup>68</sup> bei dem die Auszahlungen in den Gleichgewichten des Basis-spiels ungleich verteilt sind. Entsteht hier eine Institution auf einem dieser Gleichgewichte, dann wird eine Seite auf Dauer benachteiligt.<sup>69</sup>

### c. Die Interaktion von Recht und sozialen Normen

- 256 Zwischen Recht und sozialen Normen gibt es verschiedene Formen der Interaktion.<sup>70</sup> Soziale Normen können etwa ein Substitut für Recht sein, ein Regelungsproblem darstellen, aber auch das Recht unterstützen. So beobachten wir, dass es im sozialen Leben häufig untunlich ist und eher Anlass für Verärgerung und Konflikt gibt, wenn Menschen sich ausdrücklich auf ihre Rechte oder die Pflichten des Gegenübers berufen. Ein Grund

<sup>67</sup> „Norms of Partiality“ bei *Ullmann-Margalit*, *The Emergence of Norms*, 1977, S. 134 ff.; für eine Diskussion anhand anthropologischer Evidenz vgl. *Ensminger/Knight*, *Changing Social Norms: Common Property, Bridewealth, and Clean Exogamy*, *Current Anthropology* 38 (1997), 1 ff.

<sup>68</sup> Hierzu Rz. 211 ff.

<sup>69</sup> *Ullmann-Margalit/Sunstein*, *Inequality and Indignation*, *Philosophy and Public Affairs* 30 (2002), 337 ff.; *Wax*, *Expressive Law and Oppressive Norms: A Comment on Richard McAdams “A Focal Point Theory of Expressive Law”*, *Virginia Law Review* 86 (2000), 1731 ff.

<sup>70</sup> Ausführlich *Magen*, *Fairness, Eigennutz und die Rolle des Rechts*, in: *Engel et al.* (Hg.), *Recht und Verhalten*, 2007, 261 ff., sowie *ders.*, *Gerechtigkeit als Proprium des Rechts*, 2010, S. 145 ff.

für die Ablehnung des Rechts kann darin liegen, dass die betreffende Interaktion (ausdrücklich oder stillschweigend) von sozialen Normen reguliert wird, die bestimmen, wie „man“ sich in diesem Kontext verhält. Das Recht wird dann, jedenfalls im Regelfall, durch die soziale Norm substituiert.

Derartige Normen können für das Recht unproblematisch und als Betätigung grundrechtlicher Freiheit sogar legitim sein. Sie können für das Recht aber auch ein Gemeinwohlproblem und Anlass für gesetzliche Intervention sein. Nehmen wir das (sozial erwünschte) **Kartell-Dilemma** vom Beginn dieser Einführung, bei dem zwei Duopolisten durch Steuerung ihrer Absatzmengen kollusiv den Monopolgewinn einstreichen möchten. Wieso untersagt z.B. Art. 101 AEUV Vereinbarungen und abgestimmte Verhaltensweisen zwischen Unternehmen, wenn Kollusion kein Gleichgewicht und damit nicht erreichbar ist? Der erweiterte begriffliche Apparat öffnet eine Erklärungsmöglichkeit: Wenn Duopolisten nicht ein einmaliges, sondern ein unbestimmt oft wiederholtes Gefangenens-Dilemma spielen, dann ist Kollusion sehr wohl ein Gleichgewicht, wenn die Duopolisten ihr Verhalten wechselseitig beobachten und abstimmen können.<sup>71</sup> **Kollusion** ist in unserer Terminologie also eine soziale Norm, die das Recht aber wegen ihrer negativen Folgen für die Verbraucher und die Volkswirtschaft missbilligt. Das Recht kann dann versuchen, die unerwünschte soziale Norm zu unterbinden, indem es Kollusion durch Bußgelder soweit verteuert, dass sie im Gesamtspiel kein oder zumindest kein attraktives Gleichgewicht mehr ist.<sup>72</sup>

Rechts- und soziale Normen können umgekehrt aber auch gleichläufig sein und sich gegenseitig stützen. Nehmen wir das Problem der gegenseitigen Rücksichtnahme im Straßenverkehr (§ 1 StVO). Straßenverkehr ist eine wiederholte Interaktion vieler Personen, bei der gegenseitige Rücksichtnahme je nach Situation ein reines Koordinationsproblem, ein gemischtes Koordinations- und Konfliktspiel oder ein Kooperationsproblem

<sup>71</sup> *Wagner-von Papp*, Marktinformationsverfahren, 2004, S. 92 ff.

<sup>72</sup> Wenn z.B. der Erwartungswert der Kartellbuße die Auszahlungen bei Kollusion (kleine Menge, kleine Menge) von (3, 3) auf (0, 0) drückt, dann liegt die beiderseitige Kollusion unterhalb der *Maximin*-Auszahlung von 2 und ist damit auch im Gesamtspiel kein Gleichgewicht mehr.

|          |              | D2 wählt     |             |
|----------|--------------|--------------|-------------|
|          |              | kleine Menge | große Menge |
| D1 wählt | kleine Menge | ↓ 0, 0 →     | ↓ 1, 4      |
|          | große Menge  | 4, 1 →       | <u>2, 2</u> |

darstellen kann. Allerdings treffen die einzelnen Verkehrsteilnehmer zu selten aufeinander, als dass gegenseitige Rücksichtnahme nach dem Folk-Theorem als Gleichgewicht in einem wiederholten Spiel stabil sein könnte. Ohne das Straßenverkehrsrecht und seine Überwachung und Durchsetzung durch Polizei und Ordnungsbehörden würde ein relevantes Maß an Rücksichtnahme deshalb kaum bestehen können. Das schließt aber nicht aus, dass neben den rechtlichen Sanktionen auch informale Sanktionen wie Ermahnungen, Lichthupen usw. einen wichtigen zusätzlichen Beitrag zur Sicherheit und Leichtigkeit des Straßenverkehrs leisten können.

## § 5 – Vertragstheorie und ökonomische Analyse des Vertragsrechts

**Literatur:** *G.A. Akerlof*, The Market for ‘Lemons’: Quality Uncertainty and the Market Mechanism, *Quarterly Journal of Economics* 84 (1970), 488–500; *K.J. Arrow*, The Economics of Agency, Institute for Mathematical Studies in the Social Sciences, 1984; *G. Calabresi*, The costs of accidents, 1970; *ders.*, Transaction Costs, Ressource Allocation and Liability Rules: A Comment, *Journal of Law & Economics* 11 (1968), 67–73; *R. Coase*, The Problem of Social Cost, *Journal of Law & Economics* 3 (1960), 1–44; *R. Craswell*, Passing On the Costs of Legal Rules: Efficiency and Distribution in Buyer-Seller Relationships, *Stanford Law Review* 43 (1991), 361–398; *M. Eisenberg*, The Limits of Cognition and the Limits of Contract, *Stanford Law Review* 47 (1995), 211–259; *B. Hermanlin/A. Katz/R. Craswell*, Contract Law, in: A. Polinsky/R. Zerbe *et al.* (Hg.), *Handbook of Law and Economics* Vol. I, Ch. 1; *S. Medema/R. Zerbe*, The Coase Theorem, in: B. Bouckaert/G. de Geest *et al.* (Hg.), *Encyclopedia of Law and Economics*, Vol. I, 836–892, *M. Rothschild/J. Stiglitz*, Equilibrium in Competitive Insurance Markets: An Essay on the Economics of Imperfect Information, *Quarterly Journal of Economics* 90 (1976), 629–649; *R. Scott*, Conflict and Cooperation in Long-Term Contracts, *California Law Review* 75 (1987), 2005–2054; *S. Shavell*, Contracts, in: P. Newman *et al.* (Hg.), *The New Palgrave Dictionary of Economics and the Law* (1998), Vol. 1, A-D, 436–445; *M. Spence*, Job Market Signaling, *Quarterly Journal of Economics* 87 (1973), 355–374; *J. Stiglitz*, The Theory of ‘Screening’, Education, and the Distribution of Income, *American Economic Review* 65 (1975), 283–300; *J. Stiglitz/A. Weiss*: Credit Rationing in Markets with Imperfect Information, *American Economic Review* 71 (1981), 393–410; *O.E. Williamson*, Transaction-Cost Economics: The Governance of Contractual Relations, *Journal of Law & Economics* 87 (1979), 233–261; *O.E. Williamson*, *The Economic Institutions of Capitalism: Firms, Markets, Relational Contracting*, 1985.

Als wirtschaftswissenschaftliche (Teil-)Disziplin beschäftigt sich die Vertragstheorie mit der Frage, wie sich ökonomische Akteure im Rahmen bestimmter vertraglicher Arrangements verhalten und welche Wohlfahrtsimplikationen sich hieraus ergeben. Es lässt sich auch stärker normativ gewendet formulieren: Die Vertragsökonomik ist am Entwurf effizienter Vertragsdesigns interessiert. Bei deren Entwicklung berücksichtigt sie die

259

individuellen Anreize der (prospektiven) Vertragsparteien. Diese Parteien begreift sie als (beschränkt) rationale Maximierer ihres eigenen Nutzens. Vor allem aufgrund unvollständiger Information einer oder mehrerer Parteien, aber auch aufgrund anderer Ursachen, gelingt es den Parteien nicht immer, selbst eine optimale Vertragslösung herbeizuführen, also einen Vertrag zu schließen, der ihre gemeinsame Wohlfahrt maximiert. In diesen Fällen stellt sich die Frage, ob das (Vertrags-)Recht dabei helfen kann, den resultierenden Wohlfahrtsverlust zu mindern.

## I. Warum Verträge?

- 260 Der Zentralbegriff der Vertragsökonomik ist – wenig überraschend – der Vertrag. Juristen kennen den Vertrag als eine auf dem Willen der Parteien beruhende Vereinbarung, die typischerweise den Austausch von Leistungen zum Gegenstand hat und vor Gericht durchgesetzt werden kann. Die erhebliche Bedeutung, die der Vertrag und das Vertragsrecht für eine moderne Marktwirtschaft haben, erscheint selbstverständlich und daher keiner weiteren Erklärung wert. Wenn wir uns allerdings ins Gedächtnis rufen, wie sich die Mikroökonomie den praktischen Austausch von Gütern regelmäßig vorstellt, gerät diese Selbstverständlichkeit ins Wanken: In der Mikroökonomie werden wohlfahrtsfördernde Austauschgeschäfte typischerweise als eine Abfolge von Spot-Transaktionen gedacht, in denen Güter, Dienstleistungen und Geld zu einem bestimmten Zeitpunkt simultan (juristisch: „Zug-um-Zug“) zwischen den Parteien getauscht werden.

### 1. Austauschgeschäfte in einer idealen Welt: Das *Coase*-Theorem

- 261 Begreift man Austauschgeschäfte als kostenlose Spot-Transaktionen zwischen zwei allwissenden Parteien, die sämtliche Eigenschaften und Bedingungen des geplanten Geschäfts kennen, führt dies zu allokativ effizienten Ergebnissen. Dies gilt auch für den Fall, dass Rechte Gegenstand des betreffenden Handels sind. Zur Veranschaulichung diene folgendes Beispiel, das dem bahnbrechenden Beitrag von *Ronald Coase* über „das Problem sozialer Kosten“ entlehnt ist: A möchte eine Fabrik bauen, die schädliche Emissionen (bspw. Rauch oder Lärm) ausstößt, und zwar gleich neben dem Grundstück des B. Ganz gleich, ob man den A zur Emission ermächtigt oder stattdessen dem B das Recht zuweist, eine solche Emission untersagen zu lassen, am Ende werden sich A und B auf das (gleiche) effiziente Maß an Emissionen einigen. Denn solange die Produktion einer zusätz-

lichen Einheit an Emissionen (mehr Rauch, mehr Lärm) für A vergleichsweise mehr Nutzen birgt als sie dem B Kosten (etwa in Form von Gesundheitsschäden) verursacht, wird diese zusätzliche „Emissionseinheit“ auch produziert werden. Falls B das Recht zusteht, die Emission durch A verbieten zu lassen, wird A ihm dieses Verbotsrecht abkaufen. A verbleibt dann annahmegemäß immer noch ein Nutzenüberschuss aus der zusätzlichen Emission. Steht hingegen A das Recht zur weiteren Emission zu, so wird dieser sich das Recht wiederum von B abkaufen lassen, sobald diesem der Stopp zusätzlicher Emission mehr wert ist als A die Emission einer weiteren Einheit an Schadstoffen.

Das vorstehende Szenario ist ein Beispiel für das sogenannte **Coase-Theorem**. In seiner starken Form umfasst dieses Theorem zwei Hypothesen. Die sogenannte **Effizienzhypothese** lässt sich wie folgt zusammenfassen: Unter der Annahme, dass die Akteure rational handeln, sich die von ihnen vorgenommenen Austauschgeschäfte kostenlos durchführen lassen und auch keine rechtlichen Hindernisse für solche Geschäfte bestehen, werden (Austausch-)Transaktionen bis zu dem Moment stattfinden, in dem sich die Situation durch weitere Geschäfte nicht mehr verbessern lässt, d. h. bis zu dem Moment, in dem eine optimale Ressourcenallokation erreicht ist.<sup>1</sup> Die **Invarianzhypothese** besagt, dass diese optimale Ressourcenverteilung unabhängig davon eintritt, wie die Berechtigungen an den Ressourcen (im Beispiel: das Emissionsrecht bzw. das Recht, die Emission zu untersagen) ursprünglich verteilt sind. 262

Die Parteien (im Beispiel: A und B) werden kooperieren, um ihren aggregierten Nutzenüberschuss zu maximieren (Effizienzhypothese). Insofern ist die ursprüngliche Verteilung von Berechtigungen an der betreffenden Ressource zwar unerheblich (Invarianzhypothese). Bedeutung erlangt diese ursprüngliche Verteilung (und die relative Verhandlungsmacht der Parteien) jedoch für die Frage, wie sich der aggregierte Nutzenüberschuss auf die Parteien am Ende verteilt! Diese Verteilungsfrage ist jedoch für die Frage der Effizienz des Ergebnisses nicht, jedenfalls nicht unmittelbar von Belang. 263

Bei alledem darf man freilich nicht dem Irrtum erliegen, *Coase* hätte jemals an die Existenz einer solchen *Coase'schen* Welt ohne Transaktionskosten geglaubt. Vielmehr ging es ihm darum, die Relevanz von Rechtsre- 264

<sup>1</sup> *Calabresi*, Transaction Costs, Resource Allocation, and Liability Rules: A Comment, *Journal of Law & Economics* 11 (1968), 67 (68): „If people are rational, bargains are costless, and there are no legal impediments to bargains, transactions will [...] occur to the point where bargains can no longer improve the situation; to the point, in short, of optimal resource allocation.“

geln für die Erzielung allokativer Effizienz in der wirklichen Welt, d.h. einer Welt *mit* Transaktionskosten aufzuzeigen. Wir werden hierauf zurückkommen.<sup>2</sup>

- 265 Für die hier aufgeworfene Frage nach dem ökonomischen Sinn von Verträgen ist hingegen die Erkenntnis von Bedeutung, dass Verträge in einer solchen Idealwelt sowohl als ökonomische wie als rechtliche Institution nur insofern einen Effizienzgewinn versprechen, als sie die Parteien des Austauschgeschäfts dauerhaft an dessen Wirksamkeit binden (und sie davon abhalten, sich einfach mit Gewalt wiederzuholen, was sie gerade weggetauscht haben). Ein entsprechendes Ergebnis ließe sich in vielen Fällen allerdings auch ohne Verträge mit einem schlichten (aber durchsetzbaren) Gewaltverbot herbeiführen.

## 2. Verträge als Instrument der (Vorab-)Bindung und Koordination

- 266 Der eigentliche Wert, den durchsetzungsfähige Verträge für die Parteien haben, wird klar, wenn wir dem Austauschscenario zwei Eigenschaften hinzufügen: Die erste Eigenschaft ist eine gewisse Zeitdauer des Austauschprozesses, der also nicht zu einem bestimmten Zeitpunkt stattfindet, sondern einen länger andauernden Zeitraum in Anspruch nimmt, weil eine Partei in Vorleistung tritt oder Investitionen tätigen muss, bevor der dann simultane Austausch der Güter und Leistungen stattfindet. Die zweite Eigenschaft folgt aus der ersten und betrifft die Unsicherheit über die Absichten der (jeweiligen) Gegenpartei und deren künftiges Verhalten. Zur Veranschaulichung diene folgendes Beispiel: Nehmen wir an, dass A einen Kredit braucht. Die Bank B wird A jedoch kein Darlehen gewähren, solange sie die spätere Rückzahlung gegenüber A nicht durchsetzen kann. Wegen der erforderlichen Vorleistung der B (Auszahlung der Darlehensvaluta an A) und der Unsicherheit über die Absichten und das künftige Verhalten des A (Wird A das Darlehen zurückzahlen oder nicht?), wird B auf einer glaubhaften Vorabbindung des A bestehen. Ein gerichtlich durchsetzbarer Darlehensvertrag stellt das für diese Bindung erforderliche Mittel (*commitment device*) dar. Ein solches Instrument zur Vorabbindung kann indes auch dann erforderlich sein, wenn die Güter und Leistungen gleichzeitig ausgetauscht werden: Stellen wir uns etwa vor, A sei ein Schneider, der auf die Anfertigung von Maßanzügen spezialisiert ist. B möchte einen solchen Anzug von A erwerben. Er lehnt es aber ab, hierfür im Vorhinein zu bezahlen, weil er nicht sicher ist, ob A nicht einfach das Geld nimmt,

---

<sup>2</sup> Hierzu noch unter Rz. 269 ff.

ohne einen Anzug dafür zu liefern. A muss den Anzug daher zunächst anfertigen und dabei darauf vertrauen, dass B den vereinbarten Preis am Ende bezahlen wird. A kann den für B angefertigten Anzug auch nicht ersatzweise an jemand anderen verkaufen, da es sich um einen Maßanzug speziell für B handelt. Anders gewendet: A muss eine spezifische Investition tätigen, wenn er den Anzug anfertigt. Weil nach alledem für A also Unsicherheit in Bezug auf B's Absichten und künftiges Verhalten besteht und er seinerseits keinen Grund hat, B zu vertrauen, brauchen A und B ein Instrument glaubwürdiger Vorabbindung, kurzum: einen durchsetzbaren (Werk-)Vertrag.

Verträge können für die Parteien zudem als Instrument der Koordinierung dienen (*coordination device*): Nehmen wir etwa an, A und B wollten eine Goldmine ausbeuten. A hat das nötige Geld, B die Expertise. Daher vereinbaren sie vertraglich, dass A das erforderliche Kapital in die Unternehmung schießen soll, während B seine Expertise einbringt. Um mit dem Betrieb der Mine zu beginnen, brauchen sie aber noch einen Grubenarbeiter. A würde einen hinreichend fähigen Arbeiter finden, wenn er sich auf die Suche begäbe. Da B die „Bergbau-Community“ kennt, würde er sogar einen exzellenten Arbeiter auftreiben können. Würden nun A und B beide unabhängig voneinander einen Minenarbeiter einstellen, dann hätten sie zwei. Auch wenn das Unternehmen dann immer noch profitabel wäre, fiel der Gewinn höher aus, wenn sie sich mit einem Arbeiter begnügten. Für A und B wäre es daher sinnvoll, sich abzusprechen, d.h. ihre Bemühungen um einen Minenarbeiter zu koordinieren. Dabei könnte B zwar zusagen, die Suche zu übernehmen, A könnte sich aber nicht sicher sein, ob das nur „Gerede“ ist oder B auch tatsächlich sucht. Das Problem lässt sich lösen, indem A und B die Aufgabe „Minenarbeitersuche und -anstellung“ in ihrer vertraglichen Vereinbarung verbindlich und damit durchsetzbar einer Partei zuweisen. Der Vertrag wirkt hier also als Koordinierungsinstrument.

Der Rechtsökonom *Steven Shavell* hat diese Funktionen als *commitment* und *coordination device* im Blick, wenn er den Vertrag definiert als „eine Beschreibung von Aktionen, welche die benannten Parteien zu bestimmten Zeitpunkten durchführen sollen, und zwar in Abhängigkeit von den Bedingungen, die zu den betreffenden Zeitpunkten herrschen“.<sup>3</sup>

<sup>3</sup> *Shavell*, Contracts, in: The New Palgrave Dictionary of Economics and the Law, Vol. 1, A–D, 1998, S. 436: „A contract is a specification of the actions that named parties are supposed to take at various times, as a function of the conditions that then obtain.“



## II. Marktstörungen als Begründung für Vertragsrecht

- 269 Nach dem soeben Gesagten mag man allerdings immer noch fragen, wieso sich das Vertragsrecht nicht auf eine einzige Regel beschränkt, die da lautet: „Sämtliche Verträge, die auf der Willensübereinkunft der Parteien beruhen, werden vom Recht als wirksam anerkannt und von den Gerichten im Bedarfsfall durchgesetzt.“ Der Grund hierfür liegt darin, dass die Vertragsparteien in der realen Welt durch bestimmte Einflüsse und Umstände daran gehindert werden, die für sie optimale, d. h. den gemeinsamen Nutzen maximierende Übereinkunft zu erzielen (sogenannte Markt- oder Verhandlungsversagen).<sup>4</sup> Obwohl also die Parteien als (rationale) Nutzenmaximierer bestrebt sind, ihren gemeinsamen Nutzen mithilfe der beabsichtigten vertraglichen Vereinbarung zu maximieren, kann es am Ende dazu kommen, dass der Vertragsinhalt hinter dieser Zielsetzung zurückbleibt. Schlechtestenfalls können die erwähnten Umstände in der konkreten Verhandlungssituation dazu führen, dass überhaupt kein Vertragsabschluss zustande kommt. Diesen letztgenannten Fall veranschaulicht das folgende Beispiel: Verbraucher V möchte einen Film auf einem Blu-ray-Medium kaufen. Der Elektronik-Einzelhändler E verkauft seine Produkte – wie alle anderen Einzelhändler auch – nur unter der Bedingung, dass der Verbraucher seinen Allgemeinen Geschäftsbedingungen (AGB; das „Kleingedruckte“) zustimmt. Für den vielbeschäftigten V ist es aufwendig und mühsam, seine kostbare Zeit in die Lektüre umfangreicher AGB zu investieren, bloß um eine Blu-ray zu kaufen. Verzichtet V aber auf die Lektüre und kauft die Blu-ray trotzdem, dann läuft er Gefahr, dass E eine Klausel im Kleingedruckten „versteckt“ hat, die V schwer benachteiligt. Falls V die damit verbundenen Risiken für gewichtiger hält als den Nutzen, den er aus dem Kauf der Blu-ray zieht, wird er von dem Erwerb der Silberscheibe Abstand nehmen.
- 270 In derlei Marktstörfällen kann das Vertragsrecht u. U. effiziente Instrumente zur Verfügung stellen, um die Marktstörung oder zumindest deren Wirkung für den Vertragsschluss abzumildern.<sup>5</sup> Auch wenn das Recht im konkreten Fall nicht die wohlfahrtsmaximierende, also optimale Vertragslösung („beste Lösung“ oder *first-best solution*) herbeiführen kann, kann es möglicherweise immer noch einen Wohlfahrtsgewinn be-

<sup>4</sup> Zum Begriff des Marktversagens vgl. Rz. 153.

<sup>5</sup> Vgl. zur Regulierung von „Vertragsklauseln, die nicht im einzelnen ausgehandelt wurden“ (AGB, „Kleingedrucktes“), etwa die Richtlinie 93/13/EWG über missbräuchliche Klauseln in Verbraucherverträgen, ABl. EG v. 21.4.1993, Nr. L 95/29, die in Deutschland in den §§ 305 ff. BGB umgesetzt ist.

werkstelligen („zweitbeste Lösung“ oder *second-best solution*). Ein solcher Wohlfahrtsgewinn durch staatliche Intervention mithilfe des (Vertrags-)Rechts kann allerdings nur unter der Bedingung gelingen, dass diese Intervention in das freie Spiel des Marktes nicht ihrerseits Kosten verursacht, die den angestrebten Wohlfahrtsgewinn gänzlich aufzehren oder sogar übersteigen.<sup>6</sup> Wenn die Gerichte oder der Gesetzgeber diese Erkenntnis vernachlässigen, dann kann das Vertragsrecht unter dem Gesichtspunkt der Effizienz bzw. Wohlfahrtsmaximierung<sup>7</sup> schnell selbst Störfaktor werden. Kurzum: Im konkreten Fall kann es sich als die bestmögliche ökonomische Lösung erweisen, schlicht nichts zu tun, d. h. nicht in den (imperfekten) Marktmechanismus einzugreifen!

Mit diesen Einsichten im Gepäck können wir schnell erkennen, dass die Vertragstheorie und die ökonomische Analyse des Vertragsrechts auf das Engste miteinander verbunden sind. Sie gehen gleichsam „Hand in Hand“: So stellt die Vertragstheorie die Werkzeuge zur Verfügung, mit deren Hilfe sich das rechtliche Umfeld eines vertraglichen Arrangements unter Effizienzgesichtspunkten bewerten lässt. Sie kann damit dem Juristen und mehr noch dem Rechtspolitiker helfen, die Frage zu klären, ob die rechtliche Regelung des konkret betrachteten Entscheidungsproblems nur einen zusätzlichen Kostenfaktor bildet oder ob sie einem Marktversagen entgegenwirkt und hierdurch Wohlfahrtsgewinne erzeugt. 271

Die Ursachen eines solchen Markt- oder Verhandlungsversagens unterteilen die Ökonomen traditionell in verschiedene Kategorien. Zu diesen Kategorien zählen vor allem (1) Externalitäten, (2) unvollständige oder asymmetrisch verteilte Information (Informationsasymmetrien), (3) Marktmacht, (4) begrenzte Rationalität (kognitive Grenzen) und (5) öffentliche Güter.<sup>8</sup> Als „gemeinsamen Nenner“ für alle oder doch die meisten dieser Phänomene haben der Nobelpreisträger *Oliver E. Williamson* und gleichgesinnte Ökonomen den Begriff der **Transaktionskosten** identifiziert<sup>9</sup> und in der Folge das Forschungsfeld der Transaktionskostenökono- 272

<sup>6</sup> Vgl. in diesem Zusammenhang auch Rz. 330 zum Phänomen des Staatsversagens.

<sup>7</sup> Zum Einsatz des Vertragsrechts im Namen der Verteilungsgerechtigkeit s. noch unten in Rz. 323 ff.

<sup>8</sup> Vgl. zu öffentlichen Gütern bereits Rz. 166 sowie zu begrenzter Rationalität Rz. 503 ff.

<sup>9</sup> Vgl. zum Begriff der Transaktionskosten Rz. 163 f.

mik etabliert.<sup>10</sup> Andere Vertragstheoretiker<sup>11</sup> sehen demgegenüber in der unvollständigen Information der (prospektiven) Vertragsparteien die Hauptursache für nicht-optimales Kontrahieren. Diesen Meinungsstreit sollte man allerdings nicht überbewerten, lassen sich doch die Kosten zur Überwindung unvollständiger Information als Transaktionskosten im Sinne der Transaktionskostenökonomik begreifen.

### III. Unvollständige Information – Problem und Lösungen

- 273 Nicht nur unter Vertragstheoretikern mit einem informationsökonomischen Fokus, sondern auch unter den meisten Transaktionskostenökonomien besteht weitgehende Einigkeit, dass unvollständige Information, insbesondere **asymmetrisch verteilte Information**, das praktisch bedeutsamste Hindernis für effizientes Kontrahieren darstellt. Die Auswirkungen imperfekter Information auf die Effizienz eines vertraglichen Handels werden klarer erkennbar, wenn man Folgendes bedenkt: Die *Pareto*-Optimalität einer vertraglichen Vereinbarung setzt voraus, dass die Parteien sämtliche Faktoren und Parameter der ausgehandelten Übereinkunft exakt bepreisen können. Dies umfasst auch die Erwartungswerte<sup>12</sup> der Risiken, welche der Vertrag der einen oder anderen Partei zuweist. Das macht es wiederum erforderlich, dass alle am Vertrag beteiligten Parteien sämtliche Informationen miteinander teilen, die für die Bewertung dieser Faktoren notwendig sind. In der realen Welt besitzt eine Partei hingegen häufig private Informationen über den Vertragsgegenstand, die der anderen Partei (anfänglich) unbekannt sind. Die Unkenntnis über solche bewertungserheblichen Informationen führt dann aber regelmäßig zu falschen Preisen und in der Folge zu ineffizienten (Vertrags-)Ergebnissen.

#### 1. Das Problem adverser Selektion

- 274 In seinem Aufsatz über den „market for lemons“ hat *George A. Akerlof* am Beispiel des Gebrauchtwagenmarkts eindrucksvoll veranschaulicht,

<sup>10</sup> *Williamson*, The Economic Institutions of Capitalism: Firms, Markets, Relational Contracting, 1985, S. 17: „This book advances the proposition that the economic institutions of capitalism have the main purpose and effect of economizing on transaction costs.“

<sup>11</sup> Dies betrifft insbesondere solche Ökonomen, die (auch) in spieltheoretischen Kategorien denken bzw. von diesem Denken beeinflusst sind. Zur Spieltheorie vgl. § 4.

<sup>12</sup> Vgl. zum Begriff des Erwartungswerts bereits Rz. 74f.

welch schädliche Wirkung Informationsasymmetrien haben können.<sup>13</sup> Hierfür müssen wir uns vorstellen, dass zwei verschiedene Sorten von Gebrauchtwagen auf dem Markt verfügbar sind, nämlich solche von guter Qualität und solche von schlechter Qualität.<sup>14</sup> Der Gebrauchtwagenverkäufer kennt die Qualität seines Wagens, die potenziellen Käufer kennen sie hingegen nicht. Letztere wissen aber immerhin, dass der konkrete Wagen entweder gut oder schlecht ist. Ohne weitere Information werden die potenziellen Käufer die Wahrscheinlichkeit ( $p$ ), dass der Wagen von guter Qualität ist, mit 0,5 veranschlagen ( $p_g$ ) und dementsprechend die Wahrscheinlichkeit, dass der Wagen von schlechter Qualität ist, ebenfalls mit 0,5 ansetzen ( $p_b = 1 - p_g$ ). Nehmen wir weiter an, dass ein schlechter Gebrauchtwagen € 2.000, ein guter Wagen hingegen € 4.000 wert ist. In diesem Fall liegt der Erwartungswert des einzelnen Wagens für den Käufer bei € 3.000 ( $= 0,5 \times € 2.000 + 0,5 \times € 4.000$ ).<sup>15</sup>

Was wird nun in einem solchen Szenario geschehen? Der potenzielle Käufer wird nicht bereit sein, mehr als € 3.000 für den Gebrauchtwagen zu zahlen. Die Verkäufer guter Gebrauchtwagen werden hingegen nicht unter € 4.000 verkaufen. Da sie ihre Wagen zu diesem Preis nicht verkaufen können, werden sie aus dem Gebrauchtwagenmarkt ausscheiden. Es bleiben dann nur noch Verkäufer schlechter Gebrauchtwagen übrig. Im Ergebnis werden also keine Gebrauchtwagen von guter Qualität mehr angeboten, obwohl eine dahingehende Nachfrage auf Käuferseite durchaus bestünde. Damit ist die angestoßene Negativdynamik aber leider noch nicht an ihrem Ende. Um dies zu illustrieren, ändern wir unsere Annahmen dahingehend, dass die verbleibenden schlechten Gebrauchtwagen nicht alle von der gleichen (schlechten) Qualität sind, sondern sich gleichmäßig innerhalb der Wertspanne von € 1.000 (die schlechtesten Wagen) bis € 3.000 (die Besten der schlechten Wagen) verteilen. Anders gewendet: Der Erwartungswert der schlechten Wagen beträgt € 2.000. Nachdem nun alle Wagen mit einem Wert von mehr als € 3.000 vom Markt verschwunden sind, werden potenzielle Käufer diesen allgemeinen Qualitätsverlust am Markt beobachten können und dementsprechend ihre Einschätzung des Erwartungswerts der angebotenen Wagen auf ein Niveau von € 2.000

<sup>13</sup> Akerlof, The Market for 'Lemons': Quality Uncertainty and the Market Mechanism, Quarterly Journal of Economics 84 (1970), 488 ff.; Akerlof selbst spricht insofern von einer „Fingerübung“.

<sup>14</sup> Solche Gebrauchtwagen von schlechter Qualität werden in den USA umgangssprachlich „lemons“ genannt.

<sup>15</sup> Dazu oben Rz. 74. Wir unterstellen hier Risikoneutralität des Käufers. Allgemein zur Bedeutung von Risikopräferenzen vgl. Rz. 76.

anpassen. In der Folge werden nun auch die Anbieter von Wagen mit einem Wert von mehr als € 2.000 den Markt verlassen. Dieses (Trauer-)Spiel wird sich solange wiederholen, bis nur noch die Anbieter der schlechtesten Autos mit einem Wert von € 1.000 am Markt verbleiben. Diese Abwärtsdynamik bezeichnet man als das Problem **adverser Selektion**.

- 276 Fassen wir zusammen: Weil dem potenziellen Käufer die Qualität der Gebrauchtwagen verborgen bleibt, werden wohlfahrtsmaximierende Verkäufe von Wagen guter Qualität nicht durchgeführt. Staatliche Intervention (durch Recht) mag hier aus Effizienzgründen angezeigt sein, sofern nicht der Markt selbst Lösungen für das Problem findet.

## 2. Signaling

- 277 Das vorstehende Modell für die Dynamik adverser Selektion berücksichtigt allerdings nicht, dass zumindest einige Verkäufer und Käufer hinreichende Anreize haben mögen, um (weitere) Informationen über die Qualität des betreffenden Wagens zur Verfügung zu stellen bzw. zu sammeln und hierdurch die Informationslücke zwischen den Parteien zu schließen.<sup>16</sup>

### a. Das Konzept des *signaling*

- 278 In unserem Beispiel des Gebrauchtwagenmarkts haben die Anbieter guter Wagen erkennbar ein Interesse daran, den potenziellen Käufern ihre (private) Information über die Qualität des Autos zur Kenntnis zu bringen. Dies können sie dadurch bewerkstelligen, dass sie der Käuferseite die Qualität ihres Wagens über die Bedingungen ihres (Vertrags-)Angebots signalisieren. Freilich werden alle Anbieter von Gebrauchtwagen ihre Fahrzeuge als gut anpreisen. Das ausgesendete Signal der Anbieter tatsächlich guter Autos muss daher glaubwürdig sein. Die potenziellen Käufer müssen also in die Lage versetzt werden, aufgrund des Signals tatsächlich Anbieter guter von Anbietern schlechter Wagen unterscheiden zu können. Die Anbieter guter Fahrzeuge werden daher bestrebt sein, ein Signal zu wählen, dass für Verkäufer schlechter Wagen zu kostspielig und daher nicht imitierbar ist.

---

<sup>16</sup> *Akerlof* hat diesen Punkt durchaus gesehen; vgl. *Akerlof*, The Market for 'Lemons': Quality Uncertainty and the Market Mechanism, *Quarterly Journal of Economics* 84 (1970), 488 (499), wo er eine Garantiehafung als mögliches Signal des Verkäufers anspricht.

In unserem Gebrauchtwagenfall wäre etwa die Übernahme einer Garantie durch den Verkäufer ein im vorstehenden Sinne glaubwürdiges Signal. Nehmen wir an, der Verkäufer verspricht, den Wagen für eine bestimmte Zeitdauer nach Vertragsschluss kostenlos zu reparieren, falls ein Defekt auftreten sollte. Ein solches Garantieverprechen ist ganz offensichtlich für die Verkäufer von Autos guter Qualität deutlich billiger als für Verkäufer schlechter Wagen. Denn die Wahrscheinlichkeit, dass der Garantiefall eintritt, liegt bei den von ihnen angebotenen Wagen deutlich niedriger als bei Wagen von schlechter Qualität. Daher kann sich ein Anbieter guter Wagen ein solches Signal (das Garantieverprechen) (eher) leisten, während es für den Anbieter schlechter Wagen deutlich teurer, möglicherweise zu teuer ist. Letzterenfalls werden Verkäufer schlechter Gebrauchtwagen diese ohne Garantieverprechen (oder mit einem Garantieverprechen von deutlich geringerem Leistungsumfang) anbieten. 279

#### b. Kein *signaling* bei zu hohen Kosten

Ein Signal wird jedoch nur gesendet, wenn es sich für den Sender des Signals auch lohnt. Ist das Signal hingegen zu teuer, weil der Nutzen hinter den Kosten zurückbleibt, dann unterbleibt die Sendung des Signals durch die informierte Partei. *Joseph E. Stiglitz* hat diesen Mechanismus anhand eines Fließbandarbeiters illustriert:<sup>17</sup> Nehmen wir an, der Arbeitgeber habe keine Möglichkeit, die Produktivität des einzelnen Fließbandarbeiters unmittelbar zu überprüfen. Daher erhält jeder Arbeiter den gleichen Lohn, obwohl der Arbeitgeber bereit wäre, Arbeitern mit einer hohen Produktivität mehr zu bezahlen. Der Arbeiter könnte nun seine höhere Produktivität dadurch signalisieren, dass er sehr hart arbeitet, was ihm jedoch zusätzliche Mühe bereitet, also mit zusätzlichen Kosten  $c$  verbunden ist. Dieser (Zusatz-)Aufwand würde mit einer Lohnsteigerung von  $b$  belohnt werden. So lange jedoch  $c \geq b$  lohnt sich die Mühe für den Arbeiter nicht. Er wird dann kein Signal über seine höhere Produktivität senden. 280

#### c. *Signaling* und staatliche Intervention

*Signaling* ist ein Marktmechanismus, der dazu dient, das Problem adverser Selektion zu überwinden. Als Alternativlösung für dieses Problem käme auch eine staatliche Intervention in Betracht, und zwar gerade auch dann, wenn sich das *signaling* für die informierte Partei nicht lohnt. So 281

---

<sup>17</sup> *Stiglitz*, The Theory of 'Screening', Education, and the Distribution of Income, American Economic Review 65 (1975), 283 ff.

haben etwa einige Gliedstaaten der USA sogenannte Used Car Lemon Laws beschlossen. Diese verpflichten Gebrauchtwagenverkäufer dazu, Verbrauchern eine Garantie zu geben, die in ihrem Umfang von der Kilometerleistung des jeweiligen Fahrzeugs abhängig ist. Für Verkäufer sehr schlechter und daher sehr reparaturanfälliger Wagen ist diese Verpflichtung zu teuer, weshalb sie aus dem Gebrauchtwagenmarkt ausscheiden. Der Hauptzweck dieser gesetzlichen Bestimmungen ist freilich nicht die Überwindung von Informationsasymmetrien, sondern der Schutz von Verbrauchern vor potenziell gefährlichen Fahrzeugen.

- 282 Gesetze wie die „Used Car Lemon Laws“ schreiben typischerweise nur Mindeststandards für die Qualität der angebotenen Produkte oder die Fähigkeiten der Anbieter von Dienstleistungen vor. Daher kann es sich für die informierten Anbieter immer noch lohnen, Signale an die Nachfrageseite zu senden, welche die besonders hohe Qualität ihrer Produkte oder Dienstleistungen erkennen lassen. So schreiben beispielsweise die Art. 3 und 5 der Verbrauchsgüterkaufrichtlinie<sup>18</sup> eine zwingende Haftung des Verkäufers vor, wenn das verkaufte Verbrauchsgut zum Zeitpunkt der Lieferung nicht vertragsgemäß war und die Vertragswidrigkeit binnen zwei Jahren nach der Lieferung offenbar wird. Für einen Verkäufer, der von der überlegenen Qualität seiner Produkte überzeugt ist, mag es daher sinnvoll sein, ein entsprechendes Signal zu senden, indem er diese Haftung etwa freiwillig über die genannte Zwei-Jahres-Frist hinaus verlängert.

### 3. Screening

#### a. Screening als Mittel zur Informationsgewinnung über potentielle Vertragspartner

- 283 Der bereits erwähnte *Joseph Stiglitz* gehörte auch zu den ersten Informationsökonomien, die der Theorie des *signaling*<sup>19</sup> eine Theorie des *screening* hinzufügten.<sup>20</sup> Der Begriff des *screening* meint dabei gleichsam das Gegenzenario zum *signaling*: Im *Screening*-Szenario ist es die nicht informierte Partei, die sich darum bemüht, Informationen über den potenziellen

<sup>18</sup> Richtlinie 1999/44/EG vom 25.5.1999 zu bestimmten Aspekten des Verbrauchsgüterkaufs und der Garantien für Verbrauchsgüter, ABl. v. 7.7.1999 Nr. L 171/12.

<sup>19</sup> Die grundlegende Pionierarbeit zu *signaling* leistete *Spence*, *Job Market Signaling*, *Quarterly Journal of Economics* 87 (1973), 355 ff.

<sup>20</sup> *Stiglitz*, *The Theory of 'Screening', Education, and the Distribution of Income*, *American Economic Review* 65 (1975), 283 ff.; am 10. Oktober 2001 erhielten *Akerlof*, *Spence* und *Stiglitz* den Nobelpreis für Wirtschaftswissenschaften für ihre Analyse von Märkten mit asymmetrisch verteilten Informationen.

Vertragspartner oder die Qualität der von ihm angebotenen Güter zu erlangen. Die nicht informierte Partei „prüft“ also den potenziellen Vertragspartner.

*Rothschild* und *Stiglitz* illustrieren diesen Mechanismus anhand des 284 Versicherungsmarktes, auf dem Versicherungsunternehmen Versicherungen verkaufen:<sup>21</sup> Die Versicherungsnehmer sind risikoavers und schließen Versicherungsverträge ab, um erwartete Einbußen bei Eintritt des versicherten Risikos, etwa aufgrund möglicher Unfälle, auszugleichen. Da der Grenznutzen ihres Einkommens abnimmt<sup>22</sup>, sind sie bereit, eine Versicherungsprämie zu zahlen, die in der Summe etwas höher liegt als der erwartete Vermögensverlust aufgrund des versicherten Risikos, etwa aufgrund von Unfällen. Nehmen wir nun an, es gebe zwei Typen von Versicherungsnehmern, nämlich solche mit einem niedrigen Unfallrisiko, also einer niedrigen Wahrscheinlichkeit des Eintritts von Unfällen  $p_l$  („guter“ Typ), und solche mit einer hohen Unfallneigung  $p_h$  („schlechter“ Typ). Für die Versicherungsunternehmen stellt sich nun die Frage, welche Verträge sie wem anbieten, um ihre erwarteten Gewinne zu maximieren. Dabei stehen sie – so unsere Annahme – vor dem Problem, dass sie nicht ohne Weiteres erkennen können, zu welchem Typ der einzelne Versicherungsnehmer gehört. Sie brauchen daher weitere Informationen, um den für den konkreten Versicherungsnehmer passenden Versicherungsvertrag anbieten zu können. Daher werden die Versicherer die nötigen Informationen durch ein *screening* der Marktgegenseite zu erlangen suchen.

### b. Selbstselektion durch Vertrag

Eine Möglichkeit, ein solches *screening* vorzunehmen, besteht darin, die 285 potenziellen Vertragspartner (hier: die Versicherungsnehmer) durch eine bestimmte Preisgestaltung des Vertragsangebots zu zwingen, private Informationen (hier: über die Unfallneigung) zu offenbaren. Dieser Mechanismus nennt sich **Selbstselektion** (*self-selection*). In unserem Versicherungsbeispiel kann ein solcher Selbstselektionsmechanismus etwa wie folgt aussehen: Die Versicherer bieten zwei verschiedene Verträge an. Der erste Vertrag sieht eine volle Deckung möglicher Unfallschäden vor, während der zweite Vertrag lediglich eine teilweise Deckung gewährt. Dafür sind die zu zahlenden Versicherungsprämien bei Teildeckung niedriger als

<sup>21</sup> *Rothschild/Stiglitz*, Equilibrium in Competitive Insurance Markets: An Essay on the Economics of Imperfect Information, Quarterly Journal of Economics 90 (1976), 629 ff.

<sup>22</sup> Zum abnehmenden Grenznutzen des Einkommens bzw. von Geld vgl. Rz. 73.



bei einer vollen Deckung der Schäden. Für Versicherungsnehmer des „guten“ Typs mag es dann attraktiver sein, den Vertrag mit bloßer Teildeckung zu wählen, da bei ihnen die Wahrscheinlichkeit eines Schadenseintritts (Unfall) vergleichsweise gering ist. Demgegenüber mögen die Versicherungsnehmer des „schlechten“ Typs mehr von einer vollen Deckung haben, weil sie ein relativ hohes Unfall- und damit Schadensrisiko haben.

| Typ      | Unfall-<br>wahrschein-<br>lichkeit | Unfall-<br>kosten | Prämie:<br>volle<br>Deckung | Prämie:<br>80%ige<br>Deckung | Verlust:<br>volle<br>Deckung | Verlust:<br>80%ige<br>Deckung |
|----------|------------------------------------|-------------------|-----------------------------|------------------------------|------------------------------|-------------------------------|
| Gut      | 10 %                               | 100               | 21                          | 18                           | 21                           | 20                            |
| Schlecht | 20 %                               | 100               | 21                          | 18                           | 21                           | 22                            |

Tabelle 5.1: Selbstselektion im Versicherungsmarkt (*separating equilibrium*)

- 286 Dies lässt sich an folgendem, in Tabelle 5.1 dargestellten Zahlenbeispiel näher veranschaulichen: Den Versicherungsnehmer vom „schlechten“ Typ trifft eine Wahrscheinlichkeit  $p_h$  von 20 %, einen Unfall zu erleiden, der Kosten von 100 mit sich bringt. Hieraus folgt ein erwarteter Verlust von 20. Für einen Versicherungsvertrag mit voller Deckung verlangt der Versicherer eine Prämie von 21. Schließt der Versicherungsnehmer diesen Vertrag ab, folgt für ihn daraus ein sicherer Verlust in Höhe der Prämienzahlung, also 21 ( $0,2 \times 100 - 0,2 \times 100 - 21$ ). Entschiede sich der Versicherungsnehmer dieses Typs hingegen für den Versicherungsvertrag mit Teildeckung in Höhe von 80 % und einer (niedrigeren) Prämienzahlung von 18, dann bedeutete dies für ihn im Ergebnis einen erwarteten Verlust von 22 ( $0,8 \times 0,2 \times 100 - 0,2 \times 100 - 18$ ). Der Versicherungsnehmer vom „schlechten“ Typ bevorzugt daher den ersten Vertrag mit voller Deckung. Der Versicherungsnehmer vom „guten“ Typ erleidet hingegen nur mit einer Wahrscheinlichkeit  $p_l$  von 10 % einen Unfall, der 100 kostet. Ihn trifft daher bei Wahl des zweiten Vertrags mit 80 %iger Deckung lediglich ein erwarteter Verlust von 20 ( $0,8 \times 0,1 \times 100 - 0,1 \times 100 - 18$ ). Die Entscheidung für eine volle Deckung bedeutete hingegen auch für den „guten“ Typ einen sicheren Verlust von 21 ( $0,1 \times 100 - 0,1 \times 100 - 21$ ). Daher wird ein Versicherungsnehmer vom „guten“ Typen für den Vertrag mit teilweiser Deckung optieren, jedenfalls solange der Nutzengewinn durch den Ausschluss des verbleibenden Risikos  $< 1$  ist.
- 287 Obwohl also beide Typen von Versicherungsnehmern hierdurch Vermögensverluste erleiden, werden sie einen Versicherungsvertrag abschließen, solange er ihnen einen Nutzengewinn angesichts ihrer Risikoaversion verspricht. Allerdings ergibt sich auch in diesem Fall erfolgreicher Selbstselektion für den „guten“ Typus ein Nutzenverlust im Vergleich zu einem Alter-

nativszenario mit vollständiger Information der Beteiligten. Denn auch der „gute“ Typ präferiert aufgrund seiner Risikoaversion grundsätzlich einen Versicherungsvertrag mit voller Deckung (allerdings zu einem niedrigeren Preis als 21)!

Natürlich funktioniert ein solches *screening* durch das Angebot von Verträgen, die eine Selbstselektion der Marktgegenseite (im Beispiel: der Versicherungsnehmer vom „guten“ und „schlechten“ Typ) anstoßen, nur unter der Voraussetzung, dass der Anbieter solcher Verträge (im Beispiel: die Versicherungsgesellschaft) auch eine passende Preisgestaltung zugrundelegt. Diese Preisgestaltung muss so gewählt sein, dass alle Typen besser stehen, wenn sie den jeweils für ihren Typ vorgesehenen Vertrag wählen, sich also auf die entsprechenden Vertragsarten aufteilen (*separating equilibrium*).<sup>23</sup> Demgegenüber taugt eine konkrete Preisgestaltung nicht für ein *screening* der potenziellen Vertragspartner, soweit verschiedene Typen dieselbe Vertragsart wählen, weil es für beide (bzw. mehrere oder alle) Typen die nutzenmaximierende Option ist (*pooling equilibrium*). In der Konsequenz erlangt der uninformierte Vertragsanbieter keine zusätzlichen Informationen über die (potenziellen) Vertragspartner. Zur Veranschaulichung diene das in Tabelle 5.2 dargestellte Zahlenbeispiel:

| Typ      | Unfall-<br>wahrschein-<br>lichkeit | Unfall-<br>kosten | Prämie:<br>volle<br>Deckung | Prämie:<br>80%ige<br>Deckung | Verlust:<br>volle<br>Deckung | Verlust:<br>80%ige<br>Deckung |
|----------|------------------------------------|-------------------|-----------------------------|------------------------------|------------------------------|-------------------------------|
| Gut      | 10 %                               | 100               | 21                          | 19                           | 21                           | 21                            |
| Schlecht | 20 %                               | 100               | 21                          | 19                           | 21                           | 23                            |

Tabelle 5.2: Selbstselektion im Versicherungsmarkt (*pooling equilibrium*)

Das Zahlenbeispiel in Tabelle 5.2 unterscheidet sich von demjenigen in Tabelle 5.1 nur dadurch, dass die vom Versicherer für den Vertrag mit 80 %iger Teildeckung verlangte Prämie 19 anstatt 18 beträgt. Für den Versicherungsnehmer vom „schlechten“ Typ ändert sich hierdurch nichts. Der sichere Verlust von 21 bei Wahl des Vertrages mit voller Deckung ist für ihn im Vergleich zum Vertrag mit Teildeckung nun sogar noch attraktiver. Denn letzterer resultiert für den „schlechten“ Typus in einem erwarteten Verlust von 23 ( $0,8 \times 0,2 \times 100 - 0,2 \times 100 - 19$ ). Für den Versicherungsnehmer vom „guten“ Typus ist die Erhöhung der Prämie auf 19 hingegen entscheidend: Die Wahl des Vertrags mit 80 %iger Deckung führt für ihn

<sup>23</sup> Vgl. auch die Definition bei *Rothschild/Stiglitz*, Equilibrium in Competitive Insurance Markets: An Essay on the Economics of Imperfect Information, Quarterly Journal of Economics 90 (1976), 629 (633).

nunmehr zu einem erwarteten Verlust von 21 ( $0,8 \times 0,1 \times 100 - 0,1 \times 100 - 19$ ), während die Wahl des Vertrages mit voller Deckung unverändert einen *sicheren* Verlust von 21 bedeutet. Bei einer solchen Preisgestaltung bevorzugt daher auch der „gute“ Typ den Vertrag mit voller Deckung, da er annahmegemäß (Risikoaversion!) auch das restliche Risiko eliminieren möchte, das mit der bloßen Teildeckung einhergeht.

- 290 Stellt sich ein solches *Pooling*-Gleichgewicht (*pooling equilibrium*) ein, dann zahlen die „guten“ Typen einen vergleichsweise hohen Preis, während die „schlechten“ Typen einen in Anbetracht ihres Risikoprofils zu niedrigen Preis zahlen. Es kommt also zu einer Quersubventionierung des „schlechten“ Typs. Hieraus können im Versicherungsszenario Wohlfahrtsverluste resultieren, wenn und weil der *Pooling*-Vertrag für den „schlechten“ Typus suboptimale Anreize zur Reduzierung des Schadensrisikos setzt. Wohlfahrtsverluste ergeben sich ferner, wenn die Versicherungsprämie für bestimmte Versicherungsnehmer des „guten“ Typs zu hoch ist und sie deshalb – entgegen ihrem eigentlichen Wunsch – auf eine Versicherung ganz verzichten.

### c. Reduzierung der *Screening*-Kosten durch staatliche Intervention

- 291 Die Kosten für ein *screening* können u.U. durch staatliche Intervention (erheblich) reduziert werden. So statuiert etwa § 19 Versicherungsvertragsgesetz (VVG) eine Anzeigepflicht des Versicherungsnehmers vor Vertragsschluss über die ihm bekannten und für den Versicherer erheblichen Gefahenumstände. Kommt der Versicherungsnehmer dieser Pflicht nicht nach, dann kann der Versicherer je nach den konkreten Umständen entweder vom Vertrag zurücktreten, diesen kündigen oder eine (u. U. auch rückwirkende) Anpassung der Vertragsbedingungen verlangen.<sup>24</sup>

### 4. Marktmacht und unvollständige Information

- 292 Monopol- oder Marktmacht ist als eine nachgerade klassische Ursache von Marktversagen bekannt. Wie bereits an anderer Stelle gezeigt, liegt das Problem darin, dass der Monopolist sein Angebot unter das optimale Niveau absenkt, um seine Produzentenrente zu maximieren. Hierdurch verursacht er ein wohlfahrtsminderndes *deadweight loss*.<sup>25</sup> Man hat versucht,

<sup>24</sup> Vgl. für eine schärfere, nicht in dieser Weise durch Verbraucher- bzw. Versicherungsnehmerschutz geprägte Regelung in Sec. 21(1), 28 Australian Insurance Contracts Act 1984.

<sup>25</sup> Hierzu Rz. 149 f., Abbildung 3.19.

dieses Konzept der Marktmacht auf den Vertragsschluss bzw. auf das Vertragsrecht zu übertragen, um hiermit ineffiziente Vertragsklauseln zu erklären. Der deutsche Emigrant *Friedrich Kessler* hat diese Idee erstmals in einem Pionierbeitrag aus den 1940er-Jahren vorgestellt. Er nahm dabei vorformulierte Standardverträge in den Blick. Diese würden von Unternehmen mit starker Verhandlungsmacht „schwächeren“ Parteien angeboten, die nur noch die Wahl hätten, im Sinne eines „Friß oder stirb!“ den Vertrag im Ganzen wie angeboten zu akzeptieren oder vom Vertragsschluss gänzlich Abstand zu nehmen. Hier, so die Schlussfolgerung *Kesslers*, müsse das Recht einschreiten.<sup>26</sup>

Um diese Forderung vernünftig bewerten zu können, muss man sich erneut bewusst machen, dass Ökonomen ein Gefälle in der Verhandlungsmacht nur dann als problematisch begreifen, wenn und soweit es zu einem Wohlfahrtsverlust führt. Nur wenn ein solcher Wohlfahrtsverlust droht, stellen sich die Folgefragen, ob und wie das rechtliche Umfeld dabei helfen kann, diesen Verlust zu reduzieren. Sind daher zunächst einmal die Wohlfahrtsimplikationen solcher „Verhandlungs- oder Vertragsdisparitäten“ zu ermitteln, dann ist hierfür die Erkenntnis ganz entscheidend, dass Monopolmacht nicht die alleinige Ursache des als *deadweight loss* beschriebenen Wohlfahrtsverlusts ist. Könnte der Monopolist nämlich eine perfekte Preisdiskriminierung gegenüber der Marktgegenseite vornehmen, dann würde er für sich den maximal erreichbaren Überschuss erzielen, ohne dass ein Wohlfahrtsverlust einträte. Anders gewendet: Der mit Monopolmacht einhergehende Wohlfahrtsverlust beruht letztlich auf der unvollständigen Information des Monopolisten über die Reservationspreise der (potenziellen) Vertragspartner.<sup>27</sup>

Auf Grundlage dieser Einsicht hat der US-amerikanische Vertragsrechtswissenschaftler *Richard Craswell* zu Recht darauf hingewiesen, dass ein Monopolist zwar einen Anreiz haben mag, hohe Preise zu verlangen und hohe Gewinne zu erzielen. Aber er hat keinen Anreiz, diesen Effekt indirekt über ineffiziente Vertragsklauseln zu erreichen, wenn er stattdessen einfach den Preis für das angebotene Produkt oder die angebotene Dienstleistung erhöhen kann. Setzt ein Monopolist nun gleichwohl ineffiziente Vertragsklauseln ein, etwa in vorformulierten Standardverträgen, liegt die Annahme nahe, dass diese Klauseln dazu dienen, zwischen verschiedenen Gruppen potenzieller Vertragspartner zu diskriminieren. Solche Ver-

<sup>26</sup> *Kessler*, Contracts of Adhesion: Some Thoughts about the Freedom of Contracts, Columbia Law Review 43 (1943), 629 (632).

<sup>27</sup> Hierzu bereits Rz. 160f.

tragsklauseln könnten mit anderen Worten als *Screening*-Instrument dienen.<sup>28</sup> Greift nun das Recht in diese Praxis ein und erklärt solche Klauseln für unwirksam, dann hätte dies nur zu Folge, dass der Monopolist von allen potenziellen Vertragspartnern denselben Preis verlangt. Dabei bleibt aber völlig unklar, ob dies zu einer Wohlfahrtsverbesserung oder möglicherweise gar zu einer Verschlechterung führen würde. *Craswell* kommt daher zu dem Schluss, dass Monopolmacht keine gute Erklärung für ineffiziente Vertragsklauseln ist. Letztlich gehe es auch in diesen vermeintlichen Marktmachtfällen eher um Probleme unvollständiger bzw. asymmetrisch verteilter Information.<sup>29</sup>

#### IV. Kognitive Beschränkungen und nichtrationales Verhalten

##### 1. Kognitive Schranken als Ursache unvollständiger Information

- 295 Ebenso wie das *Coase*-Theorem unterstellt auch die Vertragstheorie traditionellerweise und in Übereinstimmung mit dem Verhaltensmodell des *homo oeconomicus* die Rationalität der Vertragsparteien.<sup>30</sup> Diese Rationalitätsannahme ist entscheidend für die Begründung von Vertragsfreiheit und die Argumentation gegen staatliche Eingriffe. Von den Vertragsparteien kann nur dann erwartet werden, dass sie ihre gemeinsame Wohlfahrt maximieren, wenn sie den (subjektiven) Erwartungsnutzen ihrer Entscheidungsoptionen korrekt berechnen können.<sup>31</sup>
- 296 Allerdings ist diese traditionelle Rationalitätsannahme heute durch zahlreiche empirische Befunde widerlegt und damit als unrealistisch entlarvt. So ist es inzwischen allgemein anerkannt, dass Menschen nur über begrenzte Kapazitäten verfügen, um Informationen zu suchen, aufzunehmen und zu verarbeiten. Diese sogenannte **beschränkte Rationalität**<sup>32</sup> (*bounded rationality*) menschlicher Akteure ist der gedankliche Ausgangspunkt für ein realistischeres Modell menschlichen Entscheidens, das sein Schöpfer, der Sozialwissenschaftler *Herbert Simon*, als *satisficing* bezeich-

<sup>28</sup> Vgl. Rz. 283 ff.

<sup>29</sup> *Craswell*, Freedom of Contract, Coase Lecture Series, 1995, S. 6f.: „[T]he focus on monopoly power is really a red herring where contract terms are concerned. If courts and laypeople tend to associate inefficient terms with monopolies, it's probably because monopoly is the only form of market failure that courts and laypeople are familiar with.“

<sup>30</sup> Zum Modell des *homo oeconomicus* allgemein bei Rz. 69 ff.

<sup>31</sup> Hierzu bereits Rz. 273 f.

<sup>32</sup> Dazu oben Rz. 86 und unten Rz. 503 ff.

net hat.<sup>33</sup> Danach zielt der menschliche Entscheider nicht auf die optimale, d.h. die nutzenmaximierende Option, da dies für ihn wegen seiner beschränkten Informationsaufnahme- und -verarbeitungskapazitäten zu kostspielig ist. Vielmehr begnügt er sich mit einem hinreichend zufriedenstellenden<sup>34</sup> Entscheidungsergebnis und spart so „Entscheidungskosten“. Der Nutzen-Kosten-Saldo dieses Vorgehens ist für den beschränkt rationalen Entscheider am günstigsten und der Entscheidungsmodus des *satisficing* für solche Entscheider daher das rational gebotene Vorgehen.

Allerdings sind menschliche Entscheider nicht nur beschränkt rational im *Simon*'schen Sinne. Sie neigen schlechterdings auch zu systematischen Entscheidungsfehlern, insbesondere bei Entscheidungen unter Unsicherheit.<sup>35</sup> Der Rechtswissenschaftler *Melvin A. Eisenberg* ging als einer der Ersten der Frage näher nach, was diese Einsichten über das tatsächliche Entscheidungsverhalten menschlicher Akteure für den Grundsatz der Vertragsfreiheit und seine Beschränkung durch Vertragsrecht bedeuten.<sup>36</sup> Er kommt dabei zu der Erkenntnis, dass die Kontrahenten die mit dem Vertrag verbundenen Risiken nicht selten falsch wahrnehmen und diese ebenso fehlbewerten wie den Nutzen derjenigen Vertragsbestimmungen, die diese Risiken der einen oder anderen Partei zuweisen. Im Ergebnis werden daher der Abschluss des Vertrages sowie die ihn ausgestaltenden Bestimmungen von einer Partei oder gar beiden falsch bepreist. Dies führt wiederum dazu, dass die Parteien eben nicht denjenigen Vertrag wählen, der ihren gemeinsamen Nutzen maximieren würde. Der hieraus entstehende Wohlfahrtsverlust bietet dann möglicherweise eine Rechtfertigung für staatliche Eingriffe durch Vertragsrecht.

Das von *Herbert Simon* entwickelte Konzept der beschränkten Rationalität haben die Vertragstheoretiker sehr bald in ihre Analysen und Modelle integriert, ohne das Grundmodell rationaler Wahl (*rational choice*) aufzugeben. Dies geschah und geschieht dadurch, dass „Lesekosten“ oder „Verstehenskosten“ einfach als eine weitere Form von Transaktionskosten angesehen werden. Aus dieser Perspektive lassen sich die beschränkte Rationalität oder, allgemeiner, die kognitiven Beschränkungen der Vertragsparteien auch als Problem unvollständiger Information begreifen. Die mei-

<sup>33</sup> Vgl. etwa *Simon*, Theories of Decision-Making in Economics and Behavioral Science, *American Economic Review* 49 (1959), 253 (262 ff.).

<sup>34</sup> Daher die Bezeichnung „*satisficing*“.

<sup>35</sup> Hierzu Rz. 503 ff.

<sup>36</sup> *Eisenberg*, The Limits of Cognition and the Limits of Contract, *Stanford Law Review* 47 (1995), 211 ff.

sten Ökonomen, insbesondere die Vertreter der Verhaltensökonomik<sup>37</sup> verstehen kognitive Schranken und Entscheidungsfehler aufgrund von Wahrnehmungsverzerrungen und der Anwendung von Heuristiken (*biases and heuristics*) hingegen als eine eigenständige Kategorie des Markt- oder Verhandlungsversagens.<sup>38</sup>

## 2. Staatliche Intervention durch paternalistisches Vertragsrecht

- 299 Die Entdeckung immer neuer „Verhaltensanomalien“ durch Psychologen und Experimentalökonomen, d.h. Abweichungen des tatsächlichen Entscheidungsverhaltens vom Rationalwahlmodell, hat nicht wenige Rechtswissenschaftler dazu veranlasst, eine Politik des Rechtspaternalismus zu propagieren: Der irrende und rational defizitäre Mensch soll danach durch das Recht vor seinen eigenen Fehlentscheidungen geschützt werden. Die hiermit zusammenhängenden Fragen werden noch ausführlich in § 8 diskutiert und behandelt. Hier soll daher abschließend lediglich ein aktuelles Beispiel zur Veranschaulichung eines solchen Rechtspaternalismus aus dem (Verbraucher-)Vertragsrecht gegeben werden:
- 300 Nach Art. 18 der Wohnimmobilienkreditrichtlinie<sup>39</sup> ist der Kreditgeber nicht nur verpflichtet, die Kreditwürdigkeit des kreditnehmenden Verbrauchers vor Abschluss des Kreditvertrags zu prüfen. Er darf dem Verbraucher den Kredit zudem nur bereitstellen, wenn „aus der Kreditwürdigkeitsprüfung hervorgeht, dass es wahrscheinlich ist, dass die Verpflichtungen im Zusammenhang mit dem Kreditvertrag in der gemäß diesem Vertrag vorgeschriebenen Weise erfüllt werden“. Führt man sich die typische Risikostruktur eines Darlehensvertrags und die Anreize der Parteien eines solchen Vertrags vor Augen, erscheint diese Regelung merkwürdig. Wenn der Darlehensgeber bereit ist, das Risiko des Kreditausfalls zu übernehmen, warum sollte er dann von Rechts wegen daran gehindert werden, den Darlehensvertrag mit dem Verbraucher durchzuführen, den dieser offensichtlich ebenfalls will?
- 301 Die Antwort hierauf wird klarer, wenn man berücksichtigt, dass der Uniongesetzgeber mit seinen Regeln zur Kreditwürdigkeitsprüfung (unter anderem) der Überschuldung von Verbraucherhaushalten entgegenwirken will. Der Richtlinienggeber hat mit anderen Worten Zweifel, ob die Verbraucher ausreichende Fähigkeiten besitzen, ihr Risiko, den Immobili-

<sup>37</sup> Zur verhaltensökonomischen Forschung vgl. § 8.

<sup>38</sup> Vgl. dazu bereits Rz. 269 ff.

<sup>39</sup> Richtlinie 2014/17/EU vom 4.2.2014 über Wohnimmobilienverträge für Verbraucher, ABl. EU v. 28.2.2014, Nr. L 60/34.

enkredit nicht zurückzahlen zu können, zutreffend einzuschätzen. Verbraucher sollen daher daran gehindert werden, Kreditverträge abzuschließen, die sie – von ihnen selbst unerkannt – überfordern würden.

## V. Anreizprobleme und unvollständige Information nach Vertragsschluss

### 1. *Moral hazard*

#### a. Das Phänomen des *moral hazard*

Informationsasymmetrien treten nicht nur zum Zeitpunkt des Vertragsschlusses auf, sondern auch später im Stadium der Vertragsdurchführung. Solche Informationsasymmetrien können Anreize zu wohlfahrtsminderndem Verhalten setzen. Das bekannteste dieser Anreizprobleme wird als *moral hazard* bezeichnet. Der Begriff soll dem Versicherungswesen entstammen und lässt sich mit einem Beispiel aus diesem Bereich leicht erklären. Hierzu betrachten wir erneut das bereits bemühte Versicherungsszenario:<sup>40</sup> Nehmen wir an, der Versicherungsnehmer mit der höheren Unfallneigung kauft eine Versicherung mit voller Schadensdeckung. Wegen seiner *Ex-ante*-Unfallwahrscheinlichkeit von 20 % bei Unfallkosten von 100, muss er eine Versicherungsprämie in Höhe von 20 (Erwartungswert seines Unfallrisikos) zuzüglich Verwaltungskosten und Unternehmerlohn zahlen. Ist der Vertrag abgeschlossen, stellt sich nun das Problem, dass der vollumfänglich versicherte Versicherungsnehmer keinen optimalen Anreiz mehr hat, Maßnahmen zur Unfallvermeidung zu treffen. In der Folge steigt seine Unfallneigung damit über das Niveau vor Vertragsschluss (das *Ex-ante*-Niveau) an. 302

Der Versicherer hat daher ein Interesse daran, dieses Problem einzuhegen, indem er dem Versicherungsnehmer etwa bestimmte risiko- bzw. schadensmindernde Verhaltenspflichten auferlegt. So könnte etwa ein Kfz-Kaskoversicherer vom Versicherungsnehmer verlangen, die Bremsen regelmäßig zu überprüfen und in Stand zu halten. Derlei Vorkehrungen des Versicherers helfen jedoch dann nicht weiter, wenn es um Aktivitäten des Versicherungsnehmers geht, die der Versicherer nicht beobachten und überprüfen kann (sogenannte *hidden actions*). So weiß das Versicherungsunternehmen zunächst einmal nicht, ob der Versicherungsnehmer 303

---

<sup>40</sup> Vgl. Rz. 286 mit Tabelle 5.1.



ein vorsichtiger und aufmerksamer Fahrer ist oder ein aggressiver und sorgloser Fahrer.<sup>41</sup>

### b. „Agentur-Verträge“ (*agency contracts*) und das Prinzipal-Agenten-Problem

- 304 Die soeben beschriebenen Anreizprobleme treten auch im Rahmen von „Agentur-Verträgen“ (*agency contracts*) auf, d.h. bei Geschäftsbesorgungsverträgen im weiteren Sinne wie etwa Arbeits- oder Dienstverträgen oder ähnlichen Arrangements wie der Bestellung zum Geschäftsführer einer Gesellschaft: Der Agent (etwa der Arbeitnehmer oder der Geschäftsführer) wird im Interessenkreis seines Prinzipals (des Arbeitgebers, des Geschäftsherrn oder der Gesellschaft) tätig und ist daher verpflichtet, im besten Interesse des Prinzipals zu handeln. Der Prinzipal bezahlt den Agenten im Gegenzug, und zwar typischerweise auf Grundlage einer Entgeltstruktur, bei der die Höhe der Bezahlung von den konkreten Handlungen des Agenten abhängt (etwa über eine Provisionsvereinbarung oder im Wege einer leistungsbezogenen Entlohnung). In der Konsequenz hat das Verhalten des Agenten sowohl Auswirkungen auf die Wohlfahrt des Prinzipals (etwa höherer Profit aufgrund einer hohen Zahl der namens und für Rechnung des Prinzipals getätigten Geschäftsabschlüsse) als auch auf die Wohlfahrt des Agenten (höhere Provisionszahlungen bei höherer Zahl von Geschäftsabschlüssen oder schlicht Weiterbeschäftigung).
- 305 Allerdings ist der Agent ungeachtet seiner Verpflichtung auf das Interesse des Prinzipals nach dem Standardverhaltensmodell zunächst einmal und vor allem ein Maximierer seines *eigenen* Nutzens. Daher entstehen im Falle von Informationsasymmetrien, genauer: bei Unsicherheit des Prinzipals über die konkreten Aktivitäten des Agenten, Anreizprobleme in der Person des Agenten. Nach *Kenneth Arrow* unterscheidet man dabei zwei (Haupt-)Kategorien der asymmetrischen Informationsverteilung zu Lasten des Prinzipals:<sup>42</sup> Zum einen entstehen Informationsasymmetrien, wenn der Prinzipal die Aktivitäten des Agenten nicht unmittelbar beobachten kann und auch keine präzisen Rückschlüsse aus dessen Arbeitsergebnissen ableiten kann (*hidden action*).<sup>43</sup> Zum anderen hat der Agent private Informationen, wenn er etwas beobachtet oder entdeckt hat, das

<sup>41</sup> In der Realität passen die Versicherer die Prämienhöhe an die Häufigkeit der Inanspruchnahme durch den Versicherungsnehmer an. Daher zahlen unfallgeneigte Versicherungsnehmer höhere Prämien.

<sup>42</sup> *Arrow*, *The Economics of Agency*, 1984, S. 3 ff.

<sup>43</sup> *Ders.*, aaO, S. 3: „The most typical hidden action is the effort of the agent.“

dem Prinzipal verborgen bleibt. In der Folge kann der Prinzipal nicht erkennen, ob der Agent diese – ihm ja unbekannte – Information im besten Interesse des Prinzipals nutzt oder nicht (*hidden information*).

Die letztgenannte Kategorie lässt sich gut am Beispiel des Vorstands einer Aktiengesellschaft illustrieren. Der Vorstand ist als Agent der Aktionäre tätig und daher verpflichtet, im besten Interesse der AG und ihrer Aktionäre zu handeln. Nehmen wir nun an, die AG betreibt ein Unternehmen, das Getränke produziert und vertreibt (nennen wir sie L-AG, im Weiteren kurz L). L kauft große Mengen Cola-Sirup von der Coca-Cola Company ein. Der Sirup wird wiederum in den Läden der L verkauft. L sucht derweil wegen der hohen Kosten des Coca-Cola-Sirups nach einem alternativen Sirup-Anbieter. In dieser Situation erlangt der Vorstand (V) der L die Information, dass die National Pepsi-Cola Company insolvent ist. Ihm wird daher in seiner Eigenschaft als Vorstand der L die von Pepsi besessene Cola-Rezeptur sowie die zugehörige Marke zum Kauf angeboten. Anstatt jedoch Rezeptur und Marke für die L zu erwerben und damit die Geschäftschance für den Prinzipal zu nutzen, kauft V Rezeptur und Marke im eigenen Namen und für eigene Rechnung, ohne dass seine Vorstandskollegen hiervon Kenntnis erlangen. V errichtet daraufhin ein eigenes Unternehmen und verkauft der L im Weiteren Pepsi-Sirup mit Gewinn.<sup>44</sup>

306

### c. Lösungsstrategien: Überwachung und Interessenangleichung durch Vergütungsanreize

Wie lassen sich die genannten Anreizprobleme beheben oder zumindest lindern? Die ökonomische Theorie schlägt hier im Wesentlichen zwei Lösungsstrategien vor, die auch kombiniert werden können. Die erste Strategie ist die Überwachung, das sogenannte *monitoring*: Hierbei überprüft der Prinzipal die Tätigkeit des Agenten, indem er etwa Berichte anfordert, unangekündigte Kontrollvisiten macht oder einen anderen Agenten mit der Aufgabe betraut, den ersten Agenten zu überwachen. Beim *monitoring* geht es dem Prinzipal also darum, eine verborgene Aktivität (*hidden action*) oder eine verborgene Information (*hidden information*) aufzudecken. Hierdurch wird die Informationslücke zwischen Prinzipal und Agent geschlossen oder doch zumindest verkleinert. Die Gelegenheiten, zu denen der Agent zum eigenen Vorteil vom Handlungspfad im besten Interesse des Prinzipals abweichen kann, werden ebenso reduziert wie die Spannbreite

307

<sup>44</sup> Der Beispielsfall ist eine vereinfachte Version des der Entscheidung des Delaware Chancery Court in der Sache *Guth v. Loft, Inc.*, 5 A.2d 503 (Del. Ch. 1939) zugrundeliegenden Sachverhalts.

an Aktivitäten, die dem Agenten bei solchen Gelegenheiten zur Verfügung stehen.

308 Ein Beispiel für einen solchen *Monitoring*-Mechanismus finden wir im deutschen Aktienrecht. Dort wird dem Vorstand als Leitungs- und Verwaltungsorgan der AG ein Aufsichtsrat als Überwachungsorgan an die Seite gestellt (vgl. § 111 AktG), um sicherzustellen, dass der Vorstand auch tatsächlich im Interesse der Gesellschaft und ihrer Aktionäre handelt.

309 Die zweite Strategie zur Bekämpfung von Anreizproblemen eines Agenten, der für einen Prinzipal tätig wird, ist die Interessenangleichung, das sogenannte *alignment of interests*. Hier wird der Agent nicht durch die Verkleinerung der Informationslücke zwischen ihm und dem Prinzipal gleichsam „gezwungen“, im Interesse des Prinzipals zu handeln. Stattdessen wird der Agent dafür „belohnt“, dass er im besten Interesse des Prinzipals handelt. Auf diese Weise wird sein eigenes Interesse mit demjenigen des Prinzipals zur (weitgehenden) Deckung gebracht. Das gebräuchlichste Instrument, um diesen Effekt zu erzielen, ist die anreizsteuernde Bezahlung (*incentive pay*). Betrachten wir etwa wieder den Vorstand einer AG. Seine Vergütung besteht regelmäßig nicht nur aus einem Fixgehalt, sondern auch aus variablen Vergütungsbestandteilen. Letztere sind typischerweise vom Erfolg des Unternehmens abhängig. Zunächst hielt man Aktienoptionen für das Mittel der Wahl, um eine Angleichung der Interessen des Vorstands an diejenigen der Gesellschaft herbeizuführen. Heutzutage halten viele einen zu großen Anteil der variablen Vergütung und insbesondere die Zuteilung von Aktienoptionen für schädlich, weil dies einem risikofreudigen Geschäftsgebaren des Vorstands Vorschub leiste. Im Anschluss an die Finanzkrise von 2008 wurde eine lebhafte Debatte darüber geführt, wie Vorstände, insbesondere Bankvorstände inzentiviert werden können und sollten, damit sie im langfristigen Interesse ihrer Unternehmen handeln.

310 Beide Lösungsstrategien sind allerdings nicht umsonst zu haben. Die mit ihnen verbundenen Kosten nennt man **Agenturkosten** (*agency costs*).<sup>45</sup> Der Begriff erfasst zudem den ursprünglichen Wohlfahrtsverlust, der durch den Einsatz nicht optimal inzentivierter Agenten entsteht.

#### d. Staatliche Intervention durch Recht

311 Der Staat versucht diese Agenturkosten durch rechtliche Intervention zu reduzieren. Das Gesellschaftsrecht stellt etwa Überwachungsstrukturen und -mechanismen zur Verfügung, die immerhin eine gewisse Wirksam-

<sup>45</sup> So erhalten etwa Aufsichtsräte regelmäßig eine Vergütung für ihre Tätigkeit.

keit darin entfalten, die Geschäftsführer und Vorstände von gesellschaftsschädlichem Verhalten abzuhalten. Eine solche *Monitoring*-Struktur ist etwa der nach deutschem Aktienrecht zwingend einzurichtende Aufsichtsrat. Ein Beispiel für ein Mittel des *interest alignment* wäre etwa die (Schadensersatz-)Haftung für den Fall, dass der Agent von den Interessen des Prinzipals abweicht. So kennt das US-amerikanische Recht der *corporation* die sogenannte *corporate opportunities doctrine*, nach der die Mitglieder des Managements einer *corporation* dafür haften, wenn sie sich die Geschäftschancen der Gesellschaft aneignen. Ganz ähnliche Regelungen finden sich etwa auch im deutschen Gesellschaftsrecht. Dahinter steht der Gedanke, dass es sich für den Agenten von vorneherein nicht lohnt, entgegen dem Interesse des Prinzipals zu handeln, wenn er die diesem dadurch entstehenden Nachteile auszugleichen hat. Ein weiteres, aus dem Aktienrecht stammendes Beispiel für eine rechtliche Regelung zur Interessensangleichung von Agent und Prinzipal findet sich in § 87 AktG. Dort werden ausführliche Vorgaben zur Ausgestaltung der Vorstandsvergütung gemacht, die im Nachgang der Finanzkrise von 2008 noch einmal verschärft worden sind. Ziel ist es hierbei zu verhindern, dass ein Vergütungssystem falsche Anreize für das Vorstandsverhalten setzt, die sich letztlich zum Schaden der Gesellschaft und ihrer Aktionäre auswirken würden.

## 2. Langzeitverträge, Opportunismus und die Kostenabwägung der Parteien

### a. Optimierungshindernisse – Transaktionskosten und beschränkte Rationalität

Die vorstehend erörterten Anreizprobleme werden typischerweise noch einmal verstärkt, wenn die Parteien Langzeitverträge eingehen, die gelegentlich auch als „relationale Verträge“ bezeichnet werden. Wie der Name bereits andeutet, geht es hierbei um Verträge, die ihrem Inhalt nach darauf gerichtet sind, das Verhalten der Parteien für eine längere, manchmal unbestimmte Zeit zu regeln. Solche Verträge werden abgeschlossen, wenn ihr Gegenstand erhebliche spezifische Investitionen sowie fortlaufende Transaktionen der Beteiligten erforderlich macht. Der von solchen Verträgen geregelte Gegenstand ist typischerweise komplex. Diese Komplexität, sowie die Fortschreibung der vertraglichen Beziehung in die ungewisse Zukunft, führen zu einem hohen Maß an Unsicherheit im Moment des Vertragsschlusses. Zusammenfassend lässt sich formulieren, dass solche relationalen Verträge abgeschlossen werden, wo die geregelten Transaktionen

312

(1) wiederholt vorgenommen werden, (2) spezifische Investitionen erfordern und (3) unter gesteigerter Unsicherheit durchgeführt werden.<sup>46</sup>

313 Aufgrund der beschränkten Rationalität der Parteien, also ihrer beschränkten Fähigkeit, Informationen aufzunehmen und zu verarbeiten,<sup>47</sup> weisen solche (Langzeit-)Verträge gezwungenermaßen ganz erhebliche Lücken auf: Selbst wenn die Parteien erkennen, dass sie den Vertrag wegen dieser Lückenhaftigkeit während dessen Laufzeit immer wieder „nachbessern“ müssen, indem sie dessen Regelungen später ändern und weiter präzisieren, werden sie zum Zeitpunkt des Vertragsschlusses (*ex ante*) auf entsprechende Regelungsbemühungen verzichten, weil sie mit prohibitiv hohen Kosten einhergehen. Dies führt im Ergebnis dazu, dass ein solcher Langzeitvertrag auch für wichtige Aspekte des konkreten Arrangements häufig kein präzises und elaboriertes Pflichtenprogramm vorhält.

314 Vor dem Hintergrund dieser notwendigen Unvollständigkeit des Vertrages hegen die Parteien die Erwartung, dass sie die im Vertrag ursprünglich vorgenommene Risikozuweisung später im Wege von Nachverhandlungen an die geänderten Umstände anpassen werden, wenn dies erforderlich wird. Die Beendigung des Vertragsverhältnisses ist demgegenüber typischerweise kein gangbarer Weg, weil sie die spezifischen Investitionen der Parteien in das gemeinsame Projekt entwerten würde. Dies gilt jedenfalls, solange sich diese Investitionen noch nicht amortisiert haben. Aufgrund dieser Zwänge sind die Parteien in dem Vertragsverhältnis „eingeschlossen“ (*Lock-in-Effekt*).

## b. Die Gefahr – Ex-post-Opportunismus

315 Die beschriebenen Eigenschaften von Langzeitverträgen wären sehr viel weniger problematisch, wenn die Parteien sich auf eine Regel einigen würden, mit der sie sich dazu verpflichten, relevant werdende Lücken im Vertragswerk zur gegebenen Zeit in einer Weise zu füllen, die zu einer Maximierung des gemeinsamen Nutzens führt.<sup>48</sup> Dieser effiziente Handlungs-

<sup>46</sup> So *Williamson*, Transaction-Cost Economics: The Governance of Contractual Relations, *Journal of Law & Economics* 22 (1979), 233 (259): „[Relational contracts are concluded] where transactions (1) are recurrent, (2) entail idiosyncratic investment, and (3) are executed under greater uncertainty.“

<sup>47</sup> Vgl. Rz. 295 ff.

<sup>48</sup> Vgl. *Williamson*, The Economic Institutions of Capitalism: Firms, Markets, Relational Contracting, 1985, S. 48: „[P]roblems during contract execution could be avoided by *ex ante* insistence upon a general clause of the following kind: I agree candidly to disclose all relevant information and thereafter to propose and cooperate in joint profit-maximizing courses of action during the contract execution interval, the

pfad der Kooperation wird jedoch durch opportunistisches Verhalten der Parteien bedroht. *Williamson* definiert diesen **Opportunismus** als „Eigennutzstreben mit Arglist“. <sup>49</sup> Ein solches Verhalten führt dazu, dass die betreffende Partei ein Verhalten an den Tag legt, das für sie selbst zwar vorteilhaft ist, jedoch Ergebnisse vereitelt, die unter dem Gesichtspunkt der gemeinsamen Wohlfahrt vorzugswürdig wären. Da kooperatives Verhalten während der Vertragsdurchführung grundsätzlich im gegenseitigen Interesse der Parteien ist, tritt opportunistisches Verhalten nur dort auf, wo die hieraus kurzfristig erzielbaren Vorteile den langfristigen Nutzen der Kooperation überwiegen. Ob und wann ein solcher *high-value opportunism* <sup>50</sup> auftritt, hängt dann davon ab, mit welcher Diskontierungsrate die Parteien in der Zukunft liegende Profite abzinsen. <sup>51</sup>

### c. Die Abwägung – Transaktionskosten versus Governance-Kosten

Nach alledem stehen die Parteien eines Langzeitvertrags mithin vor einer 316  
zweifachen Aufgabe: Zum einen sind sie schon aufgrund ihrer beschränkten Rationalität bestrebt, die Transaktionskosten vor und bei Vertragsschluss (*ex ante*) möglichst klein zu halten. Zum anderen sind sie bemüht, Sicherungsvorkehrungen gegen opportunistisches Verhalten im Stadium der Vertragsdurchführung (*ex post*) vorzunehmen. Bei der Verfolgung dieser beiden Ziele haben sie jedoch eine Abwägung zu treffen. Denn die Reduzierung von Transaktionskosten *ex ante* durch den Verzicht auf eine detaillierte Regelung bestimmter Aspekte des Vertragsverhältnisses erhöht typischerweise die Opportunismusgefahr bzw. die *ex post* anfallenden Governance-Kosten zur Eindämmung dieser Gefahr.

### d. Das Recht – Teil der Lösung oder Teil des Problems?

Welche Rolle spielt nun das Recht für dieses Abwägungsproblem bei Lang- 317  
zeitverträgen? Im Idealfall kann das Recht dabei helfen, die Summe aus den Kosten des Opportunismus und den Kosten zur Vermeidung von Opportunismus, worunter sowohl die *Ex-ante*-Transaktionskosten als auch die *Ex-post*-Governance-Kosten fallen, zu verringern. Das dispositive

---

benefits of which gains will be divided without dispute according to the sharing ration herein provided.“

<sup>49</sup> Ebd., S. 47: „self-interest seeking with guile“.

<sup>50</sup> So der von *Posner*, A Theory of Contract Law under Conditions of Radical Judicial Error, Northwestern University Law Review 94 (2000), 749 (761) gewählte Begriff.

<sup>51</sup> Dazu auch Rz. 539 ff., 562.

(Gesetzes-)Recht (sogenannte *default rules*) kann hierzu auf dreierlei Weise beitragen: Zum Ersten informiert es die Parteien über die wesentlichen (und damit regelungswürdigen) Fragestellungen und Probleme des konkret in Rede stehenden vertraglichen Arrangements (Informationsfunktion). Zum Zweiten entlastet es die Parteien von der Bürde, jeden einzelnen Aspekt des Vertragsverhältnisses selbst zu regeln, so dass sie sich auf die für sie entscheidenden Punkte konzentrieren können (Entlastungsfunktion). Zum Dritten – und auf das Engste mit der Entlastungsfunktion verknüpft – füllt das dispositive Gesetzesrecht diejenigen Lücken aus, welche nach dem Vertragsschluss der Parteien im Vertragswerk verblieben sind (Lückenfüllungsfunktion).

318 Donald J. Smythe schreibt dem Vertragsrecht (mindestens) fünf positive Effekte zu, welche die Wahrscheinlichkeit schädlichen *Ex-post*-Opportunismus in Langzeitvertragsbeziehungen verringern können: Es kann (1) die Lebensdauer der Vertragsbeziehung erhöhen, (2) die Kooperationsbereitschaft der Parteien verbessern, (3) zur Steigerung der Investitionen in das Vertragsverhältnis beitragen, (4) die Unkosten für Schlichtungs- und Vermittlungsverfahren senken und (5) den Umfang von Transaktionen verringern, die im Rahmen weniger effizienter Governance-Strukturen vorgenommen werden.<sup>52</sup>

319 Bis hierher klingt die rechtliche Intervention in langfristig angelegte vertragliche Arrangements wie eine Erfolgsstory. Das Recht kann aus ökonomischer Warte aber auch selbst zum Teil des Problems werden. So weiß man inzwischen, dass rechtliche Regelungen und Strukturen, die eigentlich dazu bestimmt sind, eine Vertragspartei vor dem opportunistischen Verhalten der anderen Partei zu schützen, unter ungünstigen Umständen von der hierdurch geschützten Partei ihrerseits für eine andere Form von Opportunismus missbraucht werden können. Diese Form des Opportunismus lässt sich als eine Ausprägung des *Moral-hazard*-Problems<sup>53</sup> begreifen: Die schützende Regelung kann wie eine Versicherung wirken, die verhindert, dass die geschützte Partei die nachteiligen Konsequenzen ihres unkooperativen Verhaltens (*defection*) zu spüren bekommt. Dies senkt die Anreize der betreffenden Partei, das kooperative Gleichgewicht in der Vertragsbeziehung aufrecht zu erhalten. Wenn nun die Parteien ihren Zwist vor die Gerichte tragen, besteht zudem die Gefahr, dass das Gericht das Verhalten der Parteien fehldeutet. So mag der Richter ein Verhalten als

<sup>52</sup> Smythe, Bounded Rationality, the Doctrine of Impracticability, and the Governance of Relational Contracts, Southern California Interdisciplinary Law Journal 13 (2004), 227 (267).

<sup>53</sup> Zum Begriff des *moral hazard* bereits Rz. 302.

opportunistisch einstufen, obwohl es tatsächlich eine angemessene Vergeltungsmaßnahme ist, welche die illoyal gewordene Partei wieder auf den Kooperationspfad zurückführen soll, und umgekehrt.

Wegen dieser Gefahren dysfunktionalen und fehlerhaft angewendeten Rechts sprechen sich einige Rechtswissenschaftler dafür aus, die Rolle des Rechts als Mittel zur Stabilisierung wohlfahrtsfördernder Kooperation in Langzeitvertragsverhältnissen (zugunsten sozialer Normen) zurückzudrängen. Das Recht – so die Idee – sollte nur noch eingesetzt werden, um in Fällen von Opportunismus erheblichen Ausmaßes (*large-scale opportunism*) „hart durchzugreifen“.

#### e. Anwendungsbeispiel – Minderheitenschutz in der geschlossenen Kapitalgesellschaft

Ein bekannter Anwendungsfall, in dem das Recht zur Eindämmung von *Ex-post*-Opportunismus in rechtlichen Langzeitbeziehungen eingesetzt wird, ist der Schutz der Minderheitsgesellschafter in der geschlossenen Kapitalgesellschaft. Den betreffenden Regeln liegt das folgende Problem zugrunde: Geschlossene Kapitalgesellschaften, etwa die deutsche GmbH, teilen die beschriebenen Eigenschaften von Langzeitverträgen. Zur Erleichterung von Entscheidungen in der Gesellschafterversammlung gilt das Mehrheitsprinzip als gesetzlicher Regelfall. Dieses Prinzip ist freilich für die Gesellschafter in der Minderheitenposition mit Gefahren verbunden. So mag die Mehrheit versucht sein, den wohlfahrtsmaximierenden Kooperationspfad zu verlassen, um einen opportunistischen Kurs einzuschlagen, der sie gegenüber der Minderheit deutlich bevorteilt. Da es keinen liquiden Markt für die Gesellschaftsanteile geschlossener Kapitalgesellschaften gibt, ist der Verkauf der Beteiligung für die benachteiligte Minderheit regelmäßig keine realistische Option. Die Minderheit ist damit im obigen Sinne<sup>54</sup> in die Gesellschaft „eingeschlossen“ (*locked in*).

Das Gesellschaftsrecht kommt der bedrängten Minderheit mit verschiedenen Schutzmaßnahmen zu Hilfe: Das englische *company law* räumt einem Gesellschafter etwa in s. 994 des Companies Act 2006 die Möglichkeit ein, sich an die Gerichte zu wenden, wenn die Angelegenheiten der Gesellschaft in einer Weise geführt werden oder geführt worden sind, die den Interessen eines Teils der Gesellschafter, zumindest aber des klageführenden Gesellschafter in unangemessener Weise schadet (sogenanntes *Unfair-prejudice*-Verfahren). Ein weiteres Beispiel stammt aus dem deutschen GmbH-Recht. Danach ist ein Gesellschafter berechtigt, aus „wich-

<sup>54</sup> Vgl. Rz. 314.



tigem Grund“ aus der Gesellschaft auszutreten. Ein solcher wichtiger Grund liegt auch vor, wenn die Beziehung zu den übrigen Gesellschaftern zerrüttet, das Vertrauensverhältnis zwischen ihnen zerstört ist. Tritt der Gesellschafter nun aus, erlangt er einen Anspruch auf eine Abfindung in Höhe des vollen Werts der aufgegebenen Beteiligung. Problematischer ist es dagegen, dass das GmbH-Recht einen solchen Anspruch auf „volle“ Abfindung auch solchen Gesellschaftern zuspricht, die aus wichtigem Grund aus der Gesellschaft ausgeschlossen worden sind. Denn diese Regelung kann insofern „nach hinten losgehen“, als sie auch den Gesellschafter versichert, der selbst den Kooperationspfad verlassen hat (*moral hazard*). Aus diesem Grund sehen die Gesellschafter in ihrer Satzung häufig eine Abfindungsbeschränkung vor, die kooperatives Verhalten fördern soll.

## VI. „Verteilungsgerechtigkeit“ durch Vertragsrecht?

- 323 Die ökonomische Betrachtung von Verträgen und von Vertragsrecht hat als normatives Ziel die Maximierung der Gesamtwohlfahrt des in Bezug genommenen Gemeinwesens (*social welfare*) vor Augen. Im Vertragskontext fällt die Maximierung dieser sozialen Wohlfahrt häufig mit der Maximierung der Gesamtwohlfahrt der (prospektiven) Vertragsparteien zusammen. Demgegenüber ist es für Vertragstheoretiker per se ohne Bedeutung, wie die im Zuge der Vertragsdurchführung erlangten Wohlfahrtsgewinne zwischen den Parteien aufgeteilt werden. Diese weitgehende Ignoranz gegenüber Fragen der Verteilung muss den Juristen zunächst überraschen, spielt doch der Gesichtspunkt der „Verteilungsgerechtigkeit“ in der rechtswissenschaftlichen Diskussion über das Vertragsrecht eine prominente Rolle. Dieser Befund wirft die Frage auf: Ist das weit verbreitete Außersichtlassen von Verteilungsfragen ein fundamentaler Webfehler der ökonomischen Betrachtung von Vertragsdesign und Vertragsrecht? Oder findet sich eine überzeugende Begründung dafür, warum Ökonomen sich stattdessen auf Fragen der Wohlfahrtsmaximierung konzentrieren?
- 324 Für eine Antwort muss man wissen, dass Ökonomen ihre Augen keineswegs davor verschließen, dass die Umverteilung von Ressourcen durch staatliche Intervention notwendig ist.<sup>55</sup> Jedoch sind sich jedenfalls die Rechtsökonomen untereinander mehrheitlich einig, dass das Privatrecht im Allgemeinen ein vergleichsweise schwerfälliges, unzuverlässiges und in-

<sup>55</sup> Dazu Rz. 47 ff., 54, 96, etwa mit Blick auf die verfassungsrechtlichen Schranken der ausschließlichen Verfolgung des Normativziels Wohlfahrtsmaximierung.

effizientes Instrument zur Umverteilung von Ressourcen ist. In Bezug auf das Vertragsrecht im Speziellen benennen sie einige weitere, bedenkenswerte Argumente für diesen Standpunkt, die der Jurist zumindest kennen sollte, ohne sie sich zwangsläufig zu eigen machen zu müssen: Wird eine Partei durch das Recht mit einer nachteiligen Vertragsbestimmung belastet, etwa ein Unternehmer, der einen Vertrag mit einem Verbraucher abschließt, dann reagiert sie hierauf regelmäßig mit einer Preiserhöhung für das angebotene Produkt oder die angebotene Dienstleistung. Im Ergebnis muss also der Verbraucher, oder allgemeiner: die aus Sicht des Gesetzgebers von der Regelung profitierende Partei, die Kosten der rechtlichen Intervention selbst tragen. Eine Umverteilung zwischen Unternehmer und Verbraucher findet in diesen Fällen überhaupt nicht statt. Zur Veranschaulichung betrachten wir erneut Art. 3 und 5 der Verbrauchsgüterkaufrichtlinie, die eine zwingende Haftung des Unternehmer-Verkäufers vorsehen, wenn das verkaufte Verbrauchsgut zum Zeitpunkt der Lieferung nicht vertragsgemäß war und die Vertragswidrigkeit binnen zwei Jahren nach der Lieferung offenbar wird. Diese Haftung begründet zusätzliche Kosten für den Verkäufer. Er wird daher die erwarteten Kosten der Haftung bei der Berechnung des Produktpreises berücksichtigen. Der Verbraucher ist nunmehr gegen die Vertragswidrigkeit des Gutes durch die Haftung versichert. Unter der Annahme von Risikoaversion entspricht dies dem Interesse der Verbraucher. Sie erhalten diese Versicherung aber nicht umsonst!

Nehmen wir nun an, die durch das zwingende Vertragsrecht belastete Vertragspartei könne die hiermit verbundenen Kosten nicht oder zumindest nicht vollständig auf die Gegenpartei abwälzen. In diesem Fall verbleiben die Kosten also bei derjenigen Partei, deren Belastung der Gesetzgeber zugunsten der Gegenpartei auch beabsichtigt hat. Aber nutzt diese Belastung der einen Partei, etwa eines Unternehmers, der Gegenpartei, etwa einem Verbraucher, auch tatsächlich? Hierzu hat man im rechtsökonomischen Schrifttum ganz richtig angemerkt, dass eine solche Vertragsbestimmung der begünstigten Partei offensichtlich nicht so wichtig ist, jedenfalls nicht wichtig genug, um bereitwillig die damit verbundenen Kosten (vollständig) zu übernehmen. *Craswell* hat diesen Gedanken mit Blick auf Verbraucherverträge verallgemeinert: Auch wenn es zunächst kontraintuitiv anmutet, nützen doch diejenigen Rechtsregeln, deren Kosten der Unternehmer weitestgehend auf die Verbraucher abwälzen kann, diesen am meisten.<sup>56</sup> 325

<sup>56</sup> *Craswell*, Passing On the Costs of Legal Rules: Efficiency and Distribution in Buyer-Seller Relationships, *Stanford Law Review* 43 (1991), 361 (372): „Paradoxical

- 326 Einige Literaturstimmen sind allerdings der Meinung, dass der Unwille der begünstigten Vertragsseite für die betreffende Vertragsbestimmung zu bezahlen, gerade der Grund dafür sei, warum der Gesetzgeber diesen Vertragsinhalt zwingend anordnet. Allerdings lässt sich dies kaum mit einem Verweis auf den Gesichtspunkt der „Verteilungsgerechtigkeit“ erklären. Entscheidend scheint vielmehr die Annahme des Gesetzgebers zu sein, dass der Unwille, für den vertraglichen Schutz bzw. die vertragliche Begünstigung zu bezahlen, auf einer fehlerhaften Entscheidungsfindung, d. h. auf den kognitiven Beschränkungen der Betroffenen beruht.<sup>57</sup> Würde der Verbraucher die Bedeutung der ihn begünstigenden Vorschrift nur richtig begreifen – so die Überlegung –, dann würde er hierfür auch bezahlen.
- 327 Zwingende Vertragsrechtsbestimmungen können allerdings Umverteilungseffekte ganz anderer Art zeitigen: Nehmen wir an, die Präferenzen der Nachfrageseite (der Verbraucher) seien in Bezug auf die betreffende Vertragsbestimmung nicht einheitlich, also heterogen. In der Konsequenz würden einige Verbraucher von der Regelung profitieren, während sie für andere allein mit Kosten verbunden wäre. In einem solchen Szenario findet eine Umverteilung statt, und zwar nicht etwa zwischen Unternehmern und Verbrauchern, sondern zwischen den verschiedenen Verbrauchergruppen. Betrachten wir zu Illustrationszwecken noch einmal die Regelung in Art. 3 und 5 der Verbrauchsgüterkaufrichtlinie. Die dort angeordnete zwingende Haftung wirkt wie eine standardisierte Pflichtversicherung der Verbraucher gegen das Risiko der Vertragswidrigkeit der verkauften Verbrauchsgüter. Der Verkäufer ist also gezwungen hinsichtlich der Haftung einen *Pooling*-Vertrag<sup>58</sup> anzubieten. Dies kann dann wiederum die Quersubventionierung von „Hochrisiko“-Verbrauchern durch „Niedrigrisiko“-Verbraucher zur Folge haben (Umverteilungseffekt). Einige „Niedrigrisiko“-Verbraucher mögen unter diesen Umständen gar von einem Kauf Abstand nehmen, weil ihnen das Produkt (inkl. Standardversicherung) zu teuer ist. Dies führt zu einem Wohlfahrtsverlust.<sup>59</sup>

---

as it may seem, the rules whose costs are most heavily passed on are also the rules that will benefit consumers the most.“

<sup>57</sup> Vgl. bereits Rz. 296.

<sup>58</sup> Zum Unterschied zwischen *Pooling*-Verträgen und Selbstselektion herbeiführenden *Separating*-Verträgen vgl. Rz. 285 ff.

<sup>59</sup> Ebd.

## § 6 – Public und Social Choice Theorie

**Literatur:** *K.J. Arrow*, Social Choice and Individual Values, 1951; *A. Downs*, An Economic Theory of Democracy, 1957; *A. Hindmoor*, Rational Choice, 2006; *D. C. Mueller*, Public Choice III, 2003; *H. Nurmi*, Voting Paradoxes and How to Deal with Them, 1999; *N. Petersen*, The German Constitutional Court and legislative capture, International Journal of Constitutional Law 12 (2014), 650–669; *H. Daudt/D. W. Rae*, The Ostrogorski Paradox: A Peculiarity of Compound Majority Decision, European Journal of Political Research 4 (1976), 391–398; *W. Riker*, Liberalism Against Populism: A Confrontation Between the Theory of Democracy and the Theory of Social Choice, 1982; *J. Schumpeter*, Capitalism, Socialism and Democracy, 1942; *A. Sen*, The Possibility of Social Choice, American Economic Review 89 (1999), 349–378; *C. R. Sunstein*, Interest Groups in American Public Law, Stanford Law Review 38 (1985), 29–87; *E. V. Towfigh*, Das Parteien-Paradox. Ein Beitrag zur Bestimmung des Verhältnisses von Demokratie und Parteien, 2015; *G. Tullock*, The Politics of Bureaucracy, 1965.

### I. Ökonomik und Staatswissenschaft

Die Ökonomie versteht sich als umfassende Sozialwissenschaft, die sich nicht nur mit wirtschaftlichen Zusammenhängen, sondern grundsätzlich mit der Erklärung aller Lebensbereiche befasst. Insofern erstaunt es wenig, dass die ökonomische Theorie auch vor der Politik nicht Halt gemacht hat. Die ökonomische Politiktheorie wird oft auch als „Neue Politische Ökonomie“ bezeichnet, und hat über den Weg der Staats-, Verfassungs- und Verwaltungstheorie vor allem in den USA auch Eingang in die Rechtswissenschaften gefunden. Die Neue Politische Ökonomie nimmt Konflikte individueller und kollektiver Rationalität bei politischen Akteuren – Wählern, Politikern, Staatsbediensteten, Verwaltungen, Parteien, Interessenverbänden und so weiter – in den Blick. Ihr Erkenntnisinteresse ist in erster Linie deskriptiv-analytisch. In unserem Zusammenhang bedeutet das, die Rationaltheorie vor allem als eine weitere Perspektive zu verstehen, um über Politik und Verwaltung nachzudenken. Diese bricht mit einem „roman-

328

tischen“ Politikverständnis,<sup>1</sup> das bis in die 1950er-Jahre vorherrschte. Bis dahin behandelten Ökonomen und Politikwissenschaftler die Akteure auf der politischen Bühne anders als solche in Märkten. In den Märkten waren die Akteure dadurch charakterisiert, dass sie als Eigennutz maximierend und mit unvollständiger Kenntnis vom Markt und den Marktbedingungen modelliert wurden. In der Politik dagegen galten dieselben Individuen als ausschließlich am öffentlichen Interesse (Gemeinwohl) orientiert und allwissend.

- 329 Eine der grundlegenden Fragen der politischen Theorie ist jene nach der Begründung des Staates und der Definition seiner Aufgaben. Eine prominente Begründung der klassischen politischen Ökonomie war, dass der Staat überall dort tätig werden müsse, wo die Märkte bei der Erfüllung öffentlicher Aufgaben versagen. Das sei in den drei Domänen der „geborenen“ Staatsaufgaben der Fall: bei der Allokation von Gütern, bei ihrer Distribution und bei der Wahrung der gesamtwirtschaftlichen Stabilität.<sup>2</sup> Bei der Güterallokation sei der Staat etwa im Bereich der öffentlichen Güter zur Intervention berufen, weil diese Güter in einem Markt nicht zur Verfügung gestellt würden,<sup>3</sup> für das Funktionieren und die Wohlfahrt einer Gesellschaft aber unentbehrlich seien. Gleiches gelte in Bereichen, in denen Transaktionskosten anfallen, die zwar den individuellen Nutzen übersteigen, die aber positive **externe Effekte**<sup>4</sup> haben; auch solche Güter würden unter einem reinen Marktregime nicht hergestellt. Darüber hinaus bedürfe es der Korrektur negativer externer Effekte, der Gewährleistung des Funktionierens der Märkte – durch ihren Schutz etwa vor Wettbewerbsverzerrungen – und gegebenenfalls des Ausgleichs von Informationsasymmetrien; alles seien Aufgaben, die der Markt nicht wahrnehmen könne.<sup>5</sup> Offensichtlich sei auch, dass aus ordnungspolitischen Gründen bisweilen die Notwendigkeit bestehen könne, die vom Markt erreichte Distribution zu korrigieren, um eine ausgeglichene Einkommensverteilung zu erreichen. Schließlich sei der Staat der einzige Akteur, der gesamtwirtschaftliche Stabilität gewährleisten könne. Der Markt selbst bietet keine Möglichkeit, Gesamtangebot und Gesamtnachfrage im Gleichge-

<sup>1</sup> Als „romantisch“ bezeichneten Vertreter der aufkommenden *Public Choice Theory* seinerzeit das klassische Bild der Wissenschaft von den politischen Akteuren, insbesondere von „Wählern“, „Politikern“ und „Bürokraten“; vgl. etwa *Tullock, The Politics of Bureaucracy*, 1965.

<sup>2</sup> *Musgrave, The Theory of Public Finance*, 1959.

<sup>3</sup> Dazu Rz. 221 ff.

<sup>4</sup> Vgl. Rz. 161 f.

<sup>5</sup> Eine gute Übersicht findet sich bei *Hindmoor, Rational Choice*, 1991, S. 132 f.

wicht zu halten. Die hieraus resultierenden zyklischen Bewegungen der Volkswirtschaften bedürften einer Steuerung durch den Staat, etwa im Bereich der Beschäftigungspolitik.

Dieser Argumentation hielten Vertreter der Entscheidungstheorie entgegen, dass die Wohlfahrtsökonomie lediglich ein Marktversagen aufgezeigt habe, aber nicht darlege, dass der Staat willens und in der Lage sei, dieses Versagen zu korrigieren. Es sei inkonsequent, wenn die klassische politische Ökonomie zwar von unperfekten Märkten, aber vom perfekten Staat ausgehe. Wenn es ein **Marktversagen** gebe, dann sei nicht auszuschließen, dass es auch ein **Staatsversagen** gebe. Nur wenn das Marktversagen größeren Schaden anrichtet als das Staatsversagen, gebe es Grund dafür, den Staat intervenieren zu lassen.

330

Damit war die Neue Politische Ökonomie geboren. Aus Ökonomik und Entscheidungstheorie wurde ein theoretischer Rahmen entwickelt, der auf den politischen Prozess angewendet werden konnte. Im Kern geht es dabei immer um die Fragen, wie Entscheidungen einzelner rationaler Akteure in Gemeinwohlfragen sich auf das Gemeinwohl auswirken, wie sich negative Auswirkungen erklären und möglichst positive Wirkungen sich sicherstellen lassen (*Public Choice Theory*), und wie die zahlreichen Einzelinteressen am sinnvollsten zu einer Kollektiventscheidung aggregiert werden können (*Social Choice Theory*). Das Ganze wird vor dem Hintergrund reflektiert, dass Marktmechanismen, die zu einer Offenbarung der wahren Präferenzen der Marktteilnehmer führen, ebenso fehlen wie die Möglichkeit, durch individuelle Gewinnanreize das Eigennutzstreben der handelnden Individuen mit übergeordneten Interessen in Einklang zu bringen: In privatwirtschaftlich organisierten Unternehmungen ist dies durch Vertragsgestaltung möglich, aber in der Politik, wie sie von der Neuen Politischen Ökonomie modelliert wird, hat kein Akteur einen Anreiz gegenzusteuern, wenn das Eigennutzstreben zulasten des Gemeinwohls geht.

331

## II. Grundlegende Annahmen der *Public Choice Theory*

Die *Public Choice Theory* trifft in ihrer Grundform drei grundlegende Annahmen:

332

- Erstens bestimmt der *Public Choice Theory* zufolge der politische Prozess, wie die **Ressourcenallokation** aussieht, nicht etwa „wohlwollende“ oder „allwissende“, gewählte Einzelne.
- Zweitens geht die *Public Choice Theory* davon aus, dass sich der politische Prozess am besten als **strategische Interaktion zwischen den betei-**

**lichten Gruppen** – in erster Linie Wähler, Politiker, Verwaltungen – verstehen lässt.

- Drittens ist der *Public Choice Theory* zufolge jeder Akteur bestrebt, seinen **individuellen Nutzen** zu maximieren. Das ist beim Wähler der uns schon vom *homo oeconomicus* bekannte (allgemeine) **persönliche Nutzen**; beim Archetyp des Politikers sind das der Theorie zufolge in erster Linie **Wählerstimmen** und beim Archetyp des Bürokraten sein **Budget**.

Aus diesen Annahmen zieht sie theoretische Schlussfolgerungen für das rationale Verhalten der drei „Klassen“ von Akteuren, mit denen sie das politische Leben und das im politischen Prozess zu beobachtende Verhalten zu erklären versucht.

## 1. Politiker

- 333 Die erste Kategorie von Akteuren im Rahmen der Neuen Politischen Ökonomie ist der „Politiker“. Unter ihn fallen in der repräsentativen Demokratie, die die Folie für den Großteil der entwickelten Theorien bildet, sowohl gewählte Repräsentanten als auch sich um eine Wahl bewerbende Kandidaten. Der Grundsatz der Nutzenmaximierung ist der *Public Choice Theory* zufolge auch auf Politiker anzuwenden – es scheint nicht überzeugend, denselben Menschen in seiner Rolle als Marktteilnehmer eigennutzenmaximierend zu modellieren (etwa den Politiker als Privatperson), ihn als Politiker aber als benevolenten Gemeinwohloptimierer zu verstehen. Wie auch sonst im Rationalmodell bedeutet dies aber nicht, dass Politiker nur ihren materiellen Nutzen maximieren würden. Sie können neben finanziellen auch vielfältige andere Anreize haben: Beispielsweise können sie die Welt verbessern wollen oder an Macht interessiert sein. Ganz gleich, welches Endziel die Politiker aber verfolgen: Sie müssen im Amt bleiben, um es erreichen zu können. Das bedeutet, dass Politiker in erster Linie die auf sie entfallenden **Wählerstimmen maximieren** müssen. Das tun sie der *Public Choice Theory* zufolge vereinfacht gesprochen, indem sie den Wählern dadurch zu gefallen versuchen, dass sie ihnen Vorteile versprechen, die zumindest die **wahrgenommenen** Kosten ihrer Amtsausübung für die Wähler überwiegen. Wie es sich auf die Politik auswirkt, wenn die Politiker neben den Wählerstimmen auch ihr eigenes materielles Wohl, ihr Prestige oder ihre Macht zu mehren suchen, versucht die **Rent-seeking-Literatur** zu erklären.<sup>6</sup> Im *Public-Choice-Modell* übt der Politiker seine Macht

<sup>6</sup> Lesenswert dazu *Mueller*, *Public Choice* III, 1979, Kap. 15.

dadurch aus, dass er der Verwaltung – und damit den Bürokraten – **Politikvorgaben** macht, **Aufsicht** und **Kontrolle** über sie übt und ihnen den Haushalt zuweist.

## 2. Wähler

Auch der Wähler versucht in erster Linie seinen eigenen Nutzen zu maximieren.<sup>7</sup> Das hat zur Folge, dass auch angenommen werden muss, dass die Wähler jene Kandidaten unterstützen, von denen sie sich den größten individuellen Nutzen versprechen. Wie im Markt ist auch für die Wähler die Beschaffung von Informationen – etwa über die Politiker und über die für die Wähler entscheidenden Sachthemen – mit Kosten verbunden (Such- und Informationskosten = **Transaktionskosten**). Das bedeutet auch, dass der Wähler auf die Wirksamkeit seiner Wahlentscheidung bedacht sein wird, da die Wahl selbst ihrerseits Kosten verursacht (etwa Informationsbeschaffung über politische Programme und Kandidaten, aber auch Fahrt zum Wahllokal usw.). Er wählt, damit sein Kandidat die Wahl gewinnt. Der „Profit“ aus der Wahlhandlung ergibt sich aus der Differenz des Erwartungsnutzens beim Sieg des einen oder des anderen Kandidaten, abzüglich der Kosten für die Wahl. Wichtig ist, dass den klassischen Modellen der *Public Choice Theory* zufolge die einzige Möglichkeit des Wählers, auf den politischen Prozesse Einfluss zu nehmen, seine Stimme bei Wahlen ist.

Die Wahrscheinlichkeit, dass eine beliebige einzelne Stimme für den Wahlausgang entscheidend ist, ist theoretisch wie empirisch äußerst gering: Dass die eigene Stimme entscheidend ist, kann nur dann angenommen werden, wenn alle anderen Stimmen genau gleich verteilt sind oder wenn der von einem selbst favorisierte Kandidat aufgrund einer einzigen fehlenden Stimme unterliegen würde, wenn man nicht wählen ginge. In der Literatur wird daher immer wieder darauf hingewiesen, es sei ebenso wahrscheinlich bei einem Autounfall auf dem Weg zum oder vom Wahllokal ums Leben zu kommen, wie dass die eigene Stimme den Wahlausgang entscheide.<sup>8</sup> Daraus folgt zweierlei: Da die Kosten der Informationsbeschaffung in aller Regel den Nutzen der Information für den Wähler überwiegen, ist es zum einen für die Wähler rational, sich zu den meisten Themen nicht, jedenfalls aber nicht umfassend zu informieren (**rationale Ignoranz**); die nicht eingeholten Informationen können dementsprechend auch keine Rolle bei der Wahlentscheidung spielen.

<sup>7</sup> Erstmals so *Downs*, *An Economic Theory of Democracy*, 1957, Kap. 11–14.

<sup>8</sup> Nachweise bei *Mueller*, *Public Choice III*, 1979, S. 305 (Anm. 4).



336 Zum anderen ist es rational, nicht wählen zu gehen:<sup>9</sup> Weil die mit der Wahl verbundenen Kosten höher sind als der erwartete Nutzen und weil angesichts von rationaler Ignoranz in aller Regel auch die Bestimmung des „richtigen“ Kandidaten mit erheblicher Unsicherheit verbunden ist. Dass dennoch so viele Bürger an Wahlen teilnehmen, wird als **Wahlparadox** (*voting paradox*) bezeichnet. Bedeutende Vertreter der Neuen Politischen Ökonomie halten das für irrationales Verhalten,<sup>10</sup> während andere nach (rationalen) Erklärungen suchen. Verschiedene Ansätze können dabei unterschieden werden: Etwa eine Modifikation des Kalküls des Wählers, mit dem Ergebnis, dass Wählen rational ist; eine Modifikation des Eigennutztheorems<sup>11</sup>, so dass auch der **soziale Nutzen** gewichtet wird; oder eine Modifikation der Rationalitätsannahme<sup>12</sup>, so dass auch die Wahl solcher Optionen rational ist, die nicht den Eigennutz im herkömmlichen Sinne maximieren (z. B. Minimierung der maximalen Reue, **Minimax-regret-Strategie**). Neben diesen gibt es noch eine Reihe weiterer Erklärungsversuche dafür, warum Menschen wählen gehen; die Frage ist bis heute nicht beantwortet.<sup>13</sup> Dagegen scheint empirische Evidenz tatsächlich zu belegen, dass die Erklärung, wie Menschen ihre Wahlentscheidung treffen – nämlich im Wesentlichen rational im Sinne der Neuen Politischen Ökonomie – zutrifft.<sup>14</sup> Wir werden darauf später, im Rahmen des **Medianwähler-Theorems**,<sup>15</sup> noch näher eingehen.

<sup>9</sup> Eine didaktisch ansprechend aufbereitete, wenngleich etwas oberflächliche Präsentation dieses Themas findet sich unter <http://www.youtube.com/watch?v=21uJU-ZuIcEo> (Metzler/Kurz, Tullock: Voting Schmoting).

<sup>10</sup> Vgl. Tullock, *Toward a Mathematics of Politics*, 1967.

<sup>11</sup> Tullock, *Toward a Mathematics of Politics*, 1967, S. 110 („taste for voting“).

<sup>12</sup> Vgl. Ledyard, *The Paradox of Voting and Candidate Competition: A General Equilibrium Analysis*, in: Hornwich/Quirk (Hg.), *Essays in Contemporary Fields of Economics. In Honor of Emanuel T. Weiler (1914–1979)*, 1981; ders., *The Pure Theory of Large Two-Candidate Elections*, *Public Choice* 44 (1984), 7 ff.; Palfrey/Rosenthal, *A Strategic Calculus of Voting*, *Public Choice* 43 (1983), 7 ff.; dies., *Voter Participation and Strategic Uncertainty*, *American Political Science Review* 79 (1985), 62 ff.; außerdem: Ferejohn/Fiorina, *The Paradox of Not Voting: A Decision Theoretic Analysis*, *American Political Science Review* 68 (1974), 525 ff.

<sup>13</sup> Empirische Untersuchungen finden sich bei Aldrich, *When is it rational to vote?*, in: Mueller (Hg.), *Perspectives on Public Choice*, 1997, 373 ff.

<sup>14</sup> Merrill/Grofman, *A Unified Theory of Voting*, 1999; vgl. Fiorina, *Voting Behaviour*, in: Mueller (Hg.), *Perspectives on Public Choice*, 1997, 391 ff.

<sup>15</sup> Hierzu Rz. 342 ff.

### 3. Bürokraten

Wir haben bislang das Modell des Politikers und des Wählers kennen gelernt – ersterer durch seine Abhängigkeit von Wählerstimmen als Marionette des letzteren, und damit mit Blick auf die Bereitstellung öffentlicher Güter und die Förderung des Gemeinwohls vor allem die „Nachfrageseite“ in den Blick genommen. Auf der anderen Seite stehen nun jene, die die von den Bürgern geforderten Politiken ausführen sollen. Auch sie versuchen dabei der *Public Choice Theory* zufolge ihren Eigennutzen zu maximieren. 337

In der Welt der *Public Choice Theory* werden Verwaltungen als außerhalb des Marktes agierende Organisationen betrachtet, die zur Umsetzung von Politiken den Wählern **Güter und Dienstleistungen** anbieten und dafür von den Politikern **Haushaltsmittel** zugewiesen bekommen. Da das Rationalmodell aber grundsätzlich auf das Handeln von Individuen abstellt,<sup>16</sup> muss auch die *Public Choice Theory* an individuelles Verhalten anknüpfen. Ihr ist daher eine Verwaltung im Sinne einer monolithischen politischen Einheit fremd, vielmehr ist die Verwaltung nur die Summe der sie bildenden Mitarbeiter, die mit der wenig schmeichelhaft klingenden Bezeichnung „Bürokraten“ betitelt werden. Sie geht auf den französischen Nationalökonom *Vincent de Gourmay* (1712–1759) zurück, der den Begriff als Gegenpol zu dem berühmteren Ausspruch „*Laissez faire, laissez passer, le monde va de lui-même*.“<sup>17</sup> konzipierte.<sup>18</sup> Seit jeher ist der Begriff negativ konnotiert, und seine beinahe durchgehende Verwendung in der *Public Choice Theory* lässt erahnen, welche Haltung jedenfalls die Begründer der *Public Choice Theory* gegenüber ihrem Forschungsgegenstand einnahmen. Da er sich aber in der Literatur als *terminus technicus* eingebürgert hat, soll er auch im Folgenden verwendet werden. 338

Für die Theorie wichtige Charakteristika der Verwaltung sind, dass diese **hierarchisch gegliedert** und **nicht auf Profit ausgerichtet** ist. Aus diesen beiden Eigenschaften ergeben sich auch die wichtigsten Unterschiede zur Organisation am Markt: Zum einen müssen aufgrund der Hierarchie alle Informationen zentral beim Vorgesetzten bzw. letztlich beim Verwaltungsleiter (*senior bureaucrat*) zusammenlaufen. Zum anderen sind, anders als am Markt, nicht alle Akteure auf das gleiche messbare Ziel – den Profit – ausgerichtet, so dass aufgrund **divergierender Interessen** unter den ver- 339

<sup>16</sup> Vgl. Rz. 69 ff.

<sup>17</sup> Zu Deutsch etwa: „Lass es geschehen, lass es vorüber gehen, die Welt dreht sich von selbst.“

<sup>18</sup> *De Gourmay*, Observations sur l'état de la Compagnie des Indes, 1759.

schiedenen Bürokraten **Zielkonflikte** auftreten können, die ihrerseits Wohlfahrtsverluste nach sich ziehen.

- 340 Was sind nun die Interessen der Bürokraten? Auch sie maximieren ihren Eigennutz, aber anders als bei Akteuren am Markt können sie nicht ihren Profit maximieren, da Verwaltungen *qua definitionem* nicht gewinnorientiert sind. Aber der Eigennutz kann bekanntlich auch in anderen Gütern liegen: Das kann, vor allem bei niederen Verwaltungsbediensteten, ein sicherer Arbeitsplatz sein; eine höhere Vergütung; eine bessere Ausstattung; ein Gewinn an Macht und Einflussnahmemöglichkeiten; öffentliche Anerkennung und Status; oder, bei höheren Verwaltungsbeamten, die möglichst einfache Führung der Verwaltung.<sup>19</sup> Zumindest bis auf das letzte<sup>20</sup> lassen sich all diese Ziele maximieren, indem der jeweils zugewiesene **Haushalt** maximiert wird. Ein zentrales Postulat der *Public Choice Theory* ist daher, dass Bürokraten bemüht sind, ihr **Budget zu maximieren**; wie sie dies tun, werden wir später<sup>21</sup> eingehender betrachten.

### III. Fehlanreize in repräsentativen Systemen

- 341 Die *Public Choice Theory* war schon ihres Entstehungszusammenhangs wegen in erster Linie ein Instrument, um das Versagen staatlicher Einrichtungen aufzuzeigen und zu erklären. Ihre besondere Stärke und Anziehungskraft verdankt sie intuitiv einleuchtenden, anreizorientierten Erklärungen von im politischen Leben (vor allem in repräsentativen Demokratien) allzu oft zu beobachtenden Problemen. Entlang der Hauptakteure der *Public Choice Theory* sollen im Folgenden drei prominente Beispiele für Erklärungsmodelle der Neuen Institutionenökonomie vorgestellt werden: Das **Medianwähler-Theorem** (Politiker/Parteien), kleine Wählergruppen mit **Sonderinteressen** (Wähler) und das Problem der **Budgetmaximierung** in der Verwaltung (Bürokraten).

<sup>19</sup> Vgl. etwa *Niskanen*, Bureaucracy and Representative Government, 1971, S. 38; *Downs*, Inside Bureaucracy, 1967, S. 81 ff.

<sup>20</sup> Auch hier präsentiert *Niskanen* aber ein subtiles Argument, das selbst für jene Bürokraten, die nach einer einfachen Verwaltungsführung streben, Budgetmaximierung als probates Mittel erscheinen lässt; vgl. a. a. O.

<sup>21</sup> Hierzu Rz. 361 ff.

## 1. Das Medianwähler-Theorem

Es gehört heute zu den Binsenweisheiten der Politik, dass Wahlen in der Mitte gewonnen werden. Allein mit Klientelpolitik kann keine Partei mehr die Mehrheit bei Wahlen erringen – sie muss vielmehr um den **Medianwähler** werben, den Wähler, der genauso viele Menschen links wie rechts des politischen Spektrums von sich weiß. Die theoretische Erklärung dieses Phänomens geht auf das **Medianwähler-Theorem** von *Anthony Downs* zurück.<sup>22</sup> Inspiriert wurde *Downs* dabei von einem Aufsatz des Ökonomen *Harold Hotelling*;<sup>23</sup> *Hotelling* beschäftigt sich in seinem Beitrag mit Wettbewerb in einem eindimensionalen Raum. Stellen wir uns ein Dorf vor, das im Wesentlichen an einer langen Hauptstraße entlang gebaut worden ist und in dem sich zwei Tankstellen ansiedeln möchten. Wo werden diese beiden sich ansiedeln, wenn man davon ausgeht, dass die Entfernung zur Tankstelle der Hauptfaktor für die Dorfbewohner ist, sich für eine der beiden zu entscheiden? Man könnte zunächst annehmen, dass sie das Dorf in zwei imaginäre Hälften teilen und sich dann jeweils in der Mitte einer der beiden Hälften ansiedeln. In der Realität findet man beide Tankstellen jedoch oft dicht beieinander in der Ortsmitte. Warum? Wenn sich eine der beiden Tankstellen in der Mitte der, sagen wir, östlichen Hälfte des Dorfes platzieren würde, dann wäre es für die zweite Tankstelle rational, sich direkt westlich davon zu platzieren. Sie würde dann nämlich alle Kunden abgreifen, die westlich von ihr wohnen, also drei Viertel des Dorfes, was ihr einen erheblichen Wettbewerbsvorteil einbrächte. Um dies zu vermeiden und dem Konkurrenten keinen Vorteil zu gewähren, werden beide Tankstellen in die Mitte des Dorfes ziehen.

### a. Das Modell

Kann man diese Konstellation auf die Politik übertragen? *Downs* argumentiert, dass dies möglich sei. Er orientiert sich dabei an der demokratietheoretischen Konzeption von *Joseph Schumpeter*. *Schumpeter* glaubt – ähnlich wie die Neue Politische Ökonomie –, dass Demokratie sich nicht wesentlich von einem wirtschaftlichen Markt unterscheide. Statt um Marktanteile konkurrierten politische Parteien um Wählerstimmen und richteten ihre Programme immer so aus, dass sie möglichst viele Wähler gewinnen können. Dies ist auch die zentrale Annahme in *Downs'* Theorie, jedoch nicht die einzige. Die weiteren Annahmen, die *Downs* trifft, sind,

<sup>22</sup> *Downs*, An Economic Theory of Democracy, 1957.

<sup>23</sup> *Hotelling*, Stability in Competition, Economic Journal 39 (1929), 41 ff.

dass der politische Wettbewerb in einem **eindimensionalen Spektrum** stattfindet, dass Parteien sich innerhalb dieses Spektrums **frei bewegen** können, dass es **lediglich zwei Parteien** gibt, dass Wähler immer für die Partei stimmen, die ihnen im politischen Spektrum **am nächsten** liegt, und schließlich, dass **Informationen vollständig** und **Wählerpräferenzen konstant** sind.

- 344 Legt man diese Annahmen zugrunde, dann kommt man zu dem Ergebnis, dass sich die Parteien immer beim Medianwähler treffen. Rückt nämlich eine Partei im politischen Spektrum ein wenig nach rechts oder links, dann kann die andere Partei sofort nachziehen und die Wähler nahe der Mitte einsammeln, die von der anderen Partei verlassen worden sind. Dabei ist es für die Theorie unerheblich, wie letztlich die Wählerverteilung aussieht. Bei einer Normalverteilung würden sich die Parteien auch tatsächlich in der Mitte des politischen Spektrums treffen. Haben wir dagegen ein Spektrum, bei dem sich auf einer Seite eine stärkere Häufung von Wählern befindet, dann rückt der Medianwähler nach rechts oder links, so dass die Parteien ebenfalls rechts oder links von der Mitte des Spektrums um den Medianwähler konkurrieren (s. Abbildung 6.1).

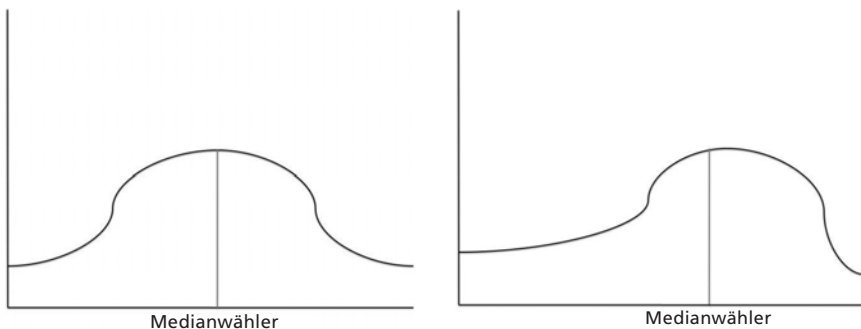


Abbildung 6.1: Position des Medianwählers bei unterschiedlicher Verteilung des politischen Spektrums

- 345 Nun drängt sich an dem Medianwähler-Theorem schnell Kritik auf, die sich zum einen auf die Annahmen, zum anderen auf das Ergebnis bezieht. Die Annahmen erscheinen unrealistisch. So ist das politische Spektrum in der Realität wohl kaum lediglich eindimensional und Politiker haben wahrscheinlich auch bestimmte politische Überzeugungen, die sie nicht im Streben nach einem Maximum an Wählerstimmen verkaufen würden. Zudem scheint das Modell mit der Annahme, dass lediglich zwei Parteien miteinander konkurrieren, stark auf den angelsächsischen Kontext zuge-

schnitten zu sein. Was das Ergebnis angeht, sieht es kaum besser aus: Zwar mag man beobachten, dass die großen Parteien in Deutschland sich in den letzten Jahren sehr stark angenähert haben. Dennoch gibt es immer noch merkliche Differenzen. Die Politik ist nicht austauschbar, sondern die Vorschläge folgen zumeist klassischen Mustern, deren Aufbau und Argumentation in der Tradition der Partei verwurzelt sind.

Allerdings nimmt ein Modell für sich nicht in Anspruch, dass es die Wirklichkeit exakt beschreiben würde. Ein Modell abstrahiert, indem es sich auf bestimmte Faktoren konzentriert, und ist daher notwendigerweise nicht immer vollständig exakt in seinen Vorhersagen. Daher sollte man ein Modell nie als die reine Wahrheit betrachten. Als Ausgangspunkt für Erklärungen hat es jedoch großen analytischen Wert. Dies gilt insbesondere auch dann, wenn wir im Folgenden einmal die Annahmen, die *Downs* getroffen hat, etwas abschwächen und schauen, wie sich dann die Vorhersagen des Modells verändern. 346

#### b. Abschwächen der Prämissen des Modells

Die zentrale Prämisse der **kompetitiven Demokratietheorie**, auf der auch *Downs'* Modell beruht, ist, dass Politiker grundsätzlich Wählerstimmen zu maximieren suchen. Ist diese Annahme realistisch? Wenn man sie strikt fasst und sagt, dass Politiker allein Wählerstimmen maximieren und keine anderweitige Motivation haben, tut man unseren gewählten Repräsentanten wohl Unrecht. Den meisten Politikern wird man unterstellen können, dass sie eine politische Überzeugung haben, bestimmte Ideale, die sie nicht im Austausch mit einem Zuwachs an Wählerstimmen aufgeben würden. Dies erklärt, warum die politischen Parteien sich nicht exakt beim Medianwähler treffen, sondern durchaus noch unterschiedliche Positionen vertreten. Aber ist die Prämisse deswegen unrealistisch? Wenn wir sie abschwächen und sagen, dass **Politiker auch Wählerstimmen maximieren**, scheinen wir der Realität schon sehr viel näher zu kommen. Wir können etwa beobachten, dass politische Programme keinesfalls Konstanten sind, sondern dass sie gesellschaftlichen Gegebenheiten angepasst werden, wenn eine Partei erfolglos ist. *Tony Blairs'* „New Labour“, mit der er seine Partei 1997 nach 18 Jahren Abstinenz von der Macht zurück an die Regierung führte, ist dafür ein gutes Beispiel. Dieser Kurswechsel, in gewisser Hinsicht vergleichbar mit der „Agenda 2010“ der deutschen Sozialdemokraten, war sicherlich nicht nur durch wandelnde politische Anschauungen innerhalb der Partei, sondern ebenso durch den Willen, die Macht zurückzuerobern, motiviert. 347

- 348 Daran schließt sich die zweite Prämisse an, dass politische Parteien ihre Position im Spektrum frei wählen können. Zwar sind Parteien in gewissem Maße flexibel, ihre Positionen zu ändern. Allerdings ist diese **Flexibilität eingeschränkt**. Dies liegt zum einen an der politischen Überzeugung der Parteieliten, vor allem aber an der politischen Überzeugung der Parteibasis. Jede Partei lebt von ihrer Basis, die die Ideale ihrer Partei vertritt und sich mit diesen identifiziert. Diese Basis steht zumeist entweder links oder rechts im politischen Spektrum. Die Partei darf ihre Basis nicht verprellen, weshalb sie bei der Gestaltung der Politik in gewisser Hinsicht einen begrenzten Spielraum hat, was ebenfalls erklärt, warum sich die Parteien oft nicht genau in der Mitte treffen.
- 349 Die weitere Annahme des Modells, dass das politische **Spektrum eindimensional** ist, ist seit langem relativ verbreitet. Klassischerweise wird in der Politik zwischen „rechts“ und „links“ unterschieden, wobei der Unterschied gradueller Art ist, jedes Individuum also von dem äußersten linken bis zum äußersten rechten Punkt im Spektrum jede beliebige Position einnehmen kann. Allerdings erscheint es schwierig, alle politischen Positionen einwandfrei auf einem Rechts-Links-Schema einzuordnen. So mag es etwa im linken politischen Spektrum Personen geben, die glauben, dass die bevorzugte Gesellschaftsordnung am besten durch einen autoritären Staat durchgesetzt werden kann – nicht umsonst waren die kommunistischen Regime nach dem ersten Weltkrieg autoritärer Natur. Andere dagegen mögen anarchistische Formen des Regierens für sehr viel aussichtsreicher halten. Eine ähnliche Form der Differenzierung zwischen eher autoritären und liberalen Systemen kann auch auf der rechten Seite des Spektrums vornehmen, so dass es plausibler erscheint, den politischen Raum als einen **mehrdimensionalen** zu betrachten. Möchte man ihn beispielsweise zweidimensional betrachten, könnte man auf der X-Achse zwischen rechts und links, auf der Y-Achse zwischen eher autoritär und eher liberal unterscheiden. Auch in einem solchen zweidimensionalen Raum kann man **Indifferenzkurven**<sup>24</sup> zeichnen und versuchen, die optimale Position einer Partei zu bestimmen.<sup>25</sup> Allerdings ist es bei allen Räumen, die mehr als eine Dimension haben, unmöglich ein stabiles Gleichgewicht zu finden. Es kann also zu jedem Punkt in dem politischen Raum ein weiterer Punkt gefunden werden, der eine größere Mehrheit der Wähler hinter sich weiß. Weil es aber einen dritten Punkt gibt, der besser ist als Punkt 2, aber schlechter als Punkt 1, ist es unmöglich, ein stabiles Gleichgewicht zu identifizieren.

<sup>24</sup> Hierzu bereits Rz. 106 ff.

<sup>25</sup> Dazu ausführlich *Hindmoor*, *Rational Choice*, 2006, S. 34–39.

Allerdings kann man eine unabgedeckte Menge identifizieren: den politischen Raum, in dem die Chance, eine Mehrheit hinter sich zu bringen, am größten ist. Diese unabgedeckte Menge findet sich zumeist in der Nähe des Schnittpunktes der jeweiligen Dimensionsmediane (s. Abbildung 6.2). Insofern ändert die bloße Addition weiterer Dimensionen zum politischen Spektrum noch nichts an den Aussagen des Medianwähler-Theorems. Dieses Erkenntnis zeigt uns auch, wie wir das Theorem für Mehrparteiensysteme nutzbar machen. Während Zwei-Parteien-Systeme grundsätzlich in Systemen mit Mehrheitswahlrecht, wie den USA oder – mit Abstrichen – Großbritannien üblich sind, finden wir in Staaten mit Verhältniswahlrecht grundsätzlich ein Mehrparteiensystem. In diesem Mehrparteiensystem besetzen die kleineren Parteien oft Positionen, die von den großen Volksparteien nicht abgedeckt werden. Die Grünen sind etwa nicht nur eine weitere Partei innerhalb des Rechts-Links-Spektrums, sondern sie greifen spezifische Belange innerhalb des mehrdimensionalen Politikraums, etwa den Umweltschutz, heraus, die bis zur Herausbildung der Grünen von anderen Parteien nicht besetzt worden waren.

350

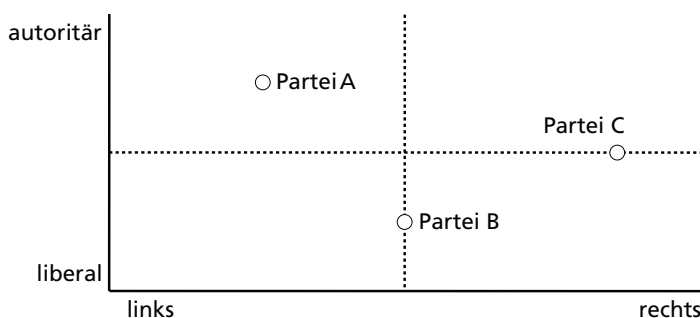


Abbildung 6.2: Wettbewerb der Parteien um den Medianwähler im zweidimensionalen Raum

Auch die Annahme, dass Wähler immer die Partei wählen, die ihnen politisch am nächsten steht, trifft nicht exakt die Realität. Für Wähler ist es wie gezeigt rational, sich nur begrenzt über Parteiprogramme zu informieren (rationale Ignoranz). Ein simpler Entscheidungsmechanismus ist daher, sich für die Partei zu entscheiden, die man immer wählt. Allerdings ist diese Trägheit begrenzt. Es gibt eine gewisse Anzahl mobiler Wähler und diese nimmt, glaubt man jüngeren Erhebungen in Westeuropa, eher zu als ab. Insofern ändert die Einschränkung dieser Prämisse nichts an den grundsätzlichen Aussagen des Modells. Bewegen sich die Parteien in Rich-

351



tung des Medianwählers, können sie damit ihre Wahl zwar nicht sicherstellen, aber doch zumindest ihre Chancen, gewählt zu werden, erhöhen.

352 Auch die Annahme der vollständigen Informationen ist nicht mehr als eine Modellannahme. Da Wahlentscheidungen nicht nur Entscheidungen über die Vergangenheit, sondern insbesondere auch über die Zukunft sind, besteht bei Wählern eine große Unsicherheit darüber, was für Folgen ihre Wahlentscheidung letztlich zeitigen wird. Um diese Unsicherheit zu verringern, sind Parteien darauf angewiesen, möglichst glaubwürdig zu sein und ihre Politik nicht ständig radikal zu ändern. Auch insofern liefert die Einschränkung der Annahme vollständiger Informationen eine Erklärung dafür, warum Parteien sich im politischen Spektrum nicht frei bewegen können, sondern einen begrenzten Bewegungsspielraum haben.

353 Die Prämisse, dass die **Präferenzen** (hier: der Wähler) **konstant** sind, gehört zu den Grundannahmen von *Rational Choice*, auch außerhalb der Neuen Politischen Ökonomie und des Medianwähler-Theorems.<sup>26</sup> Die Präferenzbildung von Individuen gehört immer noch zu einer der großen Unbekannten der sozialwissenschaftlichen Forschung. Immerhin scheint psychologische Forschung in den letzten Jahren darauf hinzudeuten, dass die Präferenzen von Individuen nicht konstant sind, sondern sich über die Zeit verändern.<sup>27</sup> Was bedeutet das für das Modell? Im Wesentlichen hat es zur Folge, dass politische Parteien einen größeren Handlungsspielraum haben. Sie müssen sich nicht allein an den vorgegebenen Präferenzen ihrer Wähler ausrichten, sondern können auch darauf hinwirken, diese Präferenzen selbst zu formen. Dies ist wahrscheinlich der schwerwiegendste Einwand gegen das Medianwähler-Theorem. Allerdings ist die Formung von Präferenzen ein komplexer und langwieriger Prozess, der sehr viel unsichereren Erfolg verspricht als die kurzfristige Ausrichtung an den politischen Anschauungen des Medianwählers. Da Politik jedoch im Wesentlichen kurzfristig ausgerichtet ist,<sup>28</sup> werden die Parteien ihre Strategien zumindest mischen. Die Orientierung am Medianwähler wird dabei einen entscheidenden Einfluss behalten.

### c. Fazit

354 Das Medianwähler-Theorem ist ein hervorragendes Beispiel für die Wirkungskraft ökonomischer Modelle. Ökonomische Modelle haben nicht

<sup>26</sup> Vgl. Rz. 66.

<sup>27</sup> Vgl. etwa *Slovic*, The construction of preference, *American Psychologist* 50 (1995), 364 ff.

<sup>28</sup> Hierzu sogleich § 6 IV.

den Anspruch, die Wirklichkeit in allen Einzelheiten zu beschreiben. Vielmehr müssen sie bestimmte Annahmen treffen, die notwendigerweise pauschalisieren. Allerdings muss man diese Annahmen nicht unbedingt als Restriktionen verstehen. Vielmehr kann es durchaus interessant sein, die Annahmen zu variieren und zu beobachten, wie sich die Aussagen des Modells mit geänderten Annahmen selbst verändern. Das Medianwähler-Theorem nimmt insofern für sich nicht in Anspruch, den politischen Prozess immer punktgenau vorherzusagen. Es dient jedoch als Erklärung für viele Phänomene, die wir in der Politik beobachten können. Dabei ist die Erklärungskraft bei Zweiparteiensystemen höher als etwa im repräsentativen System der Bundesrepublik Deutschland. Doch auch in letzterem hat es ein beträchtliches Erklärungspotential.

## 2. Sonderinteressen bei Wählern und Politikern – *rent seeking*<sup>29</sup>

Ein weiteres populäres Beispiel aus der Neuen Politischen Ökonomie ist der rationaltheoretische Blick auf die Mechanismen bei der Durchsetzung von Partikularinteressen – in der gesellschaftlichen Debatte oft als Lobbyismus bezeichnet. Das Grundproblem ist leicht erklärt: Kleine Gruppen mit besonders prononcierten Interessen machen sich die rationale Ignoranz der Mehrheit der Wähler und das Interesse der Politiker an der Stimmenmaximierung zunutze. Im Mittelpunkt dieses **Sonderinteressen-Effekts** steht eine bestimmte politische Entscheidung, die einen großen Nutzen (die sogenannte **politische Rente**) für eine kleine Gruppe von Wählern verspricht, aber mit einem externen Effekt einher geht, der geringe Kosten für eine sehr viel größere Gruppe anderer Wähler erzeugt. Die kleine Wählergruppe hat ein besonderes Interesse daran, sich zu organisieren und konzentriert zu handeln. Aus der Perspektive des Politikers dominiert das Thema die Wahlentscheidung der Gruppe; will er diese Gruppe als Wähler hinter sich bringen, so muss er ihre Forderung erfüllen.

Dagegen ist die große Gruppe von Wählern, die die Kosten dieser Entscheidung trägt, mit Blick auf das konkrete Vorhaben in der Regel rational desinteressiert, weil die Auswirkung auf den persönlichen Wohlstand unterhalb der Wahrnehmungsschwelle liegt. Aus diesem Grund haben Politiker starke Anreize, Sonderinteressen wahrzunehmen, selbst wenn das nicht im Gemeinwohlinteresse liegt. Denn sie werden von der kleinen

<sup>29</sup> Hier wird nur eine bestimmte Form des *rent seeking* – der *special interest effect* – dargestellt. Für einen Überblick der unter dieser Gruppe von Theorien verhandelten Fragen siehe etwa Hindmoor, Rational Choice, 2006, S. 155 ff.; Mueller, Public Choice III, 1979, S. 333 ff.

Gruppe belohnt, von der großen aber in aller Regel nicht bestraft. Besonders attraktiv ist der *Public Choice Theory* zufolge die Verfolgung von Sonderinteressen für Politiker dann, wenn die Lasten von Personen getragen werden, die aufgrund der politischen Konstellation keine Stimme haben: etwa weil zukünftige Generationen belastet werden (**politische Kurz-sichtigkeit** oder *shortsightedness effect*) oder Menschen angrenzender Länder oder Regionen, in denen der jeweilige Politiker oder seine Partei nicht oder zu einem anderen Zeitpunkt zur Wahl stehen.

357 Als Beispiel mag die Diskussion dienen, die nach der Bundestagswahl 2009 um die Herabsetzung des Mehrwertsteuersatzes bei Hotelübernachtungen von 19 % auf 7 % aufkam. Die Hoteliers hatten in der Zeit der Wirtschaftskrise keine Möglichkeit, ihre Preise zu erhöhen. Durch die Reduzierung des Mehrwertsteuersatzes, die in aller Regel nicht an die Übernachtungsgäste weitergegeben wurde, konnte das Beherbergungsgewerbe in einem Schritt eine Preiserhöhung um 11 % realisieren, die letztlich als Steuermindereinnahmen entweder durch Steuererhöhungen an anderer Stelle kompensiert werden musste (eine Gegenfinanzierung wurde indessen nicht beschlossen) oder durch eine erhöhte Verschuldung von zukünftigen Generationen getragen werden muss. Das Thema wurde durch die Medien aufgegriffen, weil die Regierungspartei, die die Mehrwertsteuerreduzierung durchsetzte, in erheblichem Umfang Parteispenden von Hoteliers erhalten hatte; ansonsten wäre die Steuerreduzierung der rationalen Ignoranz des Wahlvolkes wegen möglicherweise überhaupt nicht wahrgenommen worden.

358 Besonders dramatisch wird dieses Problem dadurch, dass es sich durch politische Tauschgeschäfte<sup>30</sup> vervielfacht, und externe Effekte bei Gruppen auftreten können, die an der jeweiligen Wahl nicht beteiligt sind. Tritt beispielsweise eine Partei nur in den Städten A, B und C an, können externe Effekte in D und E unberücksichtigt bleiben. Tabelle 6 zeigt, wie sich das Problem auswirken kann.

---

<sup>30</sup> In der *Public Choice* Literatur firmieren solche politische Tauschgeschäfte unter dem Begriff *log rolling* (Kuhhandel). Ein eindrucksvolles aktuelles Beispiel aus dem US-Senat kann unter <http://www.c-spanvideo.org/program/284890-1> bei 150:59 (= 2:30:59) abgerufen werden.

| Wahlbezirk    | Neue<br>Brücke in A | Sanierung Hafen-<br>anlage in B | Kaufhaus-<br>komplex in C | Gesamt |
|---------------|---------------------|---------------------------------|---------------------------|--------|
| A             | + 10 €              | – 3 €                           | – 3 €                     | + 4 €  |
| B             | – 3 €               | + 8 €                           | – 3 €                     | + 2 €  |
| C             | – 3 €               | – 3 €                           | + 6 €                     | ± 0 €  |
| D             | – 3 €               | – 3 €                           | – 3 €                     | – 9 €  |
| E             | – 3 €               | – 3 €                           | – 3 €                     | – 9 €  |
| <b>Gesamt</b> | – 2 €               | – 4 €                           | – 6 €                     | – 12 € |

Tabelle 6.1: Nutzen oder Kosten je Wähler im jeweiligen Wahlbezirk

Die *Public Choice Theory* erklärt hier, was uns auch aus der allgemeinen Lebenserfahrung einleuchtet. Aber sie gibt uns gleichzeitig einen theoretischen Rahmen, mit dessen Hilfe wir nicht nur aus anekdotischer Erfahrung, sondern aus Modellen Vorhersagen ableiten können, die dieses Verhalten beschreiben. 359

### 3. Budgetmaximierung bei den Bürokraten

Wie bereits dargelegt, ist der *Public Choice Theory* zufolge das zentrale Mittel zur Maximierung des individuellen Nutzens des Bürokraten die **Maximierung des ihm zugewiesenen Haushalts** (Budget). Hinzu kommt, dass größere Budgets grundsätzlich auch im Interesse der Bürger und Gruppen sind, für die die Verwaltung zuständig ist: Ein größerer Haushalt des Wissenschaftsministeriums liegt in aller Regel auch im Interesse der Universitäten, der Studenten, Professoren und Universitätsbediensteten. Daher laufen die Interessen der Bürokraten und der Sondergruppen, denen sie dienen, oftmals gleich; insofern kann sich hier auch eine Allianz „gegen“ die Politiker aufbauen. 360

#### a. Das Modell

Wie aber gelingt es Bürokraten, ihr Budget zu maximieren bzw. sich überhöhte Budgets zu verschaffen? Die Theorie geht davon aus, dass zwischen Politikern und Verwaltungen Verhandlungen über den Haushalt geführt werden. In diesen Verhandlungen dominiere die Verwaltung, einerseits weil es eine **Informationsasymmetrie** zu ihren Gunsten gebe, andererseits weil die Verwaltung in der Lage sei, *Take-it-or-leave-it*-Angebote zu machen. 361

- 362 Die Informationsasymmetrie erklärt sich mit der Annahme, dass die Bürokraten ihrer Sachnähe wegen genauere Vorstellungen über die mit dem Angebot einer bestimmten Leistung (Output) verbundenen minimalen Kosten haben als die Politiker. Gleichzeitig fehlt den Politikern oft ein greifbarer Maßstab, anhand dessen die Leistung einer Verwaltung bewertet werden könnte. Anders als am Markt liefert die Verwaltung nämlich in aller Regel keine zählbare Menge eines präzise beschriebenen Produkts, ihre Leistung liegt vielmehr darin, ein bestimmtes Niveau an Aktivität zu erbringen. Um eine juristische Metapher zu bemühen: Geschuldet ist nicht ein Werk, sondern ein Dienst. Wir beobachten also ein Überwachungs- oder **Monitoring-Problem**. Mit dem Informationsdefizit der Politiker hängt auch die Fähigkeit der Bürokraten zusammen, in den Verhandlungen um den Haushalt *Take-it-or-leave-it*-Angebote machen zu können.
- 363 Allerdings haben die Politiker der Theorie nach in gewissem Umfang die Fähigkeit, die Verwaltung zu kontrollieren. Zur Kontrolle der Angebote und der Leistungen der Bürokraten stehen ihnen namentlich vier Instrumente zu Gebote:
- Sie können das Gesamtniveau des Outputs festlegen. Damit verhindern sie, dass Bürokraten schlicht dadurch ihren Haushalt aufblähen, dass sie ein beliebig gigantisches Leistungsniveau für einen enormen Haushalt bieten.
  - Wenngleich die Politiker die genauen Kosten für das Output-Niveau nicht kennen, können sie doch den Nutzen des Outputs beziffern und werden daher keinem Budget zustimmen, bei dem die von der Verwaltung verursachten Kosten größer sind als der Nutzen für die Bürger.
  - Ferner werden sie keinem Budget zustimmen, bei dem der marginale Nutzen des Outputs negativ ist.
  - Schließlich können die Politiker sicherstellen, dass die Bürokraten ihre Versprechen halten und die angebotene Leistung aus dem vereinbarten Haushalt auch tatsächlich erbringen.
- 364 Am Markt – oder bei vollständiger Information der Politiker – ließe sich das wohlfahrtsoptimale (und damit für die Verwaltung erforderliche) Budget folgendermaßen ermitteln:

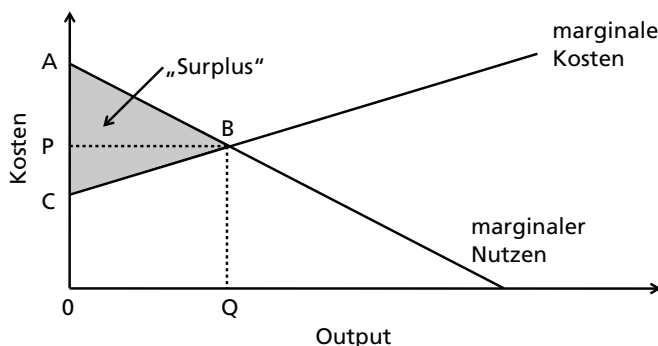


Abbildung 6.3: Wohlfahrtsoptimales Budget

365

Die Verwaltung (oder ein Unternehmen am Markt) würde die Menge  $Q$  (hier verstanden als Leistungsniveau) zum Preis  $P$  anbieten, weil bis zu diesem Punkte der marginale Nutzen die marginalen Kosten überwiegt; es handelt sich um das Marktgleichgewicht. Der dafür erforderliche Haushalt ergibt sich aus dem Viereck  $0 - C - B - Q$ . Dadurch entstünde ein wohlfahrtsökonomischer Überschuss (*surplus*), der in der Abbildung in dem Dreieck  $A - B - C$  abgetragen ist.

Unter den oben getroffenen Annahmen, dass die Bürokraten den Politikern in den Haushaltsverhandlungen überlegen sind, stellt sich aber nicht das Marktgleichgewicht ein, weil der Politiker die genaue Kostenfunktion der Verwaltung nicht kennt. Als begrenzendes Moment greift in dem Beispiel in Abbildung 6.3 nun nur, dass die Politiker keinem Budget zustimmen werden, dessen Nutzen für die allgemeine Wohlfahrt (*surplus*) geringer ist als die Kosten. Das bedeutet aber, dass die Bürokraten ihr Budget so weit aufblähen können, bis der Nutzen „aufgebraucht“ ist.

366

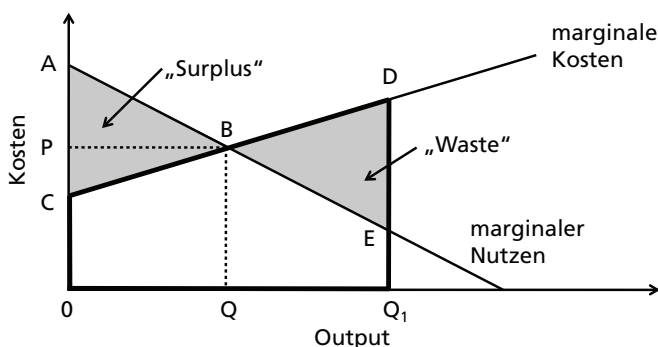


Abbildung 6.4: Verschwendung durch Budgetmaximierung seitens der Bürokraten

- 367 In dem Beispiel in Abbildung 6.4 darf mithin das „Verschwendungsdreieck“ (*waste*)  $B - D - E$  nicht größer sein als das Nutzen-Dreieck  $A - B - C$ ; es wird aber, des Strebens nach Budgetmaximierung der Bürokraten wegen, auch nicht kleiner sein: Statt das Leistungsniveau  $Q$  bereitzustellen, bieten sie – *take it or leave it!* – die Menge  $Q_1$  an; sie verdoppeln damit praktisch die Kosten für ihren Output und vergrößern ihr Budget, das jetzt durch das Viereck  $0 - C - D - Q_1$  beschrieben wird. Diesem Budgetvorschlag werden die Politiker angesichts ihrer genannten Kontrollmöglichkeiten zustimmen; es handelt sich um ein – nicht wohlfahrtsoptimales – Gleichgewicht.

### b. Einfluss und Kritik

- 368 Die Theorie des budgetmaximierenden Bürokraten hat für viel Aufsehen gesorgt. Zahllose Alternativen wurden erdacht, die verschiedene, dem Modell entgegengehaltene Schwächen ausräumen sollten – allen voran das nicht ganz einzusehende Postulat des *Take-it-or-leave-it*-Angebots. Gleichzeitig wurde mit viel Energie der Frage nachgegangen, ob und wie die Hypothesen der Theorie empirisch getestet werden könnten.<sup>31</sup>
- 369 Trotz aller Kritik hat dieser Teil der *Public Choice Theory* großen Einfluss auf Politik und Recht gehabt. In weiten Teilen gehen die Bestrebungen des *New Public Management*, des „Neuen Steuerungsmodells“ bzw. insgesamt der modernen Verwaltungswissenschaft auf Kalküle aus dieser Modellkategorie zurück. So gehen Instrumente wie das Kontraktmanagement, das „Neue Kommunale Finanzmanagement“ mit seiner Umstellung von kameralistischer auf doppische Buchführung in der öffentlichen Verwaltung, die Globalhaushalte und das Nettoprinzip, die Kosten-Leistungs-Rechnung, das *Benchmarking* und Effizienzsteigerungen durch Modifikation der Aufbau- und Ablauforganisation auf Überlegungen zurück, die bei den Anreizen von Bürokraten ansetzen. Durch das **Kontraktmanagement** etwa sollen bei den Bürokraten Anreize geschaffen werden, die Verwaltung als effizienten Dienstleister (oftmals sogar marktähnlich durch eine Organisation des Outputs in „Produktgruppen“) zu organisieren. Die Instrumente des Neuen Kommunalen Finanzmanagements und insbesondere die Globalhaushalte sollen die Leistung der Verwaltung besser messbar machen, sie anhalten, diesbezügliche Informationen für die Politik sichtbar zu machen, um dem Informationsdefizit entgegen zu wirken, und

<sup>31</sup> Eine eindrucksvolle Übersicht über die zur Effizienz behördlichen Handelns durchgeführten Untersuchungen findet sich bei *Mueller*, *Public Choice* III, 1979, S. 374–379.

sie sollen es den Bürokraten ermöglichen, ihren frei verfügbaren Haushalt durch Effizienz zu vergrößern. Viele dieser Änderungen haben Eingang ins Recht, insbesondere ins Haushaltsrecht, gefunden.

#### IV. Kollektiventscheidungen durch Wahlen und Abstimmungen: *Social Choice*

Die Arithmetik kollektiver Entscheidungen bleibt nach wie vor eines der großen ungelösten Probleme von politischer Theorie und politischer Philosophie. Wenn wir davon ausgehen, dass jedem Menschen ein grundsätzliches Freiheitsrecht zusteht, dann müssen Kollektiventscheidungen, die in diese individuelle Freiheit eingreifen, gerechtfertigt werden. Das deutsche Staatsrecht bedient sich in dieser Hinsicht der Figur des Volkswillens. Bei jeder Figur eines kollektiven Willens handelt es sich jedoch notwendigerweise um eine Fiktion, durch die die politische Anschauung der Mehrheit allen Bürgern des Staates politisch zugerechnet wird, auch wenn diese der Mehrheitsanschauung möglicherweise diametral entgegenstehen. Dabei gibt es in der politischen Theorie einen langen Streit darüber, wie der kollektive Wille am besten ermittelt wird – ob durch direktdemokratische Verfahren oder durch Institutionen repräsentativer Demokratie. 370

Die deutsche Staatsrechtslehre favorisiert dabei überwiegend die repräsentative Demokratie, wobei Repräsentation teilweise gar als „Veredelung“ des kollektiven Willens angesehen wird.<sup>32</sup> Die Prämisse ist dabei, dass sich die individuellen politischen Anschauungen in einer kollektiven Präferenz abbilden lassen, dass es also so etwas gibt wie einen kollektiven Mehrheitswillen. Die ökonomische Politiktheorie ist da sehr skeptisch. Im Folgenden wollen wir zunächst betrachten, welche Probleme bei Wahlen und Abstimmungen entstehen können, und dann zwei bekannte Theoreme näher analysieren – zum einen das *Arrows*-Theorem, zum anderen das *Ostrogorski*-Paradox. 371

##### 1. Probleme bei Wahlen und Abstimmungen

Die *Social Choice Theory* studiert Entscheidungsprobleme in Gruppen: Wie können Entscheidungen, Meinungen, Überzeugungen, Werte oder Präferenzen verschiedener Individuen zusammengenommen und in eine 372

<sup>32</sup> Böckenförde, Mittelbare/repräsentative Demokratie als eigentliche Form der Demokratie, in: Müller (Hg.): Staatsorganisation und Staatsfunktionen im Wandel: Festschrift für Kurt Eichenberger zum 60. Geburtstag, 1982, 301 ff.



**rationale Gruppenentscheidung** überführt werden? Die Anwendungsgebiete sind vielfältig; sicher eines der wichtigsten ist das Problem der Aggregation von Stimmen bei politischen Wahlen und Abstimmungen. Wir haben bereits erörtert, dass der ökonomischen Theorie zufolge die Individuen über Präferenzen verfügen, die vollständig, transitiv, ordinal, inter-personal nicht vergleichbar und konstant sind.<sup>33</sup> Für die nun folgenden Überlegungen wollen wir außerdem annehmen, dass die Wähler ehrlich sind und ihre wahren Präferenzen offenbaren, also kein taktisches Wahlverhalten an den Tag legen. Außerdem soll die „Mehrheit“ entscheiden: Zieht eine Mehrheit eine Option einer anderen vor, dann soll diese gewählt sein; das Ergebnis ist für alle verbindlich und wird durchgesetzt.

#### a. Einfache Mehrheitswahl

- 373 Die Mehrheitsentscheidung bei zwei Optionen ist einfach: Diejenige Option, die mehr Stimmen bekommt, ist gewählt. Aber schon bei drei Optionen wird das Leben sehr viel komplizierter. Schauen wir uns als Beispiel eine Entscheidung zwischen drei Kandidaten durch einfache Mehrheitswahl an, wie sie etwa bei der Erststimme bei der Bundestagswahl oder bei Wahlen im Vereinigten Königreich eingesetzt wird und bei der jeder Wähler eine Stimme hat. Als Beispiel soll und wieder das Eiskrem-Paradigma aus dem zweiten Kapitel<sup>34</sup> dienen. Danach gibt es drei Kandidaten (Vanille [V], Schokolade [S], Erdbeere [E]), und beispielsweise 21 Wähler. Die Wähler haben folgende Präferenzen:

10 Wähler: Erdbeere > Schokolade > Vanille  
 6 Wähler: Schokolade > Vanille > Erdbeere  
 5 Wähler: Vanille > Schokolade > Erdbeere

Das Ergebnis der einfachen Mehrheitswahl wäre, dass für alle Erdbeereis bestellt würde; und das, obwohl eine Mehrheit der Wähler die beiden anderen Optionen – Schokolade und Vanille – dem Erdbeereis vorgezogen hätte (vielleicht, weil sie auf einer „grundsätzlicheren“ Ebene Milcheis dem Sorbet vorziehen).

#### b. Agenda-Verfahren

- 374 Eine Alternative zur Mehrheitsentscheidung ist, paarweise zwischen den Optionen abstimmen zu lassen. Weil dieses Verfahren an das K.O.-Verfahren

<sup>33</sup> Vgl. Rz 65 ff.

<sup>34</sup> Ebd.

in Sportturnieren erinnern, wird es als „Turnier-Verfahren“ bezeichnet. Eine andere Bezeichnung als „Agenda-Verfahren“ verrät auch die größte Schwäche, die wir gleich vorführen möchten – nämlich die Abhängigkeit des Abstimmungsergebnisses von der Tagesordnung, also der Reihenfolge, in der die Optionen zur Wahl gestellt werden.<sup>35</sup> Um das Problem zu veranschaulichen, modifizieren wir das vorige Beispiel leicht:

- 10 Wähler: Erdbeere > Schokolade > Vanille  
 6 Wähler: Schokolade > Vanille > Erdbeere  
 5 Wähler: Vanille > Erdbeere > Schokolade

Nun wird paarweise abgestimmt:

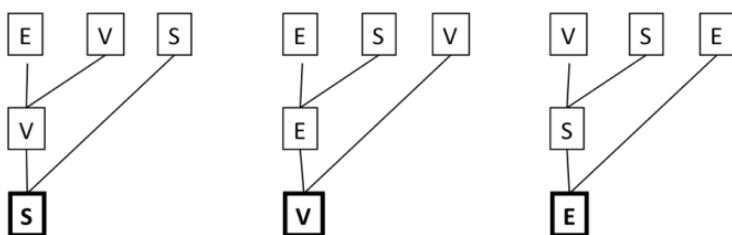


Abbildung 6.5: Abhängigkeit des Abstimmungsergebnisses von der Agenda

Dabei geschieht Erstaunliches: Je nachdem, in welcher Reihenfolge die Pärchen einander gegenübergestellt werden, kommt ein anderes Ergebnis heraus, wie Abbildung 6.5 zeigt. Tritt erst Erdbeere gegen Vanille an, und dann Vanille gegen Schokolade, so obsiegt Schokolade; tritt dagegen Erdbeere anfänglich gegen Schokolade und dann gegen Vanille an, so obsiegt – wie im ersten Schritt der ersten Abstimmung – Vanille; lässt man aber Vanille im ersten Schritt gegen Schokolade antreten, so obsiegt nach der Abstimmung zwischen Schokolade und Erdbeere – wiederum konsistent mit dem ersten Schritt der zweiten Abstimmung – Erdbeere! Drei Abstimmungen, drei Ergebnisse. Das zeigt, wie bedeutsam der Einfluss der *agenda setter*, etwa des Vorsitzenden, ist.

Das Ganze lässt sich noch auf die Spitze treiben. In unserem nächsten Beispiel sollen drei Wähler über vier Optionen abstimmen. Ihre Präferenzen seien folgendermaßen:

- Wähler 1:  $x > y > b > a$   
 Wähler 2:  $a > x > y > b$   
 Wähler 3:  $b > a > x > y$

<sup>35</sup> Riker, *Liberalism Against Populism*, 1982.

- 377 Ein geschickter Agenda-Setter kann die Abstimmung so ansetzen, dass als Ergebnis eine Option gewählt wird, die regelmäßig einer anderen unterlegen war. In unserem Beispiel gibt es eine Agenda, der zufolge die Option y gewählt wird, obwohl jeder der Wähler die Option x der Option y vorzieht:

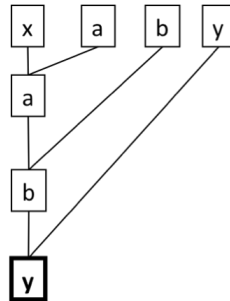


Abbildung 6.6: Agenda-Paradox – „eigentlich“ ziehen alle Wähler Option „x“ der Option „y“ vor.

### c. Condorcet-Verfahren

- 378 Um diesen Problemen zu entgehen, kann man auch unter Zuhilfenahme der vollen Präferenzordnung – und damit von mehr Information – alle Optionen paarweise vergleichen. Dieses Abstimmungsverfahren<sup>36</sup> geht auf *Marie Jean Antoine Nicolas Caritat, Marquis de Condorcet* (1743 – 1794) zurück und wird dementsprechend als **Condorcet-Verfahren** bezeichnet. Nehmen wir wieder unser Ausgangsbeispiel,

10 Wähler: Erdbeere > Schokolade > Vanille  
 6 Wähler: Schokolade > Vanille > Erdbeere  
 5 Wähler: Vanille > Schokolade > Erdbeere

so führt das *Condorcet*-Verfahren zu einer befriedigenden Lösung: Schokolade gewinnt, weil diese Option der Wahl Vanille von 16 Wählern und der Wahl Erdbeere von 11 Wählern vorgezogen wird. Allerdings kann das Condorcet-Verfahren nicht immer einen Gewinner liefern, die Präferenzen können zu einem Zirkel (*voting cycle*) führen. Als Beispiel für einen solchen Zirkel mag folgendes Beispiel dienen, an das *Arrow* in seinem Unmöglichkeitstheorem, auf das wir gleich zu sprechen kommen werden, anknüpft:

<sup>36</sup> *Marquis de Condorcet*, *Essai sur l'application de l'analyse à la probabilité des décisions rendues à la pluralité des voix*, Paris, 1785.

Wähler 1:  $a > b > c$

Wähler 2:  $b > c > a$

Wähler 3:  $c > a > b$

Mit dem *Condorcet*-Verfahren kann hier keine eindeutige Lösung gefunden werden, es ergibt sich ein Zirkel: 379

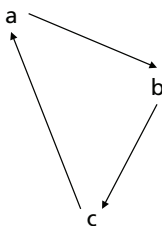


Abbildung 6.7: *Condorcet-Paradox*

#### d. *Borda*-Verfahren

Im Widerstreit mit *Condorcet* entwickelte *Jean-Charles de Borda* (1733–1799) das nach ihm benannte *Borda*-Verfahren.<sup>37</sup> Danach enthält jeder Stimmzettel eine nach Präferenzen sortierte Rangfolge. Für jede Option wird ein **Borda-Wert** errechnet, der sich aus der Summe der Ränge jeder Option auf den verschiedenen Stimmzetteln ergibt: 380

|                                    |                                         |       |
|------------------------------------|-----------------------------------------|-------|
| $A > B > C$                        | $A = 2 \cdot 1 + 1 \cdot 2$             | $= 4$ |
| $A > C > B \Rightarrow$ Borda-Wert | $B = 1 \cdot 2 + 2 \cdot 3$             | $= 8$ |
| $C > A > B$                        | $C = 1 \cdot 1 + 1 \cdot 2 + 1 \cdot 3$ | $= 6$ |

Die Option mit dem niedrigsten *Borda*-Wert ist gewählt. Das Verfahren ist einfach, und es gibt immer einen Gewinner. Aber es hat einige Nachteile. Zum einen führt es nicht immer zu den *Condorcet*-Ergebnissen, die, wie wir oben festgestellt haben, eine starke Plausibilität haben. Ein Beispiel mag dies illustrieren: 381

Es gebe drei Wähler mit folgenden Präferenzen:

2 Wähler: Vanille > Erdbeere > Schokolade > Pistazie

1 Wähler: Erdbeere > Schokolade > Pistazie > Vanille

Hier wäre nach *Condorcet* Vanille die gewählte Option, nach *Borda* in dessen Erdbeere, weil der *Borda*-Wert der Option Erdbeere (5) kleiner ist

<sup>37</sup> *De Borda*, Memoire sur les elections au scrutin, Histoire de l'Académie Royale des Sciences, 1781.

als jener von Vanille (6). Aber das *Borda*-Verfahren führt zu noch weitaus überraschenderen Ergebnissen.

- 382 Zum einen beschert es uns das *inverted order paradox*. Nehmen wir an, es gibt sieben Wähler, die zwischen vier Optionen zu wählen haben, und deren Präferenzen wie in Abbildung 6.8 dargestellt verteilt sind. Im Ergebnis hat die Option *x* den niedrigsten *Borda*-Wert und ist damit gewählt. Streicht man diese – gewinnende – Option nun und rechnet für die folgenden Optionen erneut den *Borda*-Wert aus, dann kehrt sich im Ergebnis die Reihung der übrigen Optionen um.

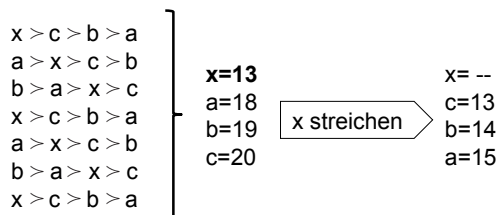


Abbildung 6.8: *Inverted Order Paradox*

- 383 Aus theoretischer Sicht noch dramatischer ist, dass irrelevante Alternativen das Ergebnis beeinflussen. Damit ist gemeint, dass das Ergebnis des Vergleichs zweier Optionen nicht allein von ihrem Verhältnis innerhalb der Präferenzordnung des Wählers abhängt, sondern vom Rang anderer Optionen. Auch hier wollen wir uns des (leicht modifizierten) Ausgangsbeispiels bedienen, um das Problem zu veranschaulichen:

10 Wähler: Erdbeere > Schokolade > Vanille  
 8 Wähler: Schokolade > Vanille > Erdbeere  
 5 Wähler: Vanille > Erdbeere > Schokolade

Ermittelt man die *Borda*-Werte, so hat die Option Erdbeere einen Wert von 44, die Option Schokolade den Wert 43 und die Option Vanille den Wert 51; Schokolade gewinnt mithin die Wahl, während Vanille verliert. Die Option Vanille kann also als irrelevante Alternative bewertet werden, sie hat keinerlei Einfluss auf den Wahlausgang. Streicht man nun diese Option und berechnet die *Borda*-Werte für die verbleibenden Optionen, so ergibt sich für Erdbeere der Wert 31 und für Schokolade der Wert 38. Plötzlich ist folglich Erdbeere der Geschmack unserer Wahl. Man muss sich das plastisch vorstellen: Wir sind also beim Eisstand und haben uns gerade für Schokolade entschieden, als der Eisverkäufer, der vielen Nachfragen überdrüssig, ein Schild „Heute leider kein Vanille-Eis!“ in die Vitrine stellt – und entscheiden uns daher dann um: für Erdbeere! Es gibt

kaum Situationen, in denen dieses Verhalten rational erscheint.<sup>38</sup> Dieser Nachteile wegen wird das *Borda*-Verfahren im politischen Raum heute nicht eingesetzt; allein Veranstaltungen wie der *Eurovision Song Contest* setzen noch Varianten dieses Verfahrens ein.

## 2. Das *Arrows*-Theorem

Die Unzulänglichkeiten verschiedener Wahlverfahren sind auch dem späteren Nobelpreisträger *Kenneth Arrow* aufgefallen. Er hat sich in seiner Dissertation zu *Social Choice and Individual Values* die Frage gestellt, ob sich systematisch ein Wahlverfahren finden lässt, das universell anwendbar ist und alle Kriterien erfüllt, die wir für eine Wahl erforderlich halten – bzw. ob sich, erkenntnistheoretisch präziser formuliert, zeigen lässt, dass es ein solches Wahlsystem nicht geben kann (daher auch die Bezeichnung seines Beweises als **Unmöglichkeitstheorem**). *Arrow* postuliert vier Erfordernisse für ein Wahlverfahren:

- Erstens muss die kollektive Präferenz **transitiv** sein. Das bedeutet, dass wenn  $a > b$  und  $b > c$ , dann muss auch  $a > c$ .
- Zweitens müssen die Präferenzen von **irrelevanten Alternativen** unabhängig sein.
- Drittens muss die Kollektiventscheidung auf den Präferenzen der Gruppenmitglieder beruhen, das heißt, wenn alle Wähler eine Option einer anderen vorziehen, sollte sich dies auch in der Kollektiventscheidung spiegeln (**schwaches Pareto-Erfordernis**).
- Viertens darf es **keinen Diktator** geben, also kein Individuum, dessen Präferenzen automatisch die Kollektiventscheidung determinieren.

Auf Grundlage dieser vier Annahmen hat *Arrow* gezeigt, dass es bei drei oder mehr Optionen, die zwei oder mehr Wählern zur Wahl gestellt werden, immer möglich ist, eine Präferenzordnung ausfindig zu machen, bei der das Finden einer konsistenten Kollektiventscheidung unmöglich ist. Jedes denkbare Wahlverfahren verletzt entweder eines der vier Kriterien oder liefert eine Rangfolge, die nicht konsistent (widerspruchsfrei) ist.

<sup>38</sup> Es gibt in ökonomischer Theorie und Philosophie Ansätze, die derlei Verhalten zu rationalisieren versuchen. Beispielsweise könnte in der Unverfügbarkeit des Vanille-Eises eine Information liegen, die für die übrigen Wahlentscheidungen relevant ist. Wenn etwa auf der Karte eines Restaurant ein schwer zuzubereitendes Gericht steht, kann ich daraus ggf. auf Fähigkeiten des Kochs schließen und daher Optionen in Betracht ziehen, die ich bei einem schlechtem Koch nicht wählen würde. Letztlich muss man aber viel Begründungsaufwand treiben, um eine solche Entscheidung zu rechtfertigen.

- 386 Beispielsweise ist es in dem uns nun schon bekannten Drei-Personen-Beispiel (Tabelle 6.2) nicht möglich, einen Stichentscheid einzuführen, der nicht eines der fünf Postulate verletzen würde.<sup>39</sup>

|          | 1. Präferenz | 2. Präferenz | 3. Präferenz |
|----------|--------------|--------------|--------------|
| Wähler 1 | a            | b            | c            |
| Wähler 2 | b            | c            | a            |
| Wähler 3 | c            | a            | b            |

Tabelle 6.2: *Voting cycle*

- 387 Nun mag man darüber streiten, wie oft solche *voting cycles* in der Realität vorkommen und ob sie in der Tat praktische Relevanz haben.<sup>40</sup> Rein stochastisch gesehen ist die Wahrscheinlichkeit relativ hoch. Für einzelne Fälle – etwa die 1991 getroffene Entscheidung zur Frage, ob Bonn oder Berlin Bundeshauptstadt sein solle – konnte gezeigt werden, dass ihnen zirkuläre Präferenzen zugrunde lagen, so dass das Ergebnis aller Wahrscheinlichkeit nach ein Artefakt der Abstimmungsreihenfolge war.<sup>41</sup> Mit Blick auf Wahlen wenden Kritiker ein, dass es bei politischen Parteien ein Kontinuum gebe, so dass Zyklen unwahrscheinlich seien: Ein Wähler, der generell links wähle, stehe wohl einer Zentrumsparterie näher als einer rechten Partei. Dieser Einwand hat zwar eine gewisse Plausibilität, kann aber nicht generalisiert werden. Gerade wenn wir davon ausgehen, dass das politische Spektrum nicht ein-, sondern mehrdimensional ist,<sup>42</sup> können wir uns durchaus vorstellen, dass Wähler sich nicht unbedingt entlang eines Kontinuums entscheiden. Ein Protestwähler, der Präferenzen für eher autoritäre Regierungen hat, mag etwa eine extrem linke Partei wählen, gleichzeitig aber eine extrem rechte Partei einer moderaten, liberalen Zentrumsparterie vorziehen. Zudem ist die Präferenz für eine bestimmte Partei immer nur ein Kompromiss.<sup>43</sup> Ein Wähler mag in einer Hinsicht der Partei A am nächsten stehen, in anderer aber die Partei B bevorzugen. Wenn man demnach *voting cycles* nicht auf Parteien, sondern auf Sachthemen bezieht, dann ist deren Auftreten sehr viel plausibler, als wenn man nur auf die Parteien schaut.
- 388 Die Erkenntnisse *Arrows* können wir auch bei den vorgestellten Wahlverfahren bestätigen:

<sup>39</sup> Arrow, A Difficulty in the Concept of Social Welfare, *Journal of Political Economy* 58 (1950), 328 ff.

<sup>40</sup> Vgl. dazu nur Hindmoor, *Rational Choice*, 2006, S. 93 ff.

<sup>41</sup> Leininger, The Fatal Vote: Berlin versus Bonn, *FinanzArch.* 50 (1993), 1 ff.

<sup>42</sup> Vgl. Rz. 349.

<sup>43</sup> Hierzu sogleich § 6 IV 3 (*Ostrogorski-Paradox*).

- Die Mehrheitswahl erfüllt zwar die postulierten Kriterien, führt aber nicht immer zu einer konsistenten Rangfolge.
- Das Agenda-Verfahren verletzt die Annahme der Unabhängigkeit von irrelevanten Alternativen und führt nicht zu einer widerspruchsfreien sozialen Rangfolge.
- Das *Condorcet*-Verfahren führt nicht immer zu einer konsistenten Rangfolge (weil es nicht immer zu einem Ergebnis führt), so dass es das Postulat der Universalität verletzt; führt man einen wie auch immer gearteten Stichentscheid ein, so wird außerdem ein „Diktator“ etabliert.
- Das *Borda*-Verfahren führt zwar zu einer widerspruchsfreien Rangfolge (das war *Bordas* stärkstes Argument gegen *Condorcet*), verletzt aber das Kriterium der Unabhängigkeit von irrelevanten Alternativen.

Daher drängt sich die Frage auf, ob wir möglicherweise unsere Erfordernisse an ein Wahlverfahren aufweichen müssen. Für das Diktator- und das *Pareto*-Kriterium wird das nicht ernsthaft diskutiert. Bei der Transitivität ließen sich ggf. Kompromisse machen (vielleicht genügt es, wenn wir nur den Gewinner einer Wahl ermitteln), aber Erweiterungen des *Arrow*-Theorems haben gezeigt, dass es auch dann noch erhebliche Schwierigkeiten gibt. Ob das Kriterium der Unabhängigkeit von irrelevanten Alternativen aufgegeben oder aufgeweicht werden soll, wird in der Tat diskutiert; es fehlen aber noch gute Begründungen, die eine gegen dieses Kriterium verstoßendes Wahlverfahren rationalisieren könnten. Hoffnung gibt es schließlich beim Kriterium der Universalität: Es gibt einen regen Forschungszweig, der Modelle entwickelt, die für bestimmte Situationen Verfahren bieten, die unsere Probleme vermeiden. 389

### 3. Das *Ostrogorski*-Paradox

Während das *Arrow*-Theorem gezeigt hat, dass die Bildung kollektiver Präferenzen teilweise unmöglich ist, geht das *Ostrogorski-Paradox* noch einen Schritt weiter: Es zeigt, dass Kollektiventscheidungen in repräsentativen Systemen die kollektiven Präferenzen teilweise verzerren.<sup>44</sup> Das Theorem stellt sich ein System vor, in dem zwei Parteien jeweils Aussagen zu drei unterschiedlichen Themen treffen. Daneben gibt es vier unterschiedliche Wählergruppen, die bei den jeweiligen Themen verschiedene Präferenzordnungen haben. Einige Wähler bevorzugen für ein Thema die Partei X, für ein anderes Thema aber die Partei Y, ganz so, wie wir es aus einer 390

<sup>44</sup> *Rae/Daudt*, The Ostrogorski Paradox: A peculiarity of compound majority decision, *European Journal of Political Research* 4 (1976), 391 ff.



repräsentativen Demokratie kennen. Selten werden wir mit dem Programm einer Partei in allen Punkten übereinstimmen. Vielmehr sind Wahlentscheidungen oft Kompromisse: Ich lege mich auf die Partei fest, die mir bei den meisten, oder den mir wichtigsten Themen am nächsten steht.

- 391 Das *Ostrogorski-Paradox* zeigt nun, dass in Fällen, in denen *en bloc* über bestimmte Fragen abgestimmt wird – wie dies in der repräsentativen Demokratie in der Regel der Fall ist –, am Ende teilweise andere Ergebnisse stehen, als wenn man einzeln über die jeweiligen Themen abgestimmt hätte. In dem in Tabelle 6.3 dargestellten Beispiel siegt Partei X, obwohl Partei Y sachbezogen jeweils die Mehrheit gehabt hätte. Auch das *Ostrogorski-Paradox* macht damit die Schwierigkeiten deutlich, die bestehen, wenn man versucht, individuelle Präferenzen in eine einheitliche kollektive Präferenz zu übersetzen.

| Wählergruppen                        | Anteil | Themenbezogene Parteipräferenz |         |         | Gewählte Partei | Wahlergebnis insgesamt             |
|--------------------------------------|--------|--------------------------------|---------|---------|-----------------|------------------------------------|
|                                      |        | Thema 1                        | Thema 2 | Thema 3 |                 |                                    |
| A                                    | 20%    | X                              | X       | Y       | X               | Partei X siegt mit 60% der Stimmen |
| B                                    | 20%    | X                              | Y       | X       | X               |                                    |
| C                                    | 20%    | Y                              | X       | X       | X               |                                    |
| D                                    | 40%    | Y                              | Y       | Y       | Y               |                                    |
| Themenbezogene Mehrheit für Partei Y |        | 60%                            | 60%     | 60%     |                 |                                    |

Tabelle 6.3: *Ostrogorski-Paradox*

#### 4. Bewertung und juristische Implikationen

- 392 Die dargestellten Schwierigkeiten bei Wahlen und Abstimmungen zeigen, dass individuelle Präferenzen nicht ohne Weiteres widerspruchsfrei in kollektive Präferenzen übersetzt werden können. Dies ist für die Rechtswissenschaft durchaus von Bedeutung. Insbesondere fordert es zu einem erneuten Nachdenken über das herrschende Demokratieverständnis in der deutschen Staatsrechtslehre heraus.<sup>45</sup> Diese versteht Art. 20 II 1 GG, demzufolge „alle Staatsgewalt [...] vom Volke aus[geht]“, so, dass jegliche Ausübung von Staatsgewalt mittels ununterbrochener Legitimationsketten

<sup>45</sup> Dazu ausführlich *Petersen*, Demokratie und Grundgesetz, JöR 58 (2010), 137 ff.; vgl. auch *Towfigh*, Das Parteien-Paradox. Ein Beitrag zur Bestimmung des Verhältnisses von Demokratie und Parteien, 2015.

auf den ‚Willen‘ des Volkes zurückgeführt werden müsse.<sup>46</sup> Was jedoch ist dieser Volkswille? Es handelt sich notwendigerweise um eine Fiktion. Der Volkswille wird in erster Linie durch das Parlament repräsentiert. Dem Parlament wird durch die Staatsrechtslehre eine zentrale Rolle zuerkannt. Dies kommt etwa in der Wesentlichkeitstheorie des BVerfG zum Ausdruck, der zufolge wesentliche Entscheidungen mit Grundrechtsrelevanz nur durch das Parlament getroffen werden können – und nicht etwa durch die Regierung, die Verwaltung oder kommunale Körperschaften. Diese Interpretation von Demokratie setzt sich in der von manchen Rechtswissenschaftlern geäußerten Skepsis gegenüber völkerrechtlichen Entscheidungsprozessen fort,<sup>47</sup> und auch die vom BVerfG in seinem Lissabon-Urteil<sup>48</sup> geübte Zurückhaltung *vis-à-vis* der europäischen Integration wird mit demokratietheoretischen Erwägungen begründet.

Die ganze Konstruktion beruht allerdings darauf, dass das Parlament 393 den Volkswillen, oder zumindest den Willen der Mehrheit, auch tatsächlich repräsentiert. Die hier dargestellten Paradoxe machen dagegen deutlich, dass kollektive Präferenzen kein statisches Konstrukt sind. Vielmehr sind sie dynamisch, was ihre Repräsentation schwierig macht. Sicherlich hängt dies davon ab, ob *voting cycles* oder die Verzerrung kollektiver Präferenzen in der Realität tatsächlich häufig vorkommen und nicht nur eine Ausnahmeerscheinung sind. Dies ist eine empirische Frage, zu der es bisher nur wenige aussagekräftige Untersuchungen gibt. Dennoch ist es nicht unplausibel, dass es sich nicht lediglich um theoretische Spielereien handelt. Wenn dies zutrifft, müsste über die Auswirkungen dieser Erkenntnisse auf die juristische Demokratietheorie neu nachgedacht werden.

<sup>46</sup> Grundlegend Böckenförde, Demokratie als Verfassungsprinzip, in: Handbuch des Staatsrechts II, 3. Aufl. 2004, § 24.

<sup>47</sup> Vgl. bspw. Bradley/Goldsmith, The Current Illegitimacy of International Human Rights Litigation, Fordham Law Review 66 (1997), 319 ff.; Stephan, International Governance and American Democracy, Chicago Journal of International Law 1 (2000), 237 ff.; Alford, Misusing International Sources to Interpret the Constitution, American Journal of International Law 98 (2004), 57 ff.

<sup>48</sup> BVerfGE 123, 267 (Lissabon-Urteil [2009]).



## § 7 – Empirische Methoden

**Literatur:** *K. Backhaus/B. Erichson/W. Plinke/R. Weiber*, Multivariate Analysemethoden: Eine anwendungsorientierte Einführung, 14. Aufl. 2016; *H. Benninghaus*, Deskriptive Statistik – Eine Einführung für Sozialwissenschaftler, 11. Aufl. 2007; *T. Cook/D. Campbell*, Quasi-Experimentation, 1979; *D. Cope*, Fundamentals of Statistical Analysis, 2005; *J. Davis*, The Logic of Causal Order, 1985; *H.E. Jackson/L. Kaplow/S.M. Shavell/W.K. Viscusi/D. Cope*, Analytical Methods for Lawyers, 2. Aufl. 2011; *L. Kish*, Representation, Randomization, and Realism, in: ders., Statistical Design for Research, 1987, 1–26; *R. M. Lawless/J. K. Robbennolt/T.S. Ulen*, Empirical Methods in Law, 2010; *S. Siegel*, Nichtparametrische statistische Methoden, 6. Aufl. 2016; *R. Singleton/B. Straits*, Approaches to Social Research, 5. Aufl. 2009; *A.H. Studenmund*, Using Econometrics, 7. Aufl. 2017; *D. Vorberg/S. Blankenberger*, Die Auswahl statistischer Tests und Maße, Psychologische Rundschau 50 (1999), 157–164; *J. Wooldridge*, Introductory Econometrics – A Modern Approach, 5. Aufl. 2013.

### I. Grundlagen und Forschungsdesign

Wie in § 1 ausgeführt wurde, kann empirische Forschung für Juristen in verschiedenen Zusammenhängen relevant werden. Dieser Abschnitt soll eine kurze Einführung in die empirische Methodik bieten. Diese soll dem Juristen erlauben, empirische Studien besser zu verstehen, deren Relevanz zu bewerten und empirische Ergebnisse eventuell gar kritisieren zu können. Empirische Forschung möchte im Wesentlichen in zweierlei Hinsicht Aussagen machen. Zum einen will sie soziale Zustände beschreiben, zum anderen versucht sie, diese zu erklären. Dabei geht es bei der Erklärung in erster Linie darum, Aussagen über bestehende **Kausalzusammenhänge** zu treffen. Warum braucht man bereits für beschreibende Forschung empirische Methoden oder gar Statistik? Der Grund dafür liegt darin, dass wir oft nicht die gesamte Wirklichkeit betrachten können, sondern lediglich einen Ausschnitt sehen, eine Stichprobe. Von dieser Stichprobe wird man allerdings nicht 1:1 auf alle denkbaren Fälle schließen können. Statistische Methoden sagen uns in diesem Fall, mit welcher Sicherheit wir Schlussfol-

394

gerungen über bestimmte Charakteristika der **Population**, also der Gesamtheit aller denkbaren Fälle, ziehen können.

395 Empirische Forschung bleibt allerdings nicht bei der Beschreibung stehen. Sie versucht auch, Kausalzusammenhänge aufzudecken. So analysiert sie etwa, ob sich die Wirtschaftsleistung eines Staates positiv auf die Stabilität einer Demokratie auswirkt, oder ob ein höherer Bildungsstand zu einem höheren Einkommen führt. In diesem Fall sind wir noch stärker auf die Statistik angewiesen als bei der bloß beschreibenden Forschung. **Statistische Testverfahren** sagen uns, ob ein statistischer Zusammenhang (eine **Korrelation**) zwischen zwei sozialen Faktoren besteht. In der sozialwissenschaftlichen Forschung werden diese Faktoren dabei als **Variablen** bezeichnet. Die abhängige (oder zu erklärende) Variable ist dabei das soziale Phänomen, dessen Auftreten erklärt werden soll, die unabhängige (oder erklärende) Variable das Phänomen, das Veränderungen in der abhängigen Variable erklären soll. Wollen wir also untersuchen, wie sich der Bildungsstand einer Person im Durchschnitt auf deren Einkommen auswirkt, dann wäre der Bildungsstand die unabhängige, das Einkommen die abhängige Variable.

396 Im Folgenden soll auf drei Punkte näher eingegangen werden. Zum einen sagt das bloße Bestehen einer Korrelation noch nichts darüber aus, ob tatsächlich ein Kausalzusammenhang besteht. Kein statistisches Testverfahren kann Aufschluss darüber geben, ob aus einer Korrelation tatsächlich auf **Kausalität** geschlossen werden kann. Daher soll in einem ersten Schritt erklärt werden, welche Regeln beim Forschungsdesign beachtet werden müssen, um Aussagen über die Kausalität treffen zu können (1). In einem zweiten Schritt soll die Messung der zu beobachtenden Variablen näher betrachtet werden (2), ehe wir schließlich kurz auf die Generalisierbarkeit der Aussagen sozialwissenschaftlicher Forschung eingehen wollen (3).

## 1. Forschungsdesign und Kausalität

### a. Kausalität bei zwei Variablen

397 Betrachten wir zunächst eine Welt, in der es nur zwei Variablen gibt, die von der Außenwelt vollkommen unabhängig sind – X und Y. Selbst wenn wir feststellen, dass X und Y miteinander korreliert sind, bedeutet dies noch nicht automatisch, dass Y auch durch X bewirkt worden ist. Vielmehr können wir uns auch vorstellen, dass der **Kausalverlauf** in die andere Richtung oder gar in beide Richtungen verläuft. Betrachten wir noch einmal das Beispiel vom Zusammenhang zwischen Wirtschaftsleistung und

Demokratie. Nehmen wir an, wir finden einen Zusammenhang zwischen beiden Faktoren: je höher die Wirtschaftsleistung, desto größer ist im Mittel die Stabilität der Demokratie. Allerdings lässt dieser Zusammenhang noch nicht den Schluss zu, dass die Stabilität der Demokratie auch durch die Wirtschaftsleistung beeinflusst worden ist. Möglicherweise ist der Kausalverlauf genau umgekehrt. Die Wirtschaft prosperiert genau deswegen, weil sie in einer Demokratie beheimatet ist. Schließlich können wir uns auch vorstellen, dass **Wechselwirkungen** zwischen beiden Faktoren bestehen, also Wirtschaftskraft zur Stärkung der Demokratie führt und dies wiederum der Wirtschaft zugutekommt.

Allerdings gibt es gewisse Regeln, die uns erlauben, in bestimmten Fällen bei zwei Variablen von der Korrelation auf die Kausalität zu schließen. Dies ist immer dann der Fall, wenn der Kausalverlauf nach unserem Verständnis von der Welt nur in eine Richtung verlaufen kann. Wenn zwei Variablen zeitlich aufeinanderfolgen, dann kann eine Korrelation zwischen beiden nur darauf zurückzuführen sein, dass die erste die zweite bewirkt hat, nicht jedoch umgekehrt. Der Bildungsgrad der Eltern mag den Bildungsgrad der Kinder beeinflussen, umgekehrt ist dies jedoch nur schwer möglich. Ähnliches gilt, wenn sich eine Variable nicht oder nur schwer verändern lässt, wie etwa die Hautfarbe oder das Geschlecht. Stellen wir also fest, dass mehr Männer als Frauen in Führungspositionen zu finden sind, dann können wir die Hypothese aufstellen, dass das Geschlecht einen Einfluss auf die Besetzung von Führungspositionen hat. Umgekehrt funktioniert diese Schlussfolgerung nicht: Nur weil jemand eine Führungsposition bekleidet, wird er damit nicht zum Mann.

Schließlich gibt es noch Fälle, in denen eine Variable zwar grundsätzlich veränderbar ist, wobei diese Veränderung aber sehr träge ist, während die andere Variable flexibel und leichter änderbar ist. In einer solchen Konstellation können wir in der Regel auch darauf schließen, dass eine Korrelation durch die träge Variable verursacht worden ist und nicht umgekehrt. Nehmen wir etwa an, wir stellen fest, dass es zwischen der Mehrheitsreligion eines Staates und dessen Staatsform eine Korrelation gibt. Zwar ist auch die Religionszugehörigkeit innerhalb der Bevölkerung Änderungen unterworfen. Theoretisch wäre es also denkbar, dass sich die Staatsform auf die Religionszugehörigkeit auswirkt. Allerdings stellen Änderungen bei der Religionszugehörigkeit sehr viel trägere und langfristige Prozesse dar als die Änderung der Staatsform, so dass wir in der Regel davon ausgehen können, dass die Religionszugehörigkeit die Staatsform beeinflusst und nicht umgekehrt.

**b. Kausalität bei mehreren Variablen**

- 400 In der Realität haben wir in der Regel nicht bloß zwei Faktoren, die unabhängig von der Außenwelt agieren. Die Regierungsform eines Staates wird nicht nur von dessen Mehrheitsreligion abhängen, sondern auch von seiner Wirtschaftskraft, der Bildung der Bevölkerung oder der ethnischen Heterogenität. Diese Faktoren brauchen uns für unsere Analyse dann nicht zu interessieren, wenn sie sich allein auf die abhängige, also die zu erklärende Variable auswirken. In diesem Fall handelt es sich um **Zufallseffekte**, die in jeder Ausprägung der unabhängigen Variable mit der gleichen Wahrscheinlichkeit auftreten. Wenn die Zahl der Beobachtungen in diesem Fall groß genug ist, dann haben diese Zufallseffekte keine Auswirkungen auf das Ergebnis. Die statistischen Testverfahren sind darauf ausgerichtet, diese Zufallseffekte herauszufiltern und uns mit einer gewissen **Wahrscheinlichkeit** anzugeben, ob die gefundenen Ergebnisse auch unabhängig von zufälligen äußeren Einflüssen Bestand hätten.
- 401 Zu Verfälschungen des Ergebnisses kann es jedoch kommen, wenn sich ein bestimmter Faktor sowohl auf die abhängige als auch auf die unabhängige Variable auswirkt. In diesem Fall sind Auswirkungen auf die abhängige Variable nicht mehr zufällig, sondern sie können das Ergebnis systematisch verfälschen und zu **Scheinrelationen** führen. So wird man etwa eine starke Korrelation zwischen den Körpergrößen von Geschwistern feststellen können. Das bedeutet jedoch nicht, dass die Größe des Bruders die der Schwester beeinflusst oder umgekehrt. Vielmehr hängen beide von einem gemeinsamen dritten Faktor ab, der Körpergröße der Eltern. Ähnlich kann man bei Bränden auch eine Korrelation zwischen der Anzahl der Feuerwehrleute und der Höhe des Sachschadens finden. Verursacht die Feuerwehr also mehr Schaden, als sie verhindert? Nach kurzem Überlegen kommen wir darauf, dass beide Phänomene auf einem gemeinsamen dritten Faktor beruhen – der Größe des Feuers, das zu löschen gewesen ist.
- 402 Es gibt grundsätzlich zwei Möglichkeiten, der Einwirkung solcher Störfaktoren Rechnung zu tragen – zum einen durch eine vorherige Kontrolle mittels des Forschungsdesigns, zum anderen durch eine nachträgliche Kontrolle bei den statistischen Testverfahren. Die effektivste Methode der Kontrolle ist dabei das Experiment, bei dem Störfaktoren durch das Forschungsdesign ausgeschaltet werden. Beim Experiment versucht man, zwei oder mehr Gruppen von Beobachtungen zu bilden, die sich allein in einem einzigen Faktor unterscheiden – der Veränderung der zu erklärenden Variable. Die Probanden werden dabei zufällig auf die einzelnen Versuchsdurchläufe verteilt – dieses Verfahren wird häufig als **Randomisierung**

bezeichnet. Ansonsten werden die Umweltbedingungen konstant gehalten, um sicherzustellen, dass es keinen Faktor gibt, der sich gleichzeitig auf die abhängige und auf die unabhängige Variable auswirkt. Alle nicht kontrollierbaren äußeren Einflussfaktoren wirken in allen Versuchsdurchläufen gleich und beeinflussen dadurch nicht den Vergleich zwischen den Experimenten.

Nehmen wir beispielsweise an, wir wollten testen, ob Gefangenen die Wiedereingliederung in die Gesellschaft leichter fällt, wenn sie für ein paar Monate eine kleine Geldrente bekommen, als wenn sie ohne Finanzierungshilfe entlassen werden.<sup>1</sup> Dazu kann man entlassene Gefangene zufällig in zwei gleich große Gruppen aufteilen: Der einen Gruppe gewährt man finanzielle Unterstützung, der anderen nicht. Am Ende kann man beobachten, bei welcher Gruppe die Rückfallquote höher war, und so feststellen, ob die finanzielle Unterstützung einen positiven Resozialisierungseffekt hatte. Sicherlich gibt es hier auch andere Faktoren, die auf den Resozialisierungserfolg des Gefangenen einen Einfluss haben: seine Persönlichkeit, die Länge seiner Haftstrafe, familiäre Unterstützung und einiges mehr. Aber bei diesen handelt es sich um Effekte, bei denen die Wahrscheinlichkeit, dass Gefangene mit einer langen Haftstrafe oder einer bestimmten Persönlichkeit in der einen oder anderen Gruppe sind, gleich hoch ist. Sie bewirken keine systematische Verfälschung des Ergebnisses. Diesen Effekten wird vielmehr durch die statistischen Testverfahren Rechnung getragen. 403

Leider ist die Durchführung eines Experiments nicht immer möglich. Manchmal kann man aufgrund **ethischer Bedenken** oder **faktischer Beschränkungen** keine Randomisierung der beobachteten Subjekte vornehmen. Möchte man etwa die Auswirkungen der Wirtschaftskraft eines Landes auf dessen Regierungsform messen, dann ist es nicht möglich, unterschiedlichen Staaten zufällig ein bestimmtes Wohlstandsniveau zuzuteilen. Vielmehr sind dies Faktoren, die vorgegeben sind und mit denen man als Wissenschaftler umgehen muss. In einem solchen Fall hilft nur eine nachträgliche Kontrolle für potentielle Störfaktoren, die in der Regel mit Hilfe von statistischen Tests vorgenommen wird. 404

Manchmal ist jedoch auch ein Kompromiss möglich. So kann es sein, dass man in der Wirklichkeit Bedingungen vorfindet, die einander sehr ähnlich sind und sich fundamental eigentlich nur in der erklärenden Variable unterscheiden. In diesem Fall spricht man von einem **Quasi-Experiment**. So hat etwa *John Henry Sloan* mit einigen Kollegen versucht, in 405

---

<sup>1</sup> Dazu *Mallar/Thornton*, Transitional Aid for Released Prisoners: Evidence from the Life Experiment, *Journal of Human Resources* 13 (1978), 208 ff.



einem Quasi-Experiment die Auswirkung von Handwaffenregulierungen auf die Kriminalitätsrate zu untersuchen.<sup>2</sup> Dabei haben die Wissenschaftler zwei Städte gewählt, die räumlich nahe beieinanderliegen und sich in ihren demographischen Faktoren fast gleichen – Seattle und Vancouver. Allerdings gibt es zwischen beiden einen entscheidenden Unterschied. Seattle liegt in den USA und hat deswegen keine Handwaffenregulierung, während Vancouver in Kanada liegt, wo der Verkauf von Handwaffen staatlich reguliert ist. Sicherlich ist dies kein perfektes Experiment, weil wir nicht vollkommen ausschließen können, dass sich Seattle und Vancouver möglicherweise in bestimmten wichtigen, aber nicht beobachteten Faktoren unterscheiden. Die Situation kommt einem Experiment aber bereits sehr nahe.

406 Wenn wir auch kein Quasi-Experiment durchführen können, müssen wir den potentiellen Störfaktoren schließlich durch das statistische Testverfahren Rechnung tragen und für **diese kontrollieren**. Eine solche Kontrolle erfolgt, indem wir alle Faktoren, die potentiell Auswirkungen sowohl auf die erklärende als auch die zu erklärende Variable haben, in unserem Regressionsmodell berücksichtigen.<sup>3</sup> Wollen wir also den Zusammenhang zwischen Wirtschaftsleistung und Regierungsform feststellen, reicht es nicht aus, allein diese beiden Faktoren zu messen und miteinander ins Verhältnis zu setzen. Vielmehr müssen wir auch die potentiellen Störfaktoren berücksichtigen, etwa den kulturellen oder religiösen Hintergrund eines Staates, Kolonialvergangenheit, Sozialkapital oder Bildungsniveau. Bevor wir also Messungen vornehmen und statistische Testverfahren durchführen, sollten wir uns erst einmal mit Papier und Bleistift hinsetzen und alle relevanten Faktoren sowie mögliche Kausalverläufe zwischen diesen identifizieren (Abb. 7.1).

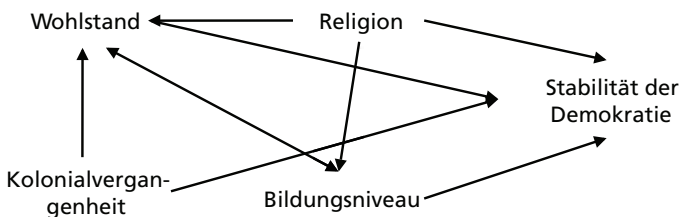


Abbildung 7.1: Pfaddiagramm, das die Kausalverläufe eines empirischen Modells darstellt

<sup>2</sup> Sloan et al., Handgun Regulations, Crime, Assaults, and Homicide, New England Journal of Medicine 319 (1988), 1256 ff.

<sup>3</sup> Hierzu Rz. 458 ff.

Kontrollieren müssen wir nur für solche Faktoren, die kausale Auswirkungen auf die erklärende (X) und die zu erklärende Variable (Y) haben. Manchmal gibt es allerdings auch Faktoren (Z), die mit beiden Variablen in Zusammenhang stehen, sich aber nur auf die zu erklärende Variable kausal auswirken. Dies ist dann der Fall, wenn X dem dritten Faktor Z im Kausalverlauf vorangeht (Abb. 7.2). Für ein kohärentes Ergebnis ist es dabei nicht notwendig, Z im Modell zu berücksichtigen. Allerdings kann es teilweise interessant sein, zwischen den direkten und den indirekten Effekten von X auf Y zu unterscheiden.

407

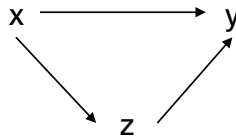


Abbildung 7.2: Unterdrückungseffekt

So können wir etwa untersuchen, wie sich der Bildungsstand eines Menschen auf seine Einstellung zum Umweltschutz auswirkt. Neben der direkten gibt es dort auch noch eine indirekte Wirkung: Der Bildungsstand wirkt sich nämlich auch auf das Einkommen einer Person aus, und dieses hat möglicherweise ebenfalls einen Einfluss auf die Einstellung zum Umweltschutz. Nehmen wir jetzt an, dass ein höherer Bildungsstand in der Regel zu einem stärkeren Umweltbewusstsein führt und sich gleichzeitig positiv auf das Einkommen auswirkt. Gleichzeitig ist es nicht unplausibel, dass ein steigendes Einkommen zu einem sinkenden Umweltbewusstsein führt. Während der direkte Einfluss von Bildung auf das Umweltbewusstsein also positiv ist, ist der indirekte über das Einkommen negativ. Wenn sich diese beiden Effekte gegenseitig aufheben (man spricht in diesem Fall von einem **Unterdrückungseffekt**), dann würde eine bloße Überprüfung des Einflusses von Bildung auf das Umweltbewusstsein ohne eine zusätzliche Kontrolle für Einkommen zu keinem Ergebnis führen. Je nach Forschungsinteresse spiegelt ein solches Ergebnis aber nur die halbe Wahrheit wider.

408

## 2. Die Messung von Daten

### a. Theoretische Überlegungen

Nachdem die einzelnen Faktoren identifiziert wurden, die für unsere Studie relevant sind, müssen die Variablen in einem zweiten Schritt gemessen

409

werden. Bevor wir jedoch die Messung vornehmen können, müssen wir noch zwei Zwischenschritte einlegen. Die Konzepte, die wir in unserem Modell haben, werden oft relativ vage sein und bedürfen daher der Konkretisierung. Wie bei der juristischen Subsumtion müssen auch in der sozialwissenschaftlichen Forschung Konzepte zunächst einmal definiert werden. Dabei müssen wir die entscheidenden Charakteristika der jeweiligen Konzepte identifizieren. Kehren wir noch einmal zu unserem Beispiel der Auswirkung der Wirtschaftsleistung auf die Stabilität der Demokratie eines Staates zurück. Sowohl bei der Wirtschaftsleistung als auch bei der Demokratie handelt es sich um Begriffe, die sich nicht selbst erklären. Verstehen wir unter Wirtschaftsleistung Wachstum, das gesamte volkswirtschaftliche Einkommen oder das Pro-Kopf-Einkommen? Liegt eine Demokratie schon dann vor, wenn die Regierung eines Staates gewählt wird, oder müssen zusätzlich noch bestimmte politische oder bürgerliche Rechte gewährt werden? Welche Anforderungen sind an die Wahlen zu stellen?

410 Der Sozialwissenschaftler wird seine Definition dabei überwiegend an Zweckmäßigkeitserwägungen orientieren. Dies ist ein möglicher Stolperstein für Juristen, die sozialwissenschaftliche Studien in ihre juristische Argumentation einbeziehen. Die Definition eines Konzeptes innerhalb einer sozialwissenschaftlichen Studie kann von derjenigen abweichen, die man im juristischen Kontext zugrunde legen möchte. Demokratie ist nicht immer gleich Demokratie. Dies kann jedoch möglicherweise dazu führen, dass die Ergebnisse einer bestimmten sozialwissenschaftlichen Studie nicht mehr auf den juristischen Kontext übertragbar sind.<sup>4</sup>

411 Schließlich müssen die Konzepte **operationalisiert** werden. Das bedeutet, dass Indikatoren gefunden werden müssen, mit denen die Konzepte auf der Grundlage ihrer Definition gemessen werden können. Schauen wir noch einmal auf unser Demokratiebeispiel und nehmen wir an, wir haben uns für einen umfassenden Demokratiebegriff entschieden, der auch die Beachtung gewisser Menschenrechte einschließt. Dann müssen wir Indikatoren finden, um die Beachtung der Menschenrechte zu messen. Das können beispielsweise Berichte von Internationalen Organisationen oder NGOs (*non-governmental organisation*) sein, die Aussagen über die Menschenrechtssituation in einem Staat treffen. Wenn wir eine statistische Analyse machen möchten, müssen die qualitativen Aussagen dieser Berichte schließlich codiert und in Zahlen übersetzt werden.

---

<sup>4</sup> Dazu näher *Petersen*, Braucht die Rechtswissenschaft eine empirische Wende?, *Der Staat* 49 (2010), 435 (452 ff.).

## b. Praktische Umsetzung

Nachdem wir uns theoretisch überlegt haben, welche Daten benötigt werden, müssen wir diese beschaffen. Dabei gibt es mehrere Möglichkeiten. Zum einen können wir auf **Felddaten**, zum anderen auf **Experimentaldaten** zurückgreifen. Beide können wir dabei entweder selbst erheben oder uns stattdessen die Erhebungen anderer zu Nutze machen. Bei Feldstudien greift man in der Regel auf bereits bestehende Daten zurück. Das können Wirtschaftsindikatoren, Daten über die Bevölkerungsstruktur oder Erhebungen zu einem vollkommen anderen Thema sein. Diese werden in der Regel von öffentlichen Institutionen, wie dem statistischen Bundesamt, der Weltbank oder der OECD, erhoben. Manchmal kann es jedoch auch bei **Feldstudien** sinnvoll sein, Daten selbst zu erheben. Dies wird oft durch Fragebögen gemacht, die an eine Auswahl von Probanden geschickt werden. Eine andere Möglichkeit ist die **Codierung von bestimmten sozialen Tatsachen**, die wir beobachten. Nehmen wir an, wir möchten untersuchen, ob sich die politische Ausrichtung von Verfassungsrichtern auf deren Entscheidungspraxis auswirkt. Dann müssten wir die Daten sowohl zur politischen Ausrichtung als auch zur Entscheidungspraxis selbst erheben, indem wir uns die für die jeweiligen Entscheidungen zuständigen Richter ansehen und entsprechend einem von uns vorher zusammengestellten Schema kodieren. 412

Experimentaldaten werden fast ausnahmslos selbst erhoben.<sup>5</sup> In Experimenten wird das Verhalten von Probanden gemessen und später ausgewertet. Die meisten Experimente finden dabei in einem **Labor** statt, in dem den Probanden bestimmte Aufgaben gegeben werden, die diese lösen müssen, oder jene in Entscheidungssituationen versetzt werden, um ihre Reaktion auf diese Situation zu testen. Wir haben gesehen, dass es in Experimenten darum geht, für alle Teilnehmer dieselben Bedingungen zu schaffen und nur ein Element zu variieren. Insofern werden die Probanden in zwei oder mehrere Gruppen aufgeteilt, wobei ihnen leicht unterschiedliche Aufgaben gegeben werden. Ist festzustellen, dass die Probanden sich jeweils unterschiedlich verhalten, kann man davon ausgehen, dass der Unterschied zwischen den verschiedenen Gruppen dafür ausschlaggebend war. Experimente müssen aber nicht notwendigerweise im Labor stattfinden. Ein Beispiel für ein Feldexperiment haben wir bereits gesehen – die 413

---

<sup>5</sup> Manchmal findet man jedoch Meta-Studien, in denen auf die Daten einer Gruppe mehrerer Experimente zurückgegriffen wird, die nicht alle unbedingt vom Autor selbst durchgeführt worden sind.

Studie über die Resozialisierungswirkung von finanzieller Unterstützung bei der Haftentlassung.

### 3. Validität der Ergebnisse

- 414 Haben wir die Messungen für eine Studie durchgeführt und die Ergebnisse ausgewertet, können wir aus unseren Beobachtungen praktische Schlussfolgerungen auf allgemeine Gesetze ziehen. Doch diese Schlussfolgerungen haben Grenzen. Diese Grenzen werden mit dem Begriff der Validität bezeichnet, wobei drei unterschiedliche Arten von Validität zu unterscheiden sind – die statistische Validität, die interne Validität und die externe Validität.<sup>6</sup> **Statistische Validität** liegt dann vor, wenn wir mit einer gewissen Sicherheit davon ausgehen können, dass eine bestimmte Beobachtung nicht auf Zufall beruht, sondern hinter dieser eine Gesetzmäßigkeit steckt. Wenn ein bestimmter Effekt, den wir beobachten, **statistisch signifikant** ist, dann können wir mit einer gewissen Wahrscheinlichkeit (in der Regel 95 %) davon ausgehen, dass der Effekt nicht zufällig ist. Umgekehrt können wir diese Schlussfolgerung allerdings nicht ziehen. Nur weil ein Ergebnis nicht statistisch signifikant ist, bedeutet dies nicht, dass die beobachtete Varianz lediglich zufällig ist. Wir können nur nicht ausschließen, dass sie nicht auf Zufall beruht.
- 415 Damit mit einiger Sicherheit von dem Beobachteten auf das Bestehen einer Gesetzmäßigkeit geschlossen werden kann, ist eine gewisse Anzahl von Beobachtungen notwendig. Das **Gesetz der großen Zahlen** besagt, dass sich die relative Häufigkeit eines Zufallsereignisses in der Regel der Wahrscheinlichkeit dieses Ereignisses annähert, je häufiger es durchgeführt wird. Die Urne in der folgenden Abbildung und die daneben stehende Graphik (Abb. 7.3) verdeutlichen die Logik hinter dem Gesetz der großen Zahlen. Angenommen, in einer Urne lägen 6000 Kugeln, die entweder blau oder grau sind. Wir möchten nun wissen, wie viel Prozent der Kugeln von grauer Farbe sind, und beginnen, verdeckt Kugeln aus der Urne zu ziehen. Die erste gezogene Kugel hat die Farbe grau (100 % der gezogenen Kugeln sind grau), die zweite Kugel hat die Farbe blau (50 % der gezogenen Kugeln sind grau) und auch die dritte gezogene Kugel hat die Farbe blau (33,3 % der gezogenen Kugeln sind grau). Je mehr Kugeln wir ziehen, umso genauer wird unser Ergebnis, und wir nähern uns der tatsächlichen Verteilung an. In unserem Beispiel befinden sich 1500 graue Kugeln in der Urne, also 25 %. Das Beispiel verdeutlicht, wie wichtig die Anzahl an Be-

---

<sup>6</sup> Vgl. dazu Cook/Campbell, Quasi-Experimentation, 1979, S. 37–94.

obachtungen für die Aussagen empirischer Forschung ist. So benötigen wir nicht alle Kugeln in der Urne, um das Verhältnis abschätzen zu können, aber eine größere Anzahl an Beobachtungen (hier ungefähr 250). Möchten wir aber nur wissen, ob mehr graue oder blaue Kugeln in der Urne sind, benötigen wir wesentlich weniger Beobachtungen (bei unserem Beispiel oben reichen 20 Ziehungen schon aus). In den statistischen Tests wird dieser Beobachtung Rechnung getragen. Je größer die Zahl der Beobachtungen insgesamt oder je größer der beobachtete Unterschied ist, desto eher werden wir ein **statistisch signifikantes Ergebnis** beobachten.

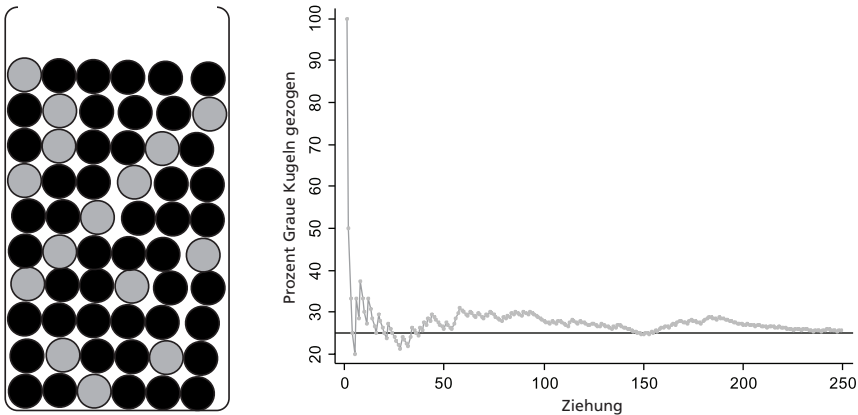


Abbildung 7.3: Beispiel einer Urne mit blauen und grauen Kugeln, Anzahl der gezogenen grauen Kugeln in einer Simulation.

Die **interne Validität** sagt uns, ob ein Ergebnis, das wir beobachten, kohärent ist. Ein Ergebnis ist insbesondere dann nicht intern valide, wenn wir bestimmte Störfaktoren, die sowohl mit der erklärenden als auch der zu erklärenden Variable korrelieren, nicht in Rechnung gestellt haben.<sup>7</sup> Mangelnde interne Validität kann zum einen auf schlechtes Forschungsdesign zurückzuführen sein, etwa weil im Vorfeld nicht alle entscheidenden Faktoren identifiziert worden sind, oder bei Experimenten die Zuteilung der Subjekte zu den Beobachtungen nicht zufällig erfolgte. Sie kann aber manchmal auch auf praktischen Problemen beruhen – etwa weil es unmöglich ist, Variablen, die potentielle Störfaktoren sein könnten, zu beobachten oder zu messen.

Schließlich müssen Ergebnisse auch **externe Validität** besitzen, um Schlussfolgerungen auf allgemeine Gesetzmäßigkeiten ziehen und sie da-

<sup>7</sup> Hierzu Rz. 400 ff.

mit auf einen anderen Kontext übertragen zu können. Eine Stichprobe kann nur für die Population extern valide Ergebnisse liefern, aus der sie gezogen wurde. Führen wir also ein Experiment mit deutschen, männlichen Arbeitnehmern durch, dann sind die Ergebnisse des Experiments über die bloßen Teilnehmer des Experiments hinaus auf alle deutschen, männlichen Arbeitnehmer generalisierbar, soweit wir die Stichprobe zufällig gezogen haben. Die Übertragbarkeit der Erkenntnisse auf weibliche Arbeitnehmer oder auf chinesische Unternehmer ist dagegen fragwürdig. Dem kann in der Praxis dadurch Rechnung getragen werden, dass entweder die Population des Ursprungsexperiments ausgeweitet oder dieses mit anderen Populationen wiederholt wird.

- 418 Zudem ist bei Experimenten zu beachten, dass diese bewusst in einer **kontrollierten Umgebung** durchgeführt werden, aber außerhalb dieser Umgebung zusätzliche Einflüsse existieren, welche die gemessenen Ergebnisse überlagern können. Bei der Frage nach der Parallelität zwischen der experimentellen Situation und der Situation im Alltagsleben (auch **Isomorphismus** genannt) muss beachtet werden, ob alle relevanten Begebenheiten im experimentellen Design berücksichtigt wurden. Sollten wichtige Eigenschaften fehlen, muss ein Jurist vorsichtig sein, die Ergebnisse sozialwissenschaftlicher Studien aus ihrem Kontext zu reißen und sie auf Sachverhalte zu übertragen, auf die sie nicht ohne Weiteres übertragbar sind.

## II. Deskriptive Statistik

- 419 Die **deskriptive Statistik** dient der Beschreibung empirischer Daten. Anders als bei der später noch zu behandelnden Teststatistik lassen sich jedoch mit der deskriptiven Statistik keine Folgerungen über Charakteristika sozialwissenschaftlicher Phänomene oder Zusammenhänge empirischer Faktoren ziehen. Am einfachsten kann man empirische Daten durch Tabellen darstellen, in welche alle gemessenen oder ermittelten Werte eingetragen werden. Im Finanzmarktteil einer Tageszeitung findet man Beispiele für solche Tabellen. Für jede geführte Aktie ist dort der Kurs vom vorherigen Tag vermerkt. Diese Tabellen sind meistens mehrere Seiten lang und – sollte man nicht Interesse an einer bestimmten Aktie haben, sondern generell am Aktienmarkt – eher uninteressant und unübersichtlich. Daher findet man in der Zeitung auch Kennzahlen, welche Informationen bündeln – zum Beispiel den DAX-Wert. Da diese Kennzahlen im Laufe der Zeit (Tag, Woche, Jahr) schwanken, werden sie häufig zusätzlich durch Graphiken, die den Zeitverlauf wiedergeben, dargestellt. Genau wie Infor-

mationen über Aktien unterschiedlich dargestellt werden, benötigt man auch für andere empirische Daten unterschiedliche Darstellungsformen. Die deskriptive Statistik erlaubt eine solche Darstellung in Form von Tabellen, Graphiken und Kennzahlen.

### 1. Statistische Variable

Bei empirischen Arbeiten werden einzelne **Merkmalsträger** aus einer Grundgesamtheit beobachtet. Die Grundgesamtheit wird häufig auch als **Population** bezeichnet. Merkmalsträger besitzen **Merkmale**, häufig auch **statistische Variable** genannt, die innerhalb der Studien gemessen werden. Merkmalsträger unterscheiden sich in der Ausprägung ihrer Merkmale. Bei einem Vergleich der Straftaten zwischen deutschen Städten bildeten alle Städte die Population, Berlin den Merkmalsträger, Straftaten das Merkmal und 482.765 begangene Straftaten die Merkmalsausprägung.<sup>8</sup>

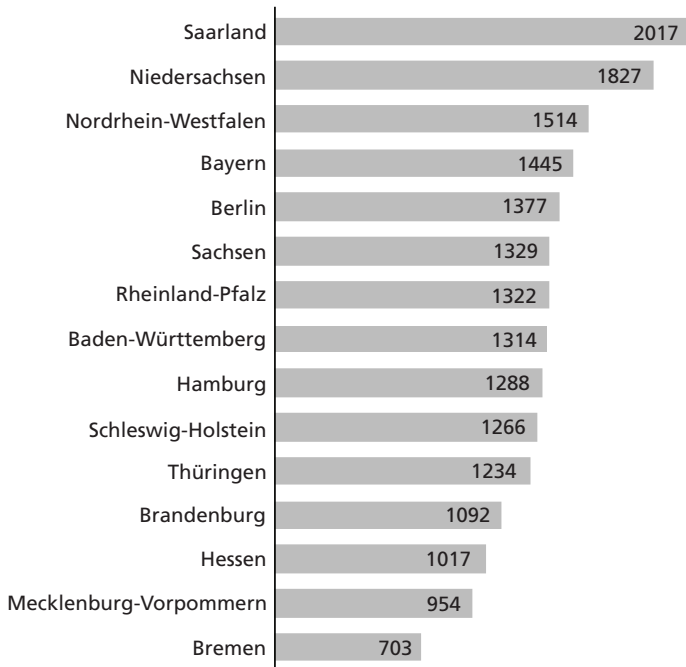
Eine gängige Unterteilung von Variablen ist die Unterscheidung zwischen qualitativen Merkmalen und quantitativen Merkmalen. **Quantitative Merkmale** sind numerisch und können in kontinuierliche und diskrete Variablen unterteilt werden. **Kontinuierliche Variablen** nehmen beliebige Werte an und sind auf einer metrischen Skala messbar. Beispiele hierfür sind Temperatur, Zeit und Größe. **Diskrete Variablen** können nur bestimmte Werte annehmen und haben zwischen diesen Werten Sprünge: Misst man die Ausprägung des Merkmales *Anzahl Kinder* einer Familie, so kann man ein, zwei, drei oder eine andere ganzzahlige Anzahl Kinder messen, aber keine 1,5 Kinder.

Eine statistische Variable ist dann ein **qualitatives Merkmal**, wenn man sie eindeutigen Kategorien zuordnen kann. Beispiele für qualitative Merkmale sind das Geschlecht mit den Ausprägungen männlich und weiblich oder der Familienstand mit den Ausprägungen ledig, verheiratet, verwitwet, geschieden, Ehe aufgehoben, eingetragene Partnerschaft, durch Tod aufgehobene Partnerschaft, aufgehobene Partnerschaft und nicht bekannt. Ein qualitatives Merkmal gehört immer genau einer Kategorie an. Daher wird es auch **kategoriales Merkmal** genannt. Qualitative Merkmale zweier Merkmalsträger sind entweder gleich oder ungleich. Ein rechnerischer Vergleich, mit dem wir beispielsweise die Größe eines Unterschiedes feststellen wollen, ist nicht möglich. Das Balkendiagramm (Abb. 7.4) ist ein Beispiel zur graphischen Darstellung von qualitativen Merkmalen. Merkmalsträger sind jugendliche Verurteilte in Deutschland, welche sich mit

---

<sup>8</sup> Bundeskriminalamt, Polizeiliche Kriminalstatistik 2008.





Quelle: Statistisches Bundesamt (2008), keine Daten für Sachsen-Anhalt

Abbildung 7.4: Beispielhafte Darstellung des qualitativen Merkmals *Bundesland*

dem Merkmal *Bundesland* unterscheiden. Balkendiagramme geben für jede Kategorie die Anzahl der Merkmalsträger an, welche in diese Kategorie passen.

## 2. Histogramme und Verteilungen

- 423 Um quantitative Daten darzustellen, werden häufig **Histogramme** verwendet. Histogramme sehen den oben beschriebenen Balkendiagrammen sehr ähnlich. Abb. 7.5 zeigt zwei Histogramme, welche das Bruttoerwerbseinkommen der Erwerbstätigen im sozioökonomischen Panel wiedergeben.<sup>9</sup>

<sup>9</sup> Das sozioökonomische Panel ist ein Datensatz, welcher durch jährliche Befragung von privaten Haushalten in Deutschland ermittelt wird. Um repräsentative Daten über Einkommen, Armut, Wohnsituation, Einwanderung und Bildung in Deutschland zu sammeln, werden jährlich etwa 12500 Haushalte befragt.

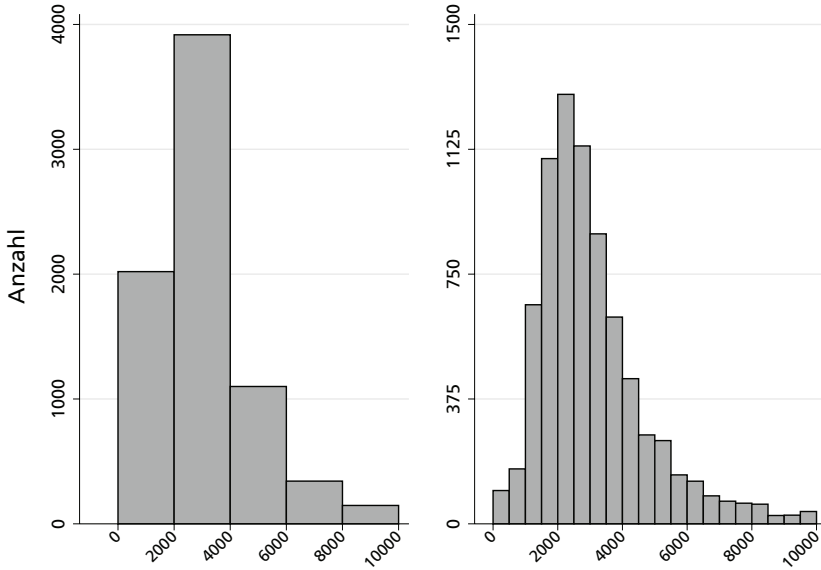


Abbildung 7.5: Histogramme mit unterschiedlichen Intervallgrößen über das Bruttoerwerbseinkommen

Wie bei den Balkendiagrammen werden auch im Histogramm Daten in Kategorien aufgeteilt. Diese Intervalle sind nicht fest vorgegeben. Sie werden nach Größe sortiert auf einer durchgehenden Achse aufgetragen. Anders als bei qualitativen Daten sind rechnerische Vergleiche zwischen den Kategorien möglich (jemand in dem Intervall € 0 – € 2.000 hat weniger Einkommen als jemand in dem Intervall € 2.000 – € 4.000). Einem Histogramm kann man durch die Höhe der Balken die Anzahl der Personen pro gewähltem Intervall entnehmen. In Abb. 7.5 werden dieselben Informationen mit zwei unterschiedlichen Einteilungen in Kategorien dargestellt. Auf der linken Seite sind € 2.000 als Intervallbreite festgelegt, auf der rechten Seite € 500. Durch eine präzisere Wahl der Intervallbreite werden zusätzliche Informationen dargestellt. So erkennt man auf der rechten Seite den starken Unterschied zwischen der Anzahl an Personen, welche € 500 – € 1.000 verdienen, und denen, welche € 1.000 – € 1.500 verdienen. Diese Informationen sind in dem Histogramm auf der linken Seite verdeckt. Wichtig ist, dass die Skalen der beiden Histogramme unterschiedlich sind, aber die Summe der Anzahl in den Intervallen von € 0 – € 2.000 auf der rechten Seite gleich der Balkenhöhe für das Intervall € 0 – € 2.000 ist.

- 425 Anhand des Histogramms kann man einfach erkennen, welche Löhne häufig und welche selten in der Beobachtung vorkommen. Aber der Hauptgrund, warum Histogramme eine sehr beliebte Darstellungsform empirischer Daten sind, ist, dass man ihre Verteilung betrachten kann. Von besonderem Interesse ist es hierbei, welche Form die Verteilung der Daten annimmt. Die **Gleichverteilung** und die **Normalverteilung** (siehe Abb. 7.6) sind zwei Verteilungsarten, mit denen Daten gern verglichen werden. Wäre das Bruttoerwerbseinkommen im sozioökonomischen Panel gleich verteilt, gäbe es für jeden beobachteten Lohn dieselbe Anzahl an Erwerbstätigen. Bei einem Histogramm, welches gleich verteilte Daten darstellt, sind alle Balken gleich groß. Eine Normalverteilung hat die Form einer Glockenkurve, bei der die größte Häufigkeit in der Mitte liegt. Wäre das Bruttoerwerbseinkommen im sozioökonomischen Panel normal verteilt, wäre der Lohn in der Mitte der Verteilung der, der von den meisten Erwerbstätigen verdient würde. Gleichzeitig würde je die gleiche Anzahl von Personen mehr und weniger als diesen Lohn verdienen.

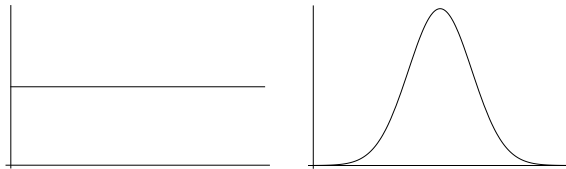


Abbildung 7.6: Skizze einer Gleichverteilung (links) und einer Normalverteilung (rechts)

- 426 Aber Verteilungen können auch andere Formen annehmen. Abb. 7.7 zeigt die verschiedenen Eigenschaften, die Verteilungen haben können. Wenn eine Verteilung die meisten Beobachtungen links der Mitte besitzt, bezeichnet man sie als **rechtsschief** (da die Daten nach rechts „auslaufen“). Wenn sie die meisten Beobachtungen rechts der Mitte besitzt, wird sie als **linksschief** bezeichnet (da die Daten nach links „auslaufen“). Besitzt eine Verteilung zwei Spitzen, bezeichnet man sie als **bimodal**.

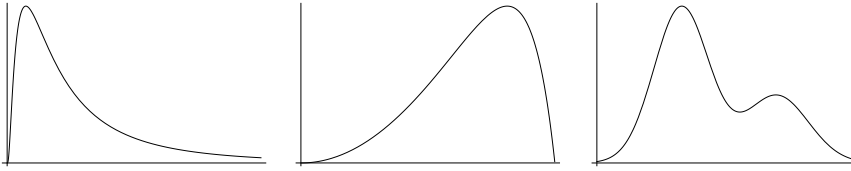


Abbildung 7.7: Eigenschaften von Verteilungen – rechtsschief (links), linksschief (mittig) sowie bimodal (rechts)

Natürlich ist das Aussehen der obigen Verteilungen idealisiert und kann durch Histogramme nur angenähert werden. Damit ein Histogramm die exakte Form dieser Verteilungen erreichen kann, müssten die Messungen sehr genau sein, die Intervalle des Histogramms extrem klein und die Anzahl der Beobachtungen extrem hoch. Daher werden normalerweise nur Annäherungen an Verteilungen beobachtet. 427

### 3. Kennzahlen

Anhand von Histogrammen können quantitative Daten sehr anschaulich betrachtet werden. Jedoch ist dies manchmal umständlich und häufig ist man gar nicht an der vollständigen Verteilung der Daten interessiert, sondern nur am Vergleich einzelner Kennzahlen. Kennzahlen geben einzelne Eigenschaften der Verteilung numerisch wieder. Es wird dabei zwischen Kennzahlen unterschieden, welche die **Lage** und die **Streuung** beschreiben. 428

**Lageparameter** geben zentrale Tendenzen einer Verteilung wieder. Die gängigsten Lageparameter sind die **Mittelwerte**, zu denen Modus, arithmetisches Mittel und Median gehören.<sup>10</sup> Der **Modus** ist der Wert, welcher am häufigsten in der Verteilung vorkommt, also z.B. der Lohn, den die meisten Erwerbstätigen einer Stichprobe erhalten. Das **arithmetische Mittel** (auch schlichtweg **Mittelwert** genannt) wird berechnet, indem man die Summe aller beobachteten Werte durch die Anzahl der Beobachtungen in der Stichprobe teilt: 429

$$\bar{x} = \frac{x_1 + x_2 + \dots + x_N}{N}$$

**Beispiel:** Wegen Geschwindigkeitsübertretung mussten fünf Autofahrer jeweils die Bußgelder € 15, € 15, € 25, € 35 und € 80 bezahlen. Wie hoch ist das arithmetische Mittel? 430

<sup>10</sup> Weitere Mittelwerte sind das geometrische Mittel, das harmonische Mittel und das quadratische Mittel.

$$\frac{15 + 15 + 25 + 35 + 80}{5} = \frac{170}{5} = 34$$

Der Mittelwert des Bußgeldes beträgt also € 34. Wenn alle fünf Fahrer € 34 zahlen, ergibt das dieselbe Summe wie die der tatsächlichen Bußgelder, also haben im Durchschnitt alle Fahrer € 34 bezahlt. Daher wird der Mittelwert in der Umgangssprache häufig auch als Durchschnitt bezeichnet.

- 431 Der **Median** zerteilt eine der Größe nach geordnete Stichprobe in zwei Hälften. 50 % aller Beobachtungen haben einen Wert, der niedriger als der (oder gleich dem) Median ist, und 50 % aller Beobachtungen haben einen Wert, der größer als der (oder gleich dem) Median ist. Formal wird der Median wie folgt berechnet:

$$x = \begin{cases} x_{(N+1)/2} & \text{wenn } N \text{ ungerade} \\ \frac{x_{N/2} + x_{(N/2)+1}}{2} & \text{wenn } N \text{ gerade} \end{cases}$$

- 432 Bei einer ungeraden Anzahl hat der Median einfach den Wert der Beobachtung, der in der geordneten Reihe in der Mitte steht. Diese Stelle kann berechnet werden, indem man auf die Stichprobengröße  $N$  eins hinzurechnet und diesen Wert durch zwei dividiert. Für eine gerade Stichprobengröße hat der Median den (fiktiven) Wert zwischen dem Wert an der Stelle  $N/2$  und der Stelle  $(N/2)+1$ .

- 433 **Beispiel:** Wegen zu schnellen Fahrens haben sieben Autofahrer Strafpunkte in Flensburg erhalten. Die Punkte sind 1, 1, 1, 3, 4, 4 und 4. Da die Anzahl der Beobachtungen in der Stichprobe ungerade ist, befindet sich der Median an der Stelle  $(7+1)/2 = 8/2 = 4$ . Der Median befindet sich also an 4. Stelle und beträgt somit 3 Punkte. Nun schauen wir uns den Fall mit einer geraden Anzahl an Beobachtungen an. Sechs Autofahrer haben die folgenden Punkte erhalten: 1, 1, 3, 4, 4 und 4. Der Median ist somit in der Mitte zwischen dem Wert an der 3. Stelle (6 dividiert durch 2) und der 4. Stelle (6 dividiert durch 2 plus 1). Der Wert an der 3. Stelle beträgt 3 Punkte und der an der 4. Stelle beträgt 4 Punkte, somit ist der Median 3,5 Punkte  $((3+4)/2)$ .

- 434 Der Median ist robuster als der Mittelwert gegenüber Ausreißern (also extremen Werten) und schiefen Verteilungen. Als Beispiel betrachten wir noch einmal die Verteilung des Bruttoerwerbseinkommens im sozioökonomischen Panel. Abb. 7.8 gibt die geschätzte Verteilung wieder. Aufgrund der sehr hohen Einkommen, welche nur wenige der Befragten verdienen, ist die Verteilung des Einkommens rechtsschief.

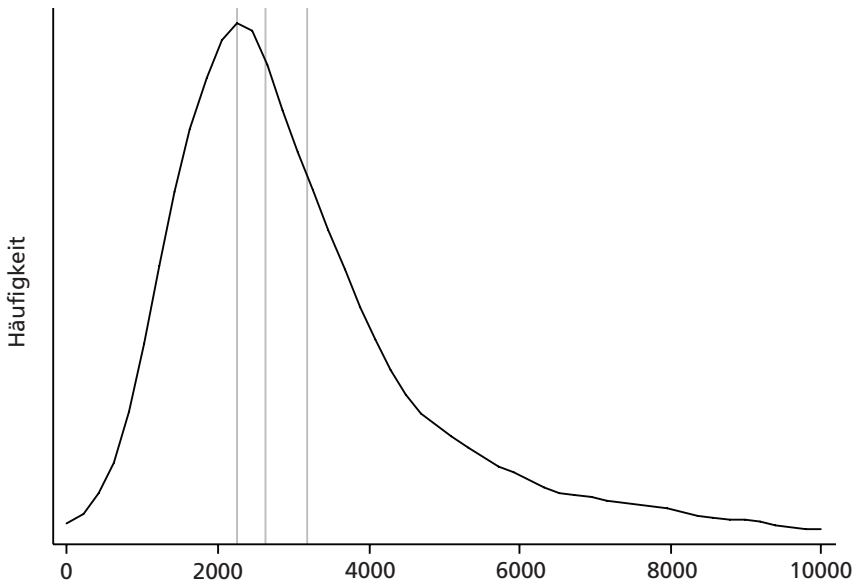


Abbildung 7.8: Geschätzte Verteilung des Bruttoerwerbseinkommens

Der Modus ist durch die linke vertikale Linie, der Median durch die mittlere vertikale Linie und das arithmetische Mittel durch die rechte vertikale Linie markiert. Der Median beträgt € 2.636 und das mittlere Bruttoerwerbseinkommen beträgt € 3.178. Der Mittelwert gibt an, wie viel der Einzelne verdienen würde, wenn alle das gleiche Einkommen bekämen. Der Median ist das Einkommen, bei dem 50 % der Befragten weniger (oder gleich viel) und 50 % der Befragten mehr (oder gleich viel) verdient haben. Dieses Beispiel verdeutlicht, warum es gerade bei Verteilungsfragen darauf ankommt zu wissen, welcher Lageparameter für welche Aussage verwendet wird.

Verteilungen besitzen aber noch zusätzliche Eigenschaften, welche nicht durch Lageparameter erfasst werden. Abb. 7.9 zeigt Histogramme von zwei imaginären Lohnverteilungen. Beide Verteilungen besitzen ungefähr dieselben Mediane, Mittelwerte und **Modalwerte**. Trotzdem sehen die zwei Verteilungen sehr verschieden aus, da sie unterschiedliche Streuungen haben. Um diese zu messen, gibt es **Streuungsmaße** – die geläufigsten bezeichnet man als Standardabweichung sowie Stichprobenvarianz.

Die **Stichprobenvarianz** wird mit den quadrierten Abständen der einzelnen Beobachtungen zum Mittelwert der Stichprobe berechnet. Je größer die Streuung der Stichprobe, desto größer ist damit auch die Varianz. Letztere ist wie folgt definiert:

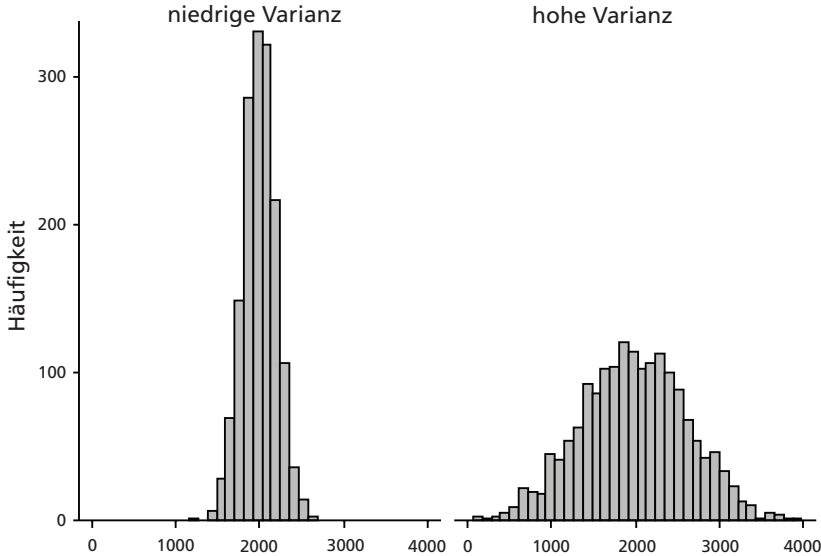


Abbildung 7.9: Histogramme zweier imaginärer Lohnverteilungen mit niedriger (links) und hoher (rechts) Varianz

$$s^2 = \frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2$$

Bei der Berechnung wird von jedem einzelnen Wert in der Stichprobe der Mittelwert  $\bar{x}$  abgezogen und das Ergebnis der Differenz quadriert. Diese Quadrate werden addiert und die Summe durch die Anzahl an Beobachtungen (minus 1) dividiert.

- 438 **Beispiel:** Betrachten wir die Bußgelder für Geschwindigkeitsübertretungen von fünf Autofahrern. Die Bußgelder betragen € 15, € 15, € 25, € 35 und € 80. In diesem Kapitel wurde bereits der Mittelwert für diese Stichprobe berechnet, er beträgt € 34. Dadurch berechnet sich die Varianz wie folgt:

$$\frac{(15-34)^2 + (15-34)^2 + (25-34)^2 + (35-34)^2 + (80-34)^2}{5-1} = 730$$

Die Varianz der Stichprobe beträgt somit 730 €<sup>2</sup>. Die Varianz wird immer in der jeweiligen Einheit zum Quadrat angegeben. Das zweite Maß, die **Standardabweichung**, hängt direkt mit der Varianz zusammen, da sie einfach als die Quadratwurzel der Varianz definiert ist:

$$s = \sqrt{\frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2}$$

Für das obige Beispiel ergibt sich somit eine Standardabweichung von 27,02 €. Die Einheit der gemessenen Standardabweichung entspricht der jeweiligen Einheit der Stichprobe. Sowohl Varianz als auch Standardabweichung werden als äquivalente Streuungsmaße verwendet. Neben der Varianz und der Standardabweichung gibt es noch den **Standardfehler**<sup>11</sup> des Mittelwertes. Der Standardfehler gibt die Güte einer Schätzung an. Grundsätzlich soll mit dem Mittelwert einer Stichprobe der (fixe) Mittelwert der Population geschätzt werden. Der Mittelwert einer Stichprobe kann von diesem Wert jedoch auf Grund von Zufallseinflüssen und Messfehlern abweichen. Je kleiner der Standardfehler des Mittelwertes ist, umso höher ist seine Genauigkeit. Der Standardfehler ist definiert als die Standardabweichung dividiert durch die Quadratwurzel der Anzahl an Beobachtungen:

$$SE = \frac{s}{\sqrt{n}}$$

Der Standardfehler wird also kleiner, wenn die Varianz bzw. die Standardabweichung kleiner wird oder aber wenn die Stichprobe größer wird. Ist die Stichprobengröße extrem groß, wird der Standardfehler sehr klein. Im Gegensatz dazu gibt die Standardabweichung die Streuung innerhalb der Stichprobe/Grundgesamtheit wieder. Diese Streuung (z.B. des Gewichtes, des Alters oder des Einkommens) bleibt auch bei genauen Messungen und großen Stichproben vorhanden.

### III. Statistische Testverfahren

Mit Hilfe der deskriptiven Statistik kann man Daten beschreiben, wichtige Informationen über eine oder mehrere Stichproben sammeln und vergleichen. Häufig sollen jedoch weiter reichende Aussagen über gesammelte Daten getroffen werden. Mit Hilfe der Stichproben möchte man entweder Aussagen über bestimmte Charakteristika der Population oder über den möglichen Zusammenhang zwischen zwei oder mehreren Variablen treffen. Im Folgenden werden zuerst einige Grundbegriffe statistischer Tests erläutert und anschließend anhand von Beispielen die Anwendungsgebiete einiger Verfahren illustriert.

<sup>11</sup> Generell ist der Standardfehler ein Streuungsmaß für Schätzungen eines unbekannten Parameters. Das kann der Mittelwert einer Population, aber zum Beispiel auch ein Regressionskoeffizient sein. Wir gehen nur auf den Standardfehler des Mittelwertes genauer ein, da andere Standardfehler analog zu interpretieren sind.



## 1. Grundbegriffe statistischer Tests

- 442 Die meisten statistischen Testverfahren werden angewandt, um Unterschiede zwischen zwei Stichproben oder Merkmalen zu qualifizieren. Hierbei wird ein Unterschied zwischen zwei Stichproben **signifikant** genannt, wenn die Wahrscheinlichkeit, dass dieser Unterschied durch Zufall entstanden ist, sehr gering ist. Ausgangspunkte eines jeden statistischen Testes sind die **Nullhypothese**  $H_0$  und die **Alternativhypothese**  $H_1$ . Die Nullhypothese ist dabei meistens, dass die Stichprobe Teil der Population ist oder dass zwischen zwei Faktoren kein Zusammenhang besteht. Die Alternativhypothese ist dagegen auf den Unterschied ausgerichtet: Entweder ist die Stichprobe nicht Teil der Population oder es besteht ein Zusammenhang zwischen zwei Faktoren. Sie bildet normalerweise die Forschungshypothese und soll bestätigt werden. Es wird immer  $H_0$  gegen  $H_1$  getestet und die Hypothesen müssen so formuliert sein, dass nur eine zutreffen kann.
- 443 Möchte man z.B. die durchschnittliche Körpergröße von Männern mit der von Frauen vergleichen, wäre die Nullhypothese, dass es keinen Unterschied gibt, und die Alternativhypothese, dass es einen solchen Unterschied gibt. Einen Test dieser Art, bei der es keine gerichtete Hypothese gibt, nennt man einen **zweiseitigen Test**. Hat man dagegen eine gerichtete Alternativhypothese, in diesem Beispiel, dass Männer im Durchschnitt größer sind als Frauen, wendet man einen **einseitigen Test** an.
- 444 Mit statistischen Tests wird man nie mit absoluter Sicherheit sagen können, ob die Null- oder die Alternativhypothese zutreffen. Vielmehr handelt es sich um Wahrscheinlichkeitsaussagen. Wie sicher man sich sein möchte, dass die Alternativhypothese tatsächlich zutrifft, wird über das **Signifikanzniveau**  $\alpha$  bestimmt. Das Signifikanzniveau sagt aus, wie hoch die Wahrscheinlichkeit ist, dass die Alternativhypothese fälschlicherweise angenommen wurde (**Fehler 1. Art**). Je kleiner  $\alpha$  ist, desto höher ist natürlich die Wahrscheinlichkeit, dass man einen umgekehrten Fehler macht, nämlich die Alternativhypothese ablehnt, obwohl sie richtig ist. In den Sozialwissenschaften geht man in der Regel von einem zu erreichenden Signifikanzniveau von 5 % aus. Daneben werden teilweise noch weitere Abstufungen gemacht. So werden Ergebnisse auf dem 10 %-Niveau oft als schwach signifikant, solche auf dem 1 %-Niveau als hoch signifikant bezeichnet. Normalerweise geben Statistikprogramme nicht das Signifikanzniveau wieder, sondern den *p-Wert* eines Testes. Der *p-Wert* ist die berechnete Wahrscheinlichkeit des angewandten Tests für einen Fehler 1. Art. Falls der *p-Wert* kleiner als das  $\alpha$  ist, also unterhalb des Signifikanzniveaus liegt, sprechen wir von einem statistisch signifikanten Ergebnis.

Auch die Nicht-Signifikanz ist mit einer Irrtumswahrscheinlichkeit be-  
legt. Wenn fälschlicherweise die Nullhypothese bestätigt wurde und die  
Alternativhypothese abgelehnt wurde, spricht man vom **Fehler 2. Art**. An-  
ders als beim Fehler 1. Art ist aber die Berechnung des Fehlers 2. Art oft  
nicht möglich und damit diese Irrtumswahrscheinlichkeit unbekannt. Da-  
her bedeutet ein nicht-signifikantes Ergebnis nicht, dass die verglichenen  
Stichproben gleich sind oder dass es keinen Effekt gibt. Die einzig zulässige  
Aussage bei einem nicht-signifikanten Ergebnis ist, dass kein signifikanter  
Unterschied oder Effekt gefunden wurde. Die folgende Tabelle fasst die  
möglichen Fehler noch einmal zusammen.

445

|                                           | H <sub>0</sub> trifft zu                                                                     | H <sub>1</sub> trifft zu                                                                |
|-------------------------------------------|----------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------|
| Test weist H <sub>0</sub><br>nicht zurück | H <sub>0</sub> ist richtig<br>und wird vom Test nicht<br>abgelehnt:<br>Richtige Entscheidung | H <sub>0</sub> ist falsch,<br>wird aber vom Test nicht<br>abgelehnt:<br>„Fehler 2. Art“ |
| Test weist H <sub>0</sub><br>zurück       | H <sub>0</sub> ist richtig,<br>wird aber vom Test<br>abgelehnt:<br>„Fehler 1. Art“           | H <sub>0</sub> ist falsch<br>und wird vom Test<br>abgelehnt:<br>Richtige Entscheidung   |

Tabelle 7.1: Fehler beim Testen von Hypothesen

2. Auswahl des Testverfahrens

Es existiert eine unüberschaubare Anzahl an statistischen Testverfahren,  
welche für unterschiedliche Fragestellungen, unterschiedliche Arten von  
Variablen und unterschiedliche Zusammenhänge zwischen den zu testen-  
den Variablen entwickelt wurden. Im Folgenden werden einige typische  
Fragestellungen erläutert, ohne dass dabei ein Anspruch auf Vollständig-  
keit erhoben wird.<sup>12</sup> Da die Berechnungen und die Funktionsweisen der  
einzelnen Tests recht unterschiedlich sein können, wird nicht weiter auf  
die Berechnung der Teststatistiken eingegangen. Diese Beispiele sollen nur  
verdeutlichen, für welche typischen Fragestellungen die Statistik wertvolle  
Werkzeuge zu bieten hat. Sie sind nicht geeignet, dem Leser die konkrete

446

<sup>12</sup> Ein vollständiger Überblick würde, wenn überhaupt möglich, den Rahmen dieses Kapitels sprengen. Die unglaublich große Anzahl an Testverfahren sowie einzelne Überschneidungen zwischen den Anwendungsgebieten der Testverfahren machen einen vollständigen und gleichzeitig kurzen Überblick unmöglich. Für eine umfassendere Entscheidungshilfe zur Anwendung von statistischen Tests und Maßen sei auf den Entscheidungsbaum von Vorberg/Blankenberger, Die Auswahl statistischer Tests und Maße, Psychologische Rundschau 50 (1999), 157 ff. verwiesen.

Anwendung dieser Werkzeuge beizubringen. Für einen ersten Einblick in die Welt der Testverfahren und ein erstes Verständnis ist dies auch nicht unbedingt nötig. Wer einen tiefen Einblick erhalten möchte, sollte auf die verwiesene Literatur zurückgreifen und sich an die Benutzung eines Statistikprogramms wagen.

- 447 Nachdem die Hypothesen und das Signifikanzniveau festgelegt wurden, wendet man sich der Wahl des Testverfahrens zu, welches für die Fragestellung und die vorhandenen Daten geeignet ist. Bereits die Forschungshypothese grenzt die zu verwendenden Testverfahren ein: Liegt nämlich eine Unterschiedshypothese vor, werden andere Verfahren verwendet, als wenn eine Zusammenhangshypothese vorliegt. Mit einer **Unterschiedshypothese** wird untersucht, ob eine Stichprobe von einer zweiten Stichprobe oder einer theoretischen Verteilung abweicht. Ein klassisches Beispiel ist der Vergleich der mittleren Körpergröße von Männern und Frauen. Liegt eine **Zusammenhangshypothese** vor, wird überprüft, ob die Werte zweier oder mehrerer Stichproben/Variablen in einer Beziehung zueinander stehen. Ein Beispiel für einen solchen Zusammenhang sind die Schuhgröße und die Körpergröße. Im nächsten Abschnitt wenden wir uns zuerst dem Testen von Unterschieden zu, um dann im Anschluss ausführlicher auf das Testen von Zusammenhängen einzugehen.

#### a. Testen von Unterschieden

- 448 Bei dem Testen von Unterschieden muss man sich zunächst auf die Art des Unterschiedes festlegen. Am häufigsten werden Unterschiede in der **zentralen Tendenz** (Mittelwert oder Median), der **Varianz** oder der **Verteilungsform** getestet. Zusätzlich muss bei der Wahl des Testverfahrens berücksichtigt werden, ob es sich bei den Stichproben um **abhängige** oder **unabhängige Stichproben** handelt. Abhängige Stichproben liegen zum Beispiel vor, wenn die Personen in beiden Stichproben dieselben sind. Solche Stichproben werden zum Beispiel häufig bei medizinischen Studien verwendet, um die Wirkung von Behandlungen oder Medikamenten zu ermitteln. In solchen Fällen besteht die erste Stichprobe aus Messwerten (z.B. Blutdruck) vor der Behandlung und die zweite Stichprobe aus Messwerten nach der Behandlung.
- 449 Oft ist es unpraktisch oder schlichtweg nicht möglich, mit abhängigen Stichproben zu arbeiten. Möchte man z.B. Unterschiede im Haarvolumen bei Männern zwischen 18 und 50 Jahren untersuchen, wäre es sehr umständlich, die erste Messung bei 18-Jährigen durchzuführen, um dann 32 Jahre bis zur zweiten Messung verstreichen zu lassen. Stattdessen würde

man auf zwei zufällige Stichproben zurückgreifen, eine mit 18-jährigen und eine mit 50-jährigen Männern. Da diese Stichproben in keiner direkten Beziehung zueinander stehen, werden sie als unabhängig betrachtet. Im Folgenden werden wir zuerst Unterschiede in der zentralen Tendenz testen, danach Unterschiede in der Varianz und zum Schluss Unterschiede in der Verteilungsform.

**aa. Unterschiede in der zentralen Tendenz.** In diesem Abschnitt werden wir zuerst das Testen abhängiger Stichproben und dann das Testen unabhängiger Stichproben anhand von Beispielen illustrieren. Nehmen wir an, dass 2001 im Zuge der Neuregelung des § 23 Abs. 1 lit. a StVO (das sogenannte Handyverbot beim Autofahren) ein Gutachten über den Einfluss von Gesprächen mit Mobiltelefonen auf die Reaktionszeit erstellt werden sollte. Mit Hilfe einer experimentellen Studie könnte man den Einfluss testen und mit Teilnehmern in einem Fahrsimulator Reaktionstests durchführen. Zuerst misst man dabei die Reaktionszeit, während sich die Teilnehmer voll auf die Straße konzentrieren, und danach, während sie dabei ein Telefonat mit ihrem Mobiltelefon führen. Da jeweils die Zeiten für dieselben Teilnehmer gemessen werden, handelt es sich um zwei abhängige Stichproben. Die folgende Tabelle gibt die gemessenen Zeiten der zehn Teilnehmer des Experiments in Millisekunden wieder. 450

| Teilnehmer               | Reaktionszeit des Teilnehmers in Millisekunden |     |     |     |     |     |     |     |     |     |
|--------------------------|------------------------------------------------|-----|-----|-----|-----|-----|-----|-----|-----|-----|
|                          | 1                                              | 2   | 3   | 4   | 5   | 6   | 7   | 8   | 9   | 10  |
| <b>Ohne Mobiltelefon</b> | 391                                            | 368 | 427 | 402 | 372 | 409 | 438 | 408 | 402 | 458 |
| <b>Mit Mobiltelefon</b>  | 446                                            | 498 | 496 | 487 | 426 | 403 | 523 | 440 | 474 | 454 |

Tabelle 7.2: Reaktionszeit mit und ohne Benutzung eines Mobiltelefons

Die Tabelle zeigt bereits, dass die Reaktionszeiten ohne Mobiltelefonat tendenziell kürzer sind als solche mit (mittlere Reaktionszeit ohne Mobiltelefon 407,5 Millisekunden, mittlere Reaktionszeit mit Mobiltelefon 464,7 Millisekunden). Jedoch gibt es auch Versuchsteilnehmer, die kürzere Reaktionszeiten hatten, wenn sie telefonierten (Teilnehmer 6 und 10). Liegt also ein systematischer Unterschied zwischen den Reaktionszeiten vor? Als Ausgangshypothese wird angenommen, dass die Reaktionszeit mit Mobiltelefon kürzer oder gleich der Reaktionszeit ohne Mobiltelefon ist. Die Alternativhypothese lautet dementsprechend, dass die Reaktionszeit ohne Mobiltelefon kürzer ist. Da die Hypothesen gerichtet sind, handelt es sich um einen einseitigen Test. Als zu erreichendes Signifikanzniveau wird das 1 %-Niveau festgelegt. Das Statistikprogramm liefert uns 451

für den Test einen p-Wert von  $p=0,006$ .<sup>13</sup> Daher kommt die Analyse zu dem Ergebnis, dass im Mittel Gespräche mit Mobiltelefonen die Reaktionszeit am Steuer hoch signifikant verlängern.

- 452 Ein Beispiel für unabhängige Stichproben sind Löhne von zufällig ausgewählten Männern und Frauen. Die folgende Tabelle gibt die Nettoarbeitslöhne von 8 Frauen und 8 Männern in € wieder. Beide Stichproben wurden zufällig aus den Daten des sozioökonomischen Panels für das Jahr 2008 gezogen.

|               |         |         |         |         |         |         |         |         |
|---------------|---------|---------|---------|---------|---------|---------|---------|---------|
| <b>Frauen</b> | € 2.400 | € 240   | € 1.300 | € 1.250 | € 1.730 | € 1.200 | € 1.500 | € 1.630 |
| <b>Männer</b> | € 2.290 | € 1.350 | € 3.540 | € 5.000 | € 1.600 | € 1.400 | € 1.400 | € 1.800 |

Tabelle 7.3: Nettoarbeitslöhne von zufälligen ausgewählten Frauen und Männern

Die Mittelwerte der zwei Stichproben unterscheiden sich stark (Frauen: 1406,25; Männer: 2297,5), allerdings ist die Standardabweichung innerhalb der Männergruppe (1314,2) wesentlich größer als die der Frauen (608,2). Liegt also ein systematischer Unterschied zwischen den beiden Stichproben vor, und gibt es einen signifikanten Unterschied zwischen den mittleren Nettoarbeitslöhnen? Getestet wird die Hypothese, dass es keinen Unterschied gibt, gegen die Alternativhypothese, dass es einen Unterschied gibt. Die Ausgangshypothese soll mit einer Signifikanz von 5 % abgelehnt werden. Da das Statistikprogramm für diesen Test einen p-Wert von  $p=0,11$  errechnet, wird die Alternativhypothese abgelehnt.<sup>14</sup> Es liegt also auf Grundlage dieser Daten kein statistisch signifikanter Unterschied zwischen den mittleren Nettoarbeitslöhnen vor.

- 453 Bei der Beurteilung statistischer Ergebnisse muss berücksichtigt werden, dass die Testverfahren nur der Absicherung dienen, dass die beobachteten Unterschiede nicht durch reinen Zufall entstanden sind. Es handelt sich dabei nie um eine Aussage über die Effektgröße, die Relevanz des Ergeb-

<sup>13</sup> In diesem Beispiel wurde der Wilcoxon-Vorzeichenrang-Test für abhängige Stichproben angewandt. Dieser Test ist einer der gängigsten für Probleme mit abhängigen Stichproben, jedoch gibt es einige Alternativen: Wurden die beide Stichproben aus normal verteilten Populationen gezogen, würde der t-Test für abhängige Messungen verwendet. Falls die Verteilungsformen unbekannt oder verschieden sind, könnte man den Fisher-Pitman Permutationstest für abhängige Stichproben anwenden.

<sup>14</sup> Um diese und ähnliche Frage zu beantworten, wird bei unabhängigen Stichproben häufig der Mann-Whitney-u-Test (auch Wilcoxon-Mann-Whitney-Test genannt) verwendet. Auch hier finden sich jedoch Alternativen: Wurden die beiden Stichproben aus normal verteilten Populationen gezogen, würde der t-Test für unabhängige Messungen verwendet, und falls die Verteilungsformen unbekannt oder verschieden sind, könnte man den Fisher-Pitman-Permutationstest für unabhängige Stichproben anwenden.

nisses oder des Unterschiedes. Stellen wir uns vor, dass wir statt der obigen Nettoarbeitslöhne zufällig die folgenden gezogen hätten:

|               |         |         |         |         |         |         |         |         |
|---------------|---------|---------|---------|---------|---------|---------|---------|---------|
| <b>Frauen</b> | € 2.400 | € 2.400 | € 2.400 | € 2.400 | € 2.400 | € 2.400 | € 2.400 | € 2.400 |
| <b>Männer</b> | € 2.405 | € 2.405 | € 2.405 | € 2.405 | € 2.405 | € 2.405 | € 2.405 | € 2.405 |

Tabelle 7.4: Fiktive Nettoarbeitslöhne von Frauen und Männern

Man erkennt sofort, dass der Unterschied zwischen den zwei Stichproben wesentlich geringer ist als im ersten Beispiel. Alle Frauen verdienen € 2.400, daher beträgt der mittlere Nettoarbeitslohn der Frauen € 2.400. Alle Männer verdienen € 2.405 und der Mittelwert entspricht somit ebenfalls € 2.405. Testet man diese Unterschiede, erhält man einen signifikanten Unterschied auf dem 1 %-Niveau. Dieses Ergebnis liegt nahe, denn dass dieser systematische Unterschied zufällig entsteht, ist sehr unwahrscheinlich. Die Frage, ob dieser Unterschied substantiell relevant und wichtig ist, kann ein statistischer Test dagegen nicht beantworten. Diese Interpretation ist vielmehr dem Forscher überlassen, der eine entsprechende Studie durchführt.

**bb. Unterschiede in der Varianz.** Neben Unterschieden in der zentralen Tendenz kann auch die Untersuchung von Unterschieden in der Varianz eine interessante Fragestellung sein. So kann man die Varianz des politischen Spektrums in verschiedenen Ländern miteinander vergleichen oder die Varianz in der Identifikation mit dem eigenem Land. Je niedriger diese Varianz ist, desto homogener ist die Identifikation mit dem Heimatland.<sup>15</sup> Die Varianz spielt aber auch in Situationen, in denen es um Genauigkeit und Präzision geht, eine große Rolle. Wenden wir uns dafür dem einfachen Beispiel eines Joghurt-Produzenten zu. Insgesamt besitzt dieser Produzent 20 Abfüllstationen, welche mit zwei unterschiedlichen Verfahren Joghurt abfüllen. Der Hersteller verkauft den Joghurt in 200-Gramm-Bechern, und im Durchschnitt füllen beide Verfahren diese Menge ab. Jedoch unterliegen die abgefüllten Mengen bei beiden Verfahren Schwankungen, weshalb der Hersteller mögliche Beschwerden von Verbrauchern abwenden und nur das Verfahren verwenden möchte, welches die niedrigsten Schwankungen bei der abgefüllten Menge Joghurt verursacht. Daher wiegt er bei jeder der 20 Abfüllstationen (wobei 10 Stationen das eine, 10 Statio-

<sup>15</sup> Beispielsweise untersuchen *Hassin/Ferguson/Shidlovski/Gross*, Subliminal exposure to national flags affect political thought and behavior, *Proceedings of the National Academy of Sciences of the United States of America*, 104 (2007), 19757 ff. die Auswirkungen von Nationalflaggen auf die Identifikation mit der eigenen Nation. Ein wichtiger Schwerpunkt ihrer Auswertung liegt dabei auf der Veränderung der Varianz.

nen das andere Abfüllverfahren benutzen) die abgefüllte Menge eines Bechers. Dies ergibt die folgenden Werte (in Gramm):

| Verfahren 1 | Verfahren 2 |
|-------------|-------------|
| 200,0366    | 199,7506    |
| 200,0482    | 200,3987    |
| 199,9556    | 200,0686    |
| 200,0431    | 199,7906    |
| 199,8577    | 199,8459    |
| 199,8774    | 199,8718    |
| 200,1188    | 200,0559    |
| 199,9589    | 200,0209    |
| 200,1025    | 199,6047    |
| 200,1093    | 199,9751    |

Tabelle 7.5: Abgefüllte Menge Joghurt mit Verfahren 1 und 2

Die im Mittel abgefüllten Mengen sind sich bei beiden Verfahren sehr ähnlich: 200,01 Gramm bei Verfahren 1 und 199,93 Gramm bei dem zweiten Verfahren. Diese Werte unterscheiden sich nicht signifikant in der zentralen Tendenz.<sup>16</sup> Beide Verfahren haben geringe, aber unterschiedliche Standardabweichungen. Bei dem ersten Verfahren beträgt die Standardabweichung 0,09 Gramm und bei dem zweiten 0,22 Gramm. Der Produzent testet die gewogenen Mengen auf einen Unterschied bezüglich der Varianz mit der Nullhypothese, dass es keinen Unterschied gibt. Die Nullhypothese soll mit einem mindestens 5%igen Signifikanzniveau abgelehnt werden. Das Statistikprogramm ermittelt einen p-Wert von  $p=0,019$ . Daher wird die Nullhypothese verworfen und der Produzent rüstet alle Abfüllanlagen auf das erste Verfahren um.<sup>17</sup>

- 455 cc. **Unterschiede in der Verteilungsform.** Eine Frage, welche sich regelmäßig bei der Arbeit mit empirischen Daten stellt, ist, wie die Daten verteilt sind. Liegen genügend Beobachtungen vor bzw. hat man eine große repräsentative Stichprobe, kann man sich die Verteilung als Graphik darstellen lassen und eine andere Verteilung damit vergleichen. Wie bereits im Abschnitt über deskriptive Statistik besprochen wurde, sind die Bruttoerwerbseinkommen in Deutschland rechtsschief und nicht normal verteilt.

<sup>16</sup> Zu überprüfen zum Beispiel mit dem *Wilcoxon-Mann-Whitney-Test*.

<sup>17</sup> Bei diesem Beispiel wurde der Varianztest nach Brown/Forsythe angewandt. Eine gängige Alternative ist der *Siegel-Tukey-Test*. Werden nicht Unterschiede in der Varianz getestet, sondern versucht, die Varianz zu erklären, finden die Methodiken der Varianzanalyse, auch ANOVA (analysis of variance) genannt, Anwendung.

Abb. 7.10 gibt noch einmal die Verteilung aus der repräsentativen Stichprobe des sozioökonomischen Panels wieder (durchgezogene Kurve) und gleichzeitig, wie eine Normalverteilung aussehen müsste (gestrichelte Kurve). Man erkennt sofort, dass diese zwei Verteilungen voneinander abweichen.

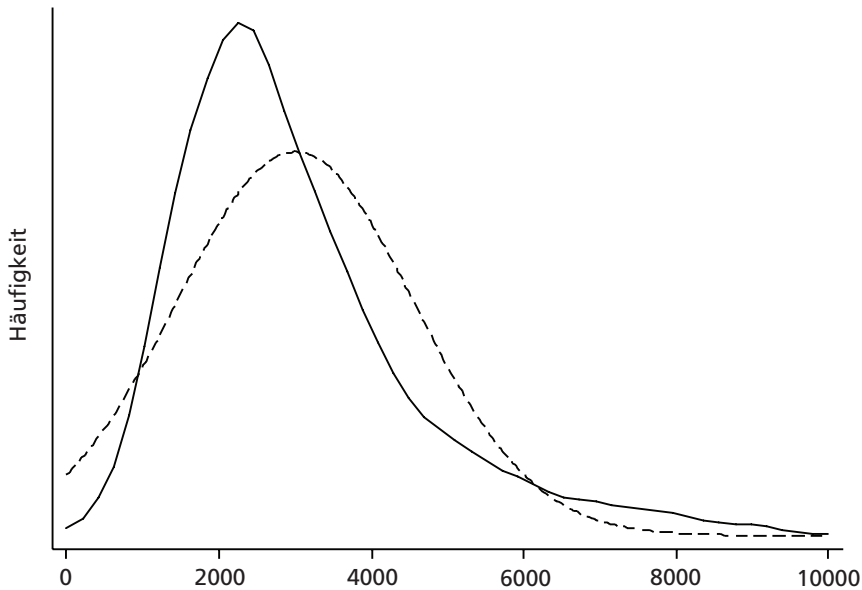


Abbildung 7.10: Verteilung aus der repräsentativen Stichprobe des sozioökonomischen Panels (durchgezogen) im Vergleich zu einer Normalverteilung (gestrichelt)

Bei großen Stichproben kann man dies oft schon auf den ersten Blick erkennen, bei kleineren Stichproben haben zufällige Ausreißer in den Daten einen höheren Einfluss. Daher wendet man auf diese häufiger Tests an, die auf die Feststellung von Unterschieden in der Verteilungsform angelegt sind.

Nehmen wir an, wir hätten nicht die Daten des sozioökonomischen Panels, sondern nur die Bruttoerwerbseinkommen von 14 zufällig ausgewählten Personen.<sup>18</sup> Die folgende Tabelle 7.6 gibt diese Einkommen wieder:

456

<sup>18</sup> Die folgenden Werte sind zufällig von 14 Personen im sozioökonomischen Panel ausgewählt worden.



|         |         |         |         |         |         |         |
|---------|---------|---------|---------|---------|---------|---------|
| € 4.701 | € 1.800 | € 2.000 | € 2.540 | € 1.400 | € 2.800 | € 6.660 |
| € 2.400 | € 2.953 | € 2.500 | € 2.288 | € 1.800 | € 5.600 | € 2.100 |

Tabelle 7.6: Bruttoerwerbseinkommen 14 zufällig ausgewählter Personen

Wir möchten nun mit Hilfe eines Verteilungstests diese Stichprobe gegen eine Normalverteilung testen. Unser Statistikprogramm ermittelt uns einen p-Wert von  $p < 0.01$  und wir können daher die Testhypothese, dass unsere Stichprobe zufällig aus einer normal verteilten Population gezogen wurde, ablehnen.<sup>19</sup> Bei diesem Beispiel wurde eine Stichprobe gegen eine theoretische Verteilung getestet. Es ist jedoch auch möglich, die Verteilung zweier Stichproben direkt gegeneinander zu testen.<sup>20</sup> Eine typische Fragestellung wäre, ob sich die Verteilung der Löhne bei Frauen von der bei Männern unterscheidet.

- 457 Bei Tests auf Unterschiede in der Verteilungsform ist zu beachten, dass es nur möglich ist festzustellen, dass die Verteilungen sich unterscheiden. Eine Aussage über die Art des Unterschiedes ist nicht möglich, denn diese Tests treffen keine Aussage darüber, ob der Unterschied in der Varianz, der zentralen Tendenz oder der Verteilungsklasse liegt.

## b. Testen von Zusammenhängen

- 458 Beim Testen von Zusammenhängen wird zwischen ungerichteten Zusammenhängen und gerichteten Zusammenhängen differenziert. **Ungerichtete Zusammenhänge** werden mit Hilfe von Korrelationen untersucht, wobei keine Aussage über eine Kausalität getroffen wird, sondern nur über gemeinsames Auftreten. Im Gegensatz dazu geht man bei **gerichteten Zusammenhängen** von einer Kausalität aus. Ein gerichteter Zusammenhang wäre z.B., dass ein höherer Bildungsabschluss später im Berufsleben zu einem höheren Einkommen führt. Diese gerichteten Zusammenhänge werden mit Hilfe von **Regressionen** untersucht. Allerdings müssen auch bei Regressionen bestimmte Regeln beachtet werden, um Schlussfolgerungen über die Kausalität ziehen zu können.<sup>21</sup> Im Folgenden werden zuerst Korrelationen näher erläutert und später wird dann detaillierter auf Regressionen eingegangen.

<sup>19</sup> Für das obige Beispiel wurde der Shapiro-Wilk-Test für Normalverteilungen verwendet. Eine häufig verwendete Alternative ist der Kolmogorov-Smirnov-Test für einzelne Stichproben.

<sup>20</sup> In diesem Fall könnte man bei stetigen Daten den *Kolmogorov-Smirnov*-Test für zwei Stichproben anwenden sowie bei diskreten Daten den *Pearson-Ch<sup>2</sup>*-Test oder den *Epps-Singleton*-Omnibus-Test.

<sup>21</sup> Hierzu Rz. 397 ff.

aa. **Ungerichteter Zusammenhang zwischen zwei Variablen: Korrelationen.** Häufig ist man an Zusammenhängen zwischen zwei gemessenen Merkmalen (z.B. Körpergröße und Schuhgröße) interessiert. Liegen positive oder negative Zusammenhänge zwischen zwei Merkmalen vor, ist von einer **Korrelation** die Rede. 459

Korrelationen lassen sich häufig optisch bereits durch Graphen erkennen, die die positiven oder negativen Zusammenhänge sichtbar machen. Die folgenden Punktwolken bzw. Streudiagramme (Abb. 7.11) verdeutlichen dies:

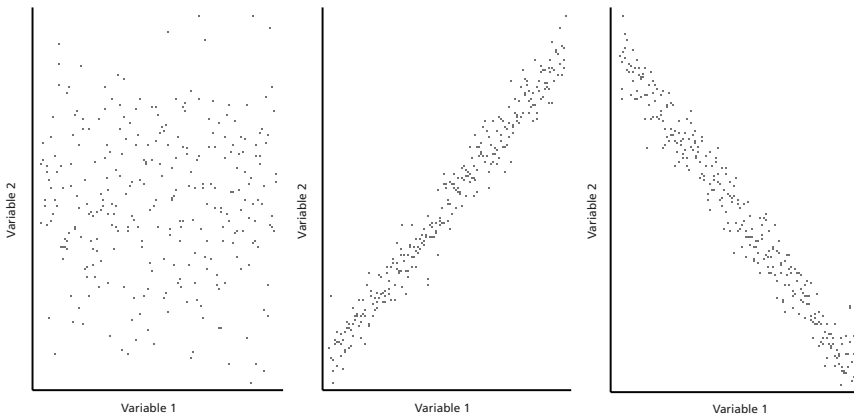


Abbildung 7.11: Punktwolke ohne Korrelation (links), mit positiver Korrelation (mittig) und negativer Korrelation (rechts)

Ein genaues Maß für eine Korrelation ist der **Korrelationskoeffizient  $r$** , 460 welcher den (linearen) Zusammenhang zwischen zwei Merkmalen widerspiegelt. Korrelationskoeffizienten können Werte zwischen -1 und 1 annehmen. Ist  $r > 0$ , wird von einer positiven Korrelation gesprochen, bei  $r < 0$  von einer negativen Korrelation. Hat eine Korrelation den Koeffizienten  $r = 0$ , spricht man von zwei unkorrelierten Merkmalen, was bedeutet, dass es keinen Zusammenhang zwischen den beiden Merkmalen gibt. Je mehr sich der Korrelationskoeffizient der -1 oder der 1 annähert, desto mehr liegen alle Punkte auf einer Geraden. In den obigen Punktwolken beträgt der Korrelationskoeffizient für die positive Korrelation  $r = 0,98$  und für die negative Korrelation  $r = -0,97$ . Statistikprogramme geben neben dem Korrelationskoeffizienten auch das Signifikanzniveau an. Wie bei den Testverfahren aus dem vorherigen Kapitel wird damit die Irrtumswahrscheinlichkeit für das fehlerhafte Annehmen eines (linearen) Zusammenhanges gemessen. Je höher die Signifikanz, desto robuster ist das Ergebnis.

461 Betrachten wir ein Beispiel für eine solche Korrelation zwischen zwei quantitativen Merkmalen:<sup>22</sup> Häufig wird in den Medien von einem Zusammenhang zwischen Arbeitslosigkeit und rechtsextremen Gewalttaten berichtet. Die folgenden Abbildungen 7.12 und 7.13 geben die Arbeitslosenquote und die Anzahl rechtsextremer Gewalttaten einmal als Balkendiagramm und einmal als Punktwolke für jedes Bundesland wieder.

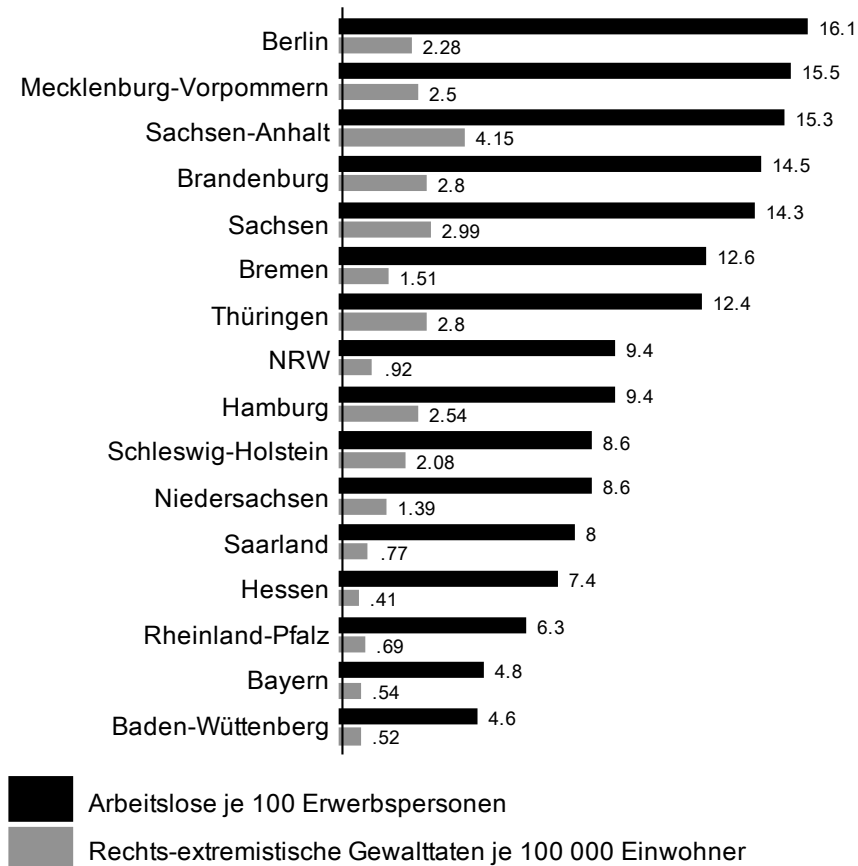


Abbildung 7.12: Arbeitslosenquote und die Anzahl rechtsextremer Gewalttaten in den Bundesländern als Balkendiagramm

<sup>22</sup> Wie bei den Tests, welche auf Unterschiede testen, gibt es auch hier unterschiedliche Verfahren je nachdem, ob die Variablen qualitativ oder quantitativ sind, und je nachdem, welche Charakteristik der Zusammenhang erfüllt (z.B. linear oder monoton). An dieser Stelle sei wieder auf den Entscheidungsbaum von *Vorberg/Blankenberger*, Die Auswahl statistischer Tests und Maße, Psychologische Rundschau 50 (1999), 157ff. verwiesen.

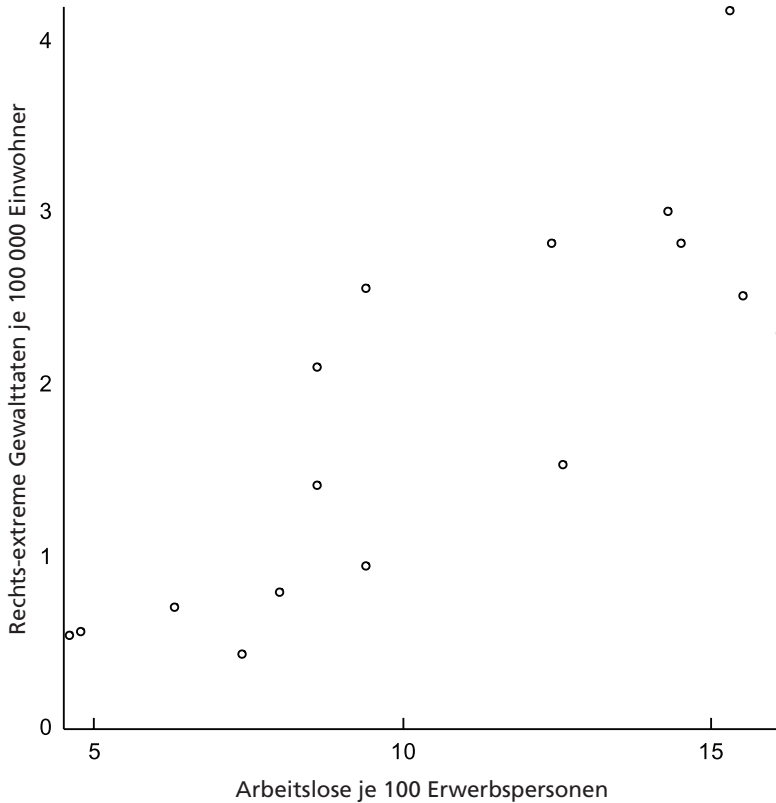


Abbildung 7.13: Arbeitslosenquote und die Anzahl rechtsextremer Gewalttaten in den Bundesländern als Streudiagramm

Durch die einheitliche Skala für Arbeitslose und rechte Gewalttaten erscheinen die Unterschiede bezüglich der Gewalttaten recht gering. Trotzdem ist eine leichte Tendenz erkennbar. Ein besserer Überblick über den Zusammenhang zwischen zwei Variablen (hier Arbeitslosigkeit und rechte Gewalt) bietet das Streudiagramm. Die Punktwolke lässt bereits einen Zusammenhang erahnen. Überprüft man den Zusammenhang mit einem Statistikprogramm, ergibt sich für die obigen Werte ein Korrelationskoeffizient von  $r=0.82$  und ein hohes Signifikanzniveau von 1 %.

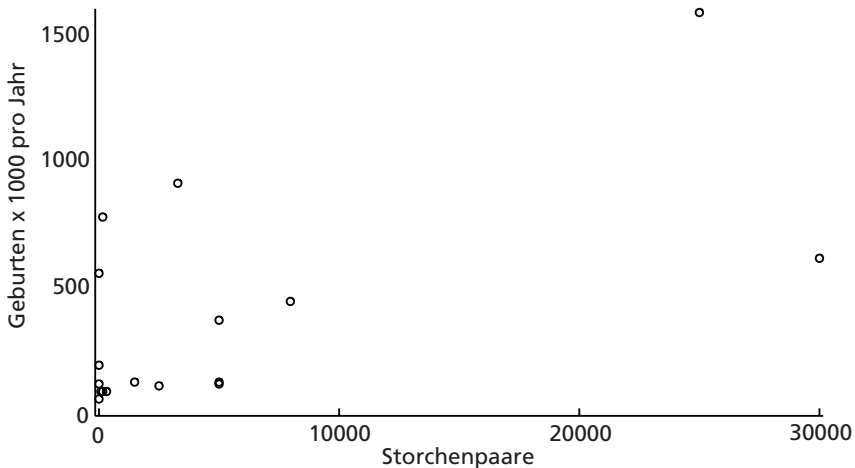
Es scheint also einen hoch signifikanten und stark positiven Zusammenhang zwischen Arbeitslosigkeit und rechten Gewalttaten zu geben. Das bedeutet, dass höhere Arbeitslosigkeit und eine höhere Quote an rechten Gewalttaten in unserer Stichprobe tendenziell gemeinsam auftreten. Jedoch ist, wie wir bereits gesehen haben, wichtig zu beachten, dass Korre-

462

463

lationen im Prinzip nur Aussagen über gemeinsames Auftreten und nicht über Kausalität treffen.<sup>23</sup> Daher ist auch in unserem Beispiel eine Aussage über Kausalität nicht möglich. Diese kann hier in zwei Richtungen laufen: Zum einen kann es sein, dass Personen auf Grund rechter Gewalt arbeitslos sind, zum anderen, dass die Arbeitslosigkeit der Grund für mehr rechts-extreme Gewalttaten ist.

- 464 Teilweise sind unterschiedliche Faktoren auch dann miteinander korreliert, wenn zwischen beiden keine direkte Kausalität besteht. Es handelt sich um das bereits angesprochene Problem der Scheinkorrelation, bei dem ein dritter Faktor die gemeinsame Ursache für die beiden miteinander korrelierten Faktoren ist. Ein gängiges Beispiel hierfür ist der Zusammenhang zwischen brütenden Storchenpaaren und Geburtenraten beim Menschen. Die folgende Abbildung 7.14 gibt die Kombination aus der Anzahl der Geburten beim Menschen und der Anzahl von Storchenpaaren für 17 Länder wieder.



Beispiel und Daten: Robert Matthews (2000)

Abbildung 7.14: Geburtenrate beim Menschen und die Anzahl von Storchenpaaren in 17 Ländern

Daten von Robert Matthews (Teaching Statistics, 22(2), 2000)

- 465 Das Streudiagramm legt einen positiven Zusammenhang zwischen der Anzahl der Geburten und der Anzahl der Störche nahe – und tatsächlich, die beiden Variablen sind auf dem 1 %-Niveau signifikant und positiv miteinander korreliert ( $r=0.62$ ). Es scheint also einen Zusammenhang zu ge-

<sup>23</sup> Hierzu Rz. 397.

ben zwischen der Anzahl von Storchepaaren und der Zahl menschlicher Geburten. Wie kommt es zu diesem Zusammenhang? Im konkreten Fall ist er auf ein drittes, intervenierendes Merkmal zurückzuführen. Beide Faktoren (Geburten und Storchepaare) sind jeweils mit diesem dritten Merkmal korreliert. Für das obige Beispiel ist es die Fläche des jeweiligen Landes. Je größer das Land, umso mehr Menschen leben dort und umso mehr Platz ist auch für Störche. Daher finden wir dort mehr Menschen, die Kinder gebären, und gleichzeitig mehr Störche. Wir werden uns diesem Beispiel gleich unten noch einmal zuwenden.

**bb. Gerichteter Zusammenhang zwischen Variablen: Regressionen.** Die 466  
bisher vorgestellten statistischen Testverfahren haben entweder Unterschiede oder Abhängigkeit von Variablen/Stichproben analysiert. Regressionsanalysen, welche zu den wichtigsten Werkzeugen der Ökonometrie gehören, gehen einen Schritt weiter. In der empirischen Wirtschaftsforschung werden Regressionsmodelle verwendet, um Wirkungszusammenhänge (Kausalität) zu analysieren und um Modelle zu testen. Die Ökonometrie verbindet somit die ökonomische Theorie (das theoretische Modell) mit der beobachteten Realität (den beobachteten Daten).

Bei Regressionen unterscheidet man zwischen der abhängigen Variable 467  
(die Variable, die man erklären möchte) und den unabhängigen Variablen (auch erklärende Variablen genannt). Die abhängige Variable wird durch die unabhängigen Variablen erklärt und ist daher von diesen abhängig. Mit Regressionen kann ermittelt werden, ob eine erklärende Variable überhaupt Einfluss auf die abhängige Variable hat (Ursachenanalyse), wie die abhängige Variable sich bei einer Änderung der unabhängigen Variable verändert (Wirkungsprognose) und wie sich die abhängige Variable im Zeitverlauf ändert (Zeitreihenanalyse). Typische Fragestellungen, die mit Hilfe von Regressionsmodellen beantwortet werden, lauten:

- Wie verändert sich die Nachfrage nach einem Produkt, wenn dessen Preis um € 1 steigt?
- Wie verändert sich die Anzahl der begangenen Verbrechen innerhalb eines Stadtteiles, wenn dort mehr Polizisten eingesetzt werden?
- Hängt das Einkommen eines Arbeitnehmers von dessen Bildung, Alter und Geschlecht ab?

Als Beispiel nehmen wir im Folgenden den Preis eines gebrauchten Autos. Die Tabelle 7.7 spiegelt die Preise von einigen Gebrauchtwagen sowie deren Ausstattungsmerkmalen wider.

| Tür | Pferdestärken | Kilometer | Euro   |
|-----|---------------|-----------|--------|
| 3   | 75            | 52.000    | 4.499  |
| 3   | 90            | 47.000    | 5.980  |
| 3   | 75            | 45.000    | 5.980  |
| 3   | 75            | 39.990    | 5.999  |
| 3   | 90            | 37.570    | 6.250  |
| 3   | 90            | 44.640    | 6.590  |
| 3   | 75            | 36.000    | 6.598  |
| 3   | 102           | 37.500    | 6.700  |
| 3   | 90            | 40.400    | 6.799  |
| 3   | 75            | 41.270    | 6.980  |
| 5   | 90            | 33.000    | 7.450  |
| 5   | 75            | 35.000    | 7.450  |
| 3   | 75            | 35.000    | 7.500  |
| 5   | 75            | 39.800    | 7.900  |
| 5   | 90            | 27.690    | 8.450  |
| 3   | 102           | 37.540    | 8.490  |
| 5   | 105           | 33.000    | 8.950  |
| 5   | 75            | 23.000    | 8.980  |
| 5   | 116           | 32.000    | 9.290  |
| 5   | 90            | 36.550    | 9.890  |
| 5   | 105           | 36.000    | 9.999  |
| 5   | 116           | 30.100    | 11.200 |
| 5   | 105           | 28.400    | 11.950 |

Tabelle 7.7: Preise von Gebrauchtwagen

- 468 Grundlage einer Regression ist immer die Spezifikation eines Modells. Dabei werden die abhängige Variable, die unabhängigen Variablen, mögliche Störgrößen sowie der funktionale Zusammenhang definiert. Nehmen wir an, dass der Preis eines Gebrauchtwagens ausschließlich von den gefahrenen Kilometern abhängt. Diese Abhängigkeit soll linear sein, d.h. dass der Preis mit jedem gefahrenen Kilometer gleich stark sinkt. Unser ökonometrisches Modell für den Preis eines Autos  $X$  sähe dann folgendermaßen aus:

$$\text{Preis Gebrauchtwagen}_x = \alpha + \beta_1 \cdot \text{Anzahl Kilometer}_x + \gamma_x$$

- 469 Auf der linken Seite steht die Variable, die erklärt werden soll, nämlich der Preis für das Auto  $X$  (die abhängige Variable). Auf der rechten Seite finden wir die Variable, welche den Preis erklären soll – hier die Anzahl der Kilometer, welche das Auto  $X$  bereits gefahren wurde (die unabhängige Variable). Beta ist der Koeffizient, welcher später mit Hilfe der Regression geschätzt wird. Er reflektiert den Einfluss der erklärenden Variable, also um

wie viel der Preis eines Autos sinkt, wenn ein Kilometer gefahren wurde. Alpha ist eine Konstante. Sie gibt den geschätzten Wert eines Autos wieder, wenn es 0 Kilometer gefahren wurde. Gamma ist eine Störgröße. Diese Störgröße ist zufällig und spiegelt z.B. zu hohe oder zu niedrige Preisvorstellungen des Verkäufers wider. Es wird angenommen, dass sich diese Abweichungen im Mittel ausgleichen. Wichtig ist die Annahme, dass Alpha und Beta für alle Fahrzeuge gleich sind, und sich nur die Anzahl der gefahrenen Kilometer sowie die Störgrößen zwischen den Autos unterscheiden.<sup>24</sup>

Eine lineare Regression mit nur einer erklärenden Variable kann graphisch sehr leicht veranschaulicht werden. Abb. 7.15 gibt den Preis und die gefahrenen Kilometer für die Gebrauchtwagen aus Tabelle 7.7 als Punktwolke wieder. Das obige Regressionsmodell schätzt einfach die Gerade, welche diese Punktwolke am besten beschreibt. Die Steigung der Geraden ist dabei der Koeffizient Beta und die Konstante ist der Wert, bei dem die y-Achse von der Geraden geschnitten wird.

470

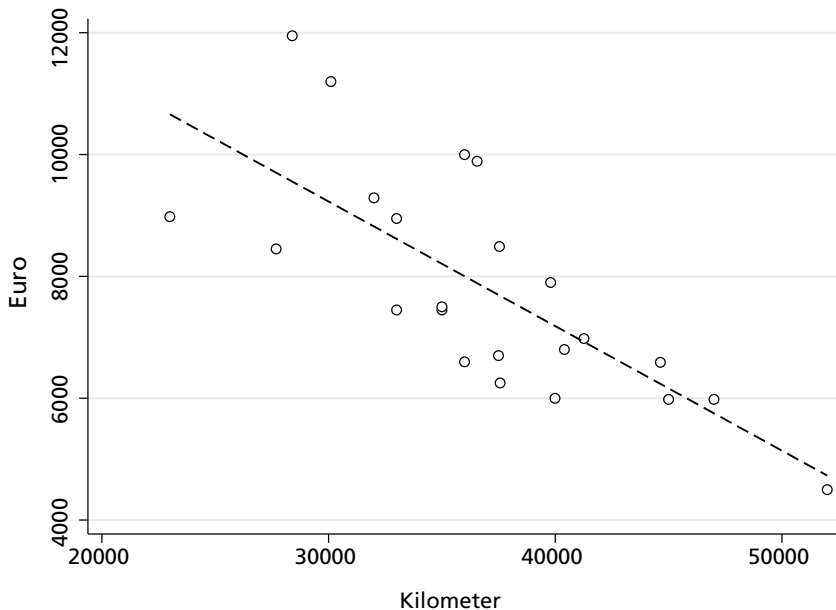


Abbildung 7.15: Streudiagramm über die gefahrenen Kilometer und den Kaufpreis von Gebrauchtwagen

<sup>24</sup> Weitere Annahmen sind unter anderem, dass es keine Korrelationen zwischen erklärenden Variablen und Störgrößen gibt sowie dass die Störgrößen eine konstante Varianz haben und normal verteilt sind.



- 471 Als Ergebnis einer Regression geben Statistikprogramme neben den Koeffizienten und der Konstanten noch weitere wichtige Maße an. Das sind zum einen Gütemaße für die Regression und zum anderen die Signifikanzen der Regressionskoeffizienten. Zu den Gütemaßen gehören das **Bestimmtheitsmaß** ( $R^2$ ) und die **F-Statistik**. Das **Bestimmtheitsmaß** misst, wie gut die geschätzte Regressionsfunktion die empirischen Daten abbildet. Es wird also bewertet, wie gut die Regression die Daten beschreibt, auf denen die Schätzung der Regression basiert. In dem obigen Beispiel beschreibt das Bestimmtheitsmaß, wie viel der Streuung in den Preisen durch die Zahl der gefahrenen Kilometer erklärt wird. Der Wert des Bestimmtheitsmaßes ist normiert und liegt immer zwischen 0 und 1. Je höher der Wert, desto mehr wird die vorhandene Streuung durch die Regressionsgleichung erklärt. Wenn die gesamte Streuung erklärt wird, ist  $R^2 = 1$ , und die Daten werden vollständig durch die Koeffizienten der erklärenden Variablen bestimmt. Wenn die vorhandene Streuung nicht durch die erklärenden Variablen erklärt werden kann, ist  $R^2 = 0$ . Möchte man mit den ermittelten Koeffizienten der erklärenden Variable nicht nur die vorliegenden Daten beschreiben, sondern auch Prognosen über die vorhandenen Daten hinweg anstellen, sollte das  $R^2$  recht groß sein. Die Güte der Prognose hängt jedoch nicht nur davon ab, wie gut die vorhandene Stichprobe beschrieben wird, sondern auch davon, wie groß diese Stichprobe ist. Die **F-Statistik** berücksichtigt dies und testet, ob generell die Regressionskoeffizienten (also der Einfluss der unabhängigen Variablen) signifikant verschieden von Null sind. Ist der F-Test signifikant, bedeutet das, dass generell ein statistisch robuster Zusammenhang zwischen mindestens einer erklärenden Variable und der abhängigen Variable besteht.
- 472 Ist der F-Test signifikant, bedeutet das jedoch nicht, dass alle Koeffizienten statistisch signifikant sind. Um den Einfluss der einzelnen erklärenden Variable zu überprüfen, werden die einzelnen Regressionskoeffizienten mithilfe der **t-Statistik** getestet. Der t-Test überprüft für jede einzelne unabhängige Variable, ob sie in einem signifikanten Zusammenhang mit der abhängigen Variable steht. Bei einem signifikanten t-Test eines Koeffizienten beeinflusst die unabhängige Variable tatsächlich die abhängige Variable.
- 473 Die Ergebnisse von Regressionen werden üblicherweise in Tabellenform dargestellt. Die folgende Tabelle 7.8 gibt die Ergebnisse dreier Regressionsmodelle wieder. Modell 1 ist das Modell aus dem obigen Beispiel, das zweite Modell ergänzt das erste um die Anzahl der Pferdestärken als erklärende Variable und das dritte Modell ergänzt das zweite um die Anzahl der Türen. Formal ausgedrückt sehen dann die drei Modelle wie folgt aus:

Modell 1: Preis Gebrauchtwagen =  $\alpha + \beta_1 \cdot \text{Kilometer}_x + \gamma_x$

Modell 2: Preis Gebrauchtwagen =  $\alpha + \beta_1 \cdot \text{Kilometer}_x + \beta_2 \cdot \text{PS} + \gamma_x$

Modell 3: Preis Gebrauchtwagen =  $\alpha + \beta_1 \cdot \text{Kilometer}_x + \beta_2 \cdot \text{PS} + \beta_3 \cdot \text{Türen} + \gamma_x$

| Variable         | Model 1   |     | Model 2   |     | Model 3   |     |
|------------------|-----------|-----|-----------|-----|-----------|-----|
| Anzahl Kilometer | -0,2045   | *** | -0,163    | *** | -0,0984   | *** |
|                  | (0,0402)  |     | (0,0349)  |     | (0,0406)  |     |
| Anzahl PS        |           |     | 56,198    | *** | 48,973    | *** |
|                  |           |     | (16,4937) |     | (15,001)  |     |
| Anzahl Türen     |           |     |           |     | 659,05    | *** |
|                  |           |     |           |     | (266,17)  |     |
| Konstante        | 15364,61  |     | 8809,97   |     | 4468,852  |     |
|                  | (1504,15) |     | (2281,19) |     | (2686,88) |     |
| p-Wert F-Test    | 0,0000    |     | 0,0000    |     | 0,0000    |     |
| R <sup>2</sup>   | 0,5525    |     | 0,7169    |     | 0,7860    |     |

\*  $p < 0.1$ , \*\*  $p < 0.05$ , \*\*\*  $p < 0.01$ , Standardfehler in Klammern

Tabelle 7.8: Lineare Regression Preis Gebrauchtwagen

Zwar sehen solche Tabellen kompliziert aus, aber man sollte sich nicht von ihnen abschrecken lassen. Zuerst werfen wir einen Blick auf den p-Wert der F-Tests und erkennen, dass er für alle drei Modelle hoch signifikant ist. Wir wissen also, dass zumindest einer der Koeffizienten auf jeden Fall den Preis beeinflusst. Dann suchen wir nach Sternen in den einzelnen Modellen. Hierbei stehen die Sterne für Signifikanzniveaus des t-Tests. Je mehr Sterne, desto sicherer können wir uns sein, dass die erklärende Variable wirklich einen Einfluss auf den Wert des Gebrauchtwagens hat. Allerdings variiert die in den Studien für unterschiedliche Signifikanzniveaus verwandte Anzahl der Sterne. Die Bedeutung ist jedoch meistens in der Legende erklärt.

Nun wenden wir uns dem Vorzeichen der Koeffizienten zu. Im dritten Modell ist der Koeffizient für Kilometer negativ, während er für Türen und Pferdestärken positiv ist. Das bedeutet, dass der Preis eines Gebrauchtwagens mit der Anzahl der gefahrenen Kilometer abnimmt, gleichzeitig aber mit der Anzahl der Türen und der Anzahl der PS zunimmt. Da alle Koeffizienten signifikant sind, kann davon ausgegangen werden, dass auch tatsächlich ein solcher Einfluss existiert. In den Klammern wird der Standardfehler des Regressionskoeffizienten wiedergegeben. Er bildet die Standardabweichung des Koeffizienten ab und hängt wie der Koeffizient von der genauen Modellspezifikation ab.

- 476 Bei unserem Beispiel ist jeder neu eingeführte Koeffizient signifikant und bleibt es auch. Dies bedeutet, dass wichtige erklärende Variablen dem Modell hinzugefügt werden und dadurch die Modelle genauer werden. Dass die Modelle genauer werden, erkennt man zusätzlich an dem „ $R^2$ “, welches wir uns nun im letzten Schritt ansehen. Mit jeder neuen Variable verbessert sich das Bestimmtheitsmaß, und die vorhandene Streuung wird umso genauer durch das Modell beschrieben. Das dritte Modell ist am besten dazu geeignet, den Preis eines Gebrauchtwagens vorherzusagen, da es ein höheres Bestimmtheitsmaß hat und mehr signifikante, erklärende Variablen enthält.
- 477 Aber auch bei Regressionsanalysen sind fehlerhafte Anwendungen und Interpretationsfehler möglich. Genau wie bei Korrelationen können Scheinkorrelationen auch bei Regressionen zu problematischen Interpretationen führen. Wenden wir uns dem Beispiel der Scheinkorrelation zwischen Störchen und Geburten zu, welches wir bereits diskutiert haben. Abbildung 7.14 gibt die Anzahl der Geburten beim Menschen sowie die Anzahl der Störchenpaare für 17 Länder wieder. Statt einer einfachen Korrelation möchten wir nun eine Regression mit diesen Daten durchführen. Die Tabelle gibt die Ergebnisse zweier linearer Regressionen wieder. Im ersten Modell wird die Anzahl der Geburten nur durch die Anzahl der Störchenpaare pro Land erklärt, im zweiten wurde die Größe des jeweiligen Landes in Quadratkilometern als erklärende Variable ergänzt. Modell 1 deutet auf einen signifikanten Zusammenhang zwischen der Anzahl der Störche und den Geburten pro Jahr hin. Zwar scheint der Effekt recht klein zu sein, aber das Signifikanzniveau ist hoch. Bei Modell 1 liegt – wie bei der einfachen Korrelation zwischen der Anzahl der Störche und der Anzahl der Geburten – das Problem einer Scheinkorrelation vor. Wenn nun aber die dritte (intervenierende) Variable, mit der jeweils die Geburten und die Anzahl der Störche korrelieren, hinzugefügt wird, verschwindet die Signifikanz für die Anzahl der Störche als erklärende Variable und der Koeffizient wird sehr klein. Des Weiteren erhöht sich die Genauigkeit des Modells,  $R^2$  steigt stark an, und der p-Wert des F-Testes verbessert sich.<sup>25</sup>

---

<sup>25</sup> Neben Scheinkorrelationen gibt es weitere Stolperfallen bei Regressionsanalysen, z.B. Multikollinearität (hohe Korrelation zwischen den erklärenden Variablen) oder Heteroskedastizität (Probleme mit unterschiedlichen Varianzen in der Stichprobe). Für weiterführende Erklärungen sei auf *Wooldridge, Introductory Econometrics*, 2013 verwiesen.

| Variable                 | Model 1              | Model 2                |
|--------------------------|----------------------|------------------------|
| Anzahl Störche           | 0,029 ***<br>(0,009) | -0,006<br>(0,006)      |
| Größe in km <sup>2</sup> |                      | 0,0015 ***<br>(0,0002) |
| Konstante                | 225,03<br>(93,56)    | -7,411<br>(56,702)     |
| p-Wert F-Test            | 0,0079               | 0,0000                 |
| R <sup>2</sup>           | 0,3847               | 0,8622                 |

\*  $p < 0.1$ , \*\*  $p < 0.05$ , \*\*\*  $p < 0.01$ , Standardfehler in Klammern

Werte des Beispiels aus Robert Matthews (Teaching Statistics, 22(2), 2000)

Tabelle 7.9: Lineare Regression Anzahl Geburten pro Jahr

Natürlich gibt es noch wesentlich mehr Anwendungsgebiete für Regressionen, mit denen dann wiederum weitere Probleme verbunden sind. Beispielsweise können erklärende und erklärte Variable in **nicht-linearen Zusammenhängen** stehen, sie können auf einen bestimmten Wertebereich beschränkt sein oder Wahrscheinlichkeiten abbilden. Auf Grund dieser Komplexität ist die Ökonometrie eine eigene Wissenschaft innerhalb der Wirtschaftswissenschaften. Wer sich näher für dieses Thema interessiert, sei auf die eingangs des Kapitels aufgeführte Literatur verwiesen.

478



## § 8 – Verhaltensökonomik

**Literatur:** R. Axelrod, *Die Evolution der Kooperation*, 7. Aufl. 2009; S. Bowles/H. Gintis, *Homo Reciprocans*, *Nature* 415 (2002), 125–128; C. Camerer/G. Loewenstein/M. Rabin (Hg.), *Advances in Behavioral Economics*, 2003; C. Camerer/S. Issacharoff/G. Loewenstein/T. O'Donoghue/M. Rabin, *Regulation for conservatives: Behavioral economics and the case for 'asymmetric paternalism'*, *University of Pennsylvania Law Review* 151 (2003), 1211–1254; C. Engel/M. Englerth/J. Lüdemann/I. Spiecker gen. Döhmman (Hg.), *Recht und Verhalten*, 2000; E. Fehr/S. Gächter, *Fairness and Retaliation: The Economics of Reciprocity*, *Economic Perspectives* 14 (2000), 159–181; G. Gigerenzer/R. Selten, *Bounded Rationality: The Adaptive Toolbox*, 2001; C. Guthrie/J. Rachlinski/A. Wistrich, *Inside the Judicial Mind*, *Cornell Law Review* 86 (2001), 777–830; C. Jolls/C. Sunstein/R. Thaler, *A Behavioral Approach to Law and Economics*, *Stanford Law Review* 50 (1998), 1471–1550; D. Kahan, *The Logic of Reciprocity: Trust, Collective Action, and the Law*, *Michigan Law Review* 95 (2003), 2477–2497; D. Kahneman, *Thinking, Fast and Slow*, 2012; D. Kahneman/A. Tversky, *Prospect Theory: An Analysis of Decision Under Risk*, *Econometrica* 47 (1979), 263–291; D. Kahneman/A. Tversky (Hg.), *Choices, Values and Frames*, 1982; A. Kemmerer/C. Möllers/M. Steinbeis/G. Wagner, *Choice Architecture in Democracies. Exploring the Legitimacy of Nudging*, 2016; R. Korobkin/T. Ulen, *Law and Behavioral Science: Removing the Rationality Assumption from Law and Economics*, *California Law Review* 88 (2000), 1051–1144; M. Rabin, *Economics and Psychology*, *American Economics Review* 83 (1998), 11–46; M. Schweitzer, *Kognitive Täuschungen vor Gericht: eine empirische Studie*, 2005; C. Sunstein (Hg.), *Behavioral Law and Economics*, 2000; R. Thaler, *Quasi Rational Economics*, 1991; R. Thaler/C. Sunstein, *Nudge: Wie man kluge Entscheidungen anstößt*, 2010.

### I. Verhaltenstheorie in der Ökonomie

Der „rationale Investor“, wie wir ihn in traditionellen Ökonomielehrbüchern antreffen, ist ein vorsichtiger Zeitgenosse. Er kennt die Finanzmärkte in- und auswendig und verfolgt nüchtern seinen aufgeklärten Eigennutz, indem er die potenziellen Risiken eines jeden Finanzprodukts, das er kauft oder verkauft, sorgfältig abwägt. Wenn wir das Verhalten in der jüngsten

479

Finanzkrise an diesem klassischen Bild messen, werden die Herausforderungen, vor denen das Modell des *homo oeconomicus* steht, nur zu deutlich: Die Nachrichtensendungen haben uns als Beobachter des Geschehens auf den Finanzmärkten konfrontiert mit draufgängerischen Akteuren, die von Gier, Maßlosigkeit und bisweilen Angst getrieben waren und die zu Hochzeiten fieberhaft komplexe Finanzprodukte handelten, die sie oft kaum verstanden. Mit diesen Einsichten erschütterte die Finanzkrise von 2007–08 nicht nur die weltweiten Märkte und das Finanzsystem, sondern auch die Wirtschaftswissenschaften. Die Wirklichkeit schien einige ihrer grundlegendsten Annahmen fundamental in Frage zu stellen, was zu einer veritablen Krise der Disziplin selbst führte. Doch jede Krise hat ihre Profiteure: Der Sturz der konventionellen Wirtschaftswissenschaft führte zu einer Hausse für jene Ökonomen, die einige der Grundannahmen der Disziplin schon seit Langem in Frage stellten.

480 Die Kritik galt zuvörderst dem klassischen ökonomischen Verhaltensmodell,<sup>1</sup> das freilich schon lange vor der Finanzkrise kritischen Anfragen ausgesetzt war.<sup>2</sup> Eine überzeugende Kritik am Verhaltensmodell der Rechtsökonomik muss belegen, dass (1) die Voraussagen des ökonomischen Paradigmas und seines Verhaltensmodells im Regelfall, also für die Mehrheit aller Individuen, falsch sind und dass (2) diese Abweichungen für das Institutionendesign wesentlich sind. Will sie aber nicht rein destruktiv sein, muss sie (3) mit einer Alternative aufwarten. Bliebe ihre einzige Einsicht, dass Menschen unberechenbar sind und nicht-rational handeln, so wäre dies zwar ein Sargnagel für das traditionelle Verhaltensmodell, aber ohne jeden Nutzen für eine an Steuerung interessierte Rechtswissenschaft. Erst die Systematik einer Abweichung vom Rationalmodell macht sie für die Zwecke ökonomischer Modellierung brauchbar. Insofern muss eine Alternative zur rationaltheoretischen Rechtsökonomik über die Behauptung von Nicht-Rationalität hinausgehen und ein eigenes, zu Voraussagen fähiges Verhaltensmodell liefern.

481 Genau dies ist das Ziel einer relativ neuen interdisziplinären Strömung, die als *Behavioral Economics* (oder **Verhaltensökonomik**) zu beträchtlicher Prominenz gelangt ist. Die Ambition der Vertreter dieses Ansatzes erschöpft sich nicht in einer Kritik des Rationalmodells. *Behavioral Economics* beansprucht vielmehr, eine Alternative zur traditionellen Ökonomik darzustellen. Statt deren Annahmen einfach zu akzeptieren, versucht der neue Ansatz, sie empirisch zu testen und durch Einsichten aus der Ver-

<sup>1</sup> Vgl. dazu oben Rz. 69 ff.

<sup>2</sup> Vgl. dazu oben Rz. 80 ff.

haltenswissenschaft zu modifizieren und zu ergänzen. Als wichtigster Verbündeter diene ihnen dabei die kognitive Psychologie. *Daniel Kahneman* wurde 2002 als erster Psychologe mit dem Ökonomie-Nobelpreis ausgezeichnet, Aufsätze von Verhaltensökonomern erscheinen in den führenden ökonomischen Fachblättern und prominente Vertreter der traditionellen Ökonomik beginnen wie selbstverständlich, ihre Modelle mit psychologischen Einsichten anzureichern. Die Verhaltensökonomik ist mittlerweile innerhalb der Ökonomik kein Orchideenfach mehr, sondern fest im disziplinären Mainstream verankert. Mit *Behavioral Finance*, der psychologisch informierten Analyse von Finanz- und Kapitalmärkten, hat sie bereits eine eigene, einflussreiche Subdisziplin entwickelt.

Das – bisherige – Ergebnis der Bemühungen um eine Alternative zur traditionellen Ökonomik und zum Rationalmodell sieht menschliches Verhalten weniger als konsequente Rationalität, aber mehr als zufällige Nicht-Rationalität: *Richard Thaler* hat dafür den Ausdruck „**Quasi-Rationalität**“<sup>3</sup> geprägt. Andere sprechen vielleicht noch prägnanter von „**berechenbarer Irrationalität**“<sup>4</sup>. 482

Es war nur eine Frage der Zeit, bis die sich auch die Rechtswissenschaft für den neuen Ansatz zu interessieren begann. Ein von *Cass Sunstein* herausgegebener Sammelband belegte erstmals die Vielfalt möglicher juristischer Anwendungen der Verhaltensökonomik und gab dieser neuen Form der ökonomischen Analyse des Rechts mit *Behavioral Law and Economics* (zu Deutsch etwa **Verhaltensrechtsökonomik**) einen Namen.<sup>5</sup> Inzwischen bieten die meisten amerikanischen Rechtsfakultäten Lehrveranstaltungen dazu an. *Behavioral Law and Economics* macht heute der traditionellen ökonomischen Analyse auf allen Gebieten ernsthafte Konkurrenz. Potenzial und mögliche Schwächen dieses alternativen Verhaltensmodells für die Rechtsökonomik sollen im Folgenden diskutiert werden. 483

## II. Methodische und konzeptionelle Grundlagen

*Behavioral Law and Economics* beinhaltet schon sprachlich drei Komponenten: (1) eine verhaltenswissenschaftliche, (2) eine ökonomische und (3) eine juristische. Was diese jeweils ausmacht und in welchem Verhältnis sie zueinander stehen, soll im Folgenden erläutert werden. 484

<sup>3</sup> *Thaler*, Quasi Rational Economics, 1991.

<sup>4</sup> *Ariely*, Predictably Irrational: The Hidden Forces that Shape Our Decisions, 2008.

<sup>5</sup> *Sunstein* (Hg.), Behavioral Law and Economics, 2000.



## 1. Verhaltenswissenschaftliche Komponente

- 485 Das verhaltenswissenschaftliche Element unterscheidet schon nominell den neuen Ansatz von der herkömmlichen ökonomischen Analyse des Rechts. Gleichwohl ist der Begriff „behavioral“ etwas ungünstig gewählt, denn er weckt Assoziationen zum psychologischen Behaviorismus nach *Burrhus Frederic Skinner*. Tatsächlich ist aber ziemlich genau das Gegenteil gemeint. Für den Behaviorismus war der menschliche Kopf eine *Black Box* und damit undurchdringlich. Dagegen ist die Verhaltensökonomik ein Kind der kognitiven Wende: Gerade diejenigen Vorgänge, welche der Behaviorismus ausblenden wollte, stehen für sie im Mittelpunkt des Interesses.
- 486 Entsprechend war die dominierende Verhaltenswissenschaft für aufgeschlossene Ökonomen lange Zeit die kognitive Psychologie. Dafür gibt es verschiedene Gründe. Einer besteht sicherlich darin, dass diese psychologische Subdisziplin Sachverhalte erhellt, über die bereits die Urväter der modernen Ökonomik wie *Adam Smith* spekulierten. Zum anderen zeichnen sich ihre Befunde durch eine größere Präzision aus als etwa die der *Freud'schen* Psychoanalyse, der die meisten Ökonomen wenig abgewinnen konnten. Am wichtigsten dürfte indes ein dritter Grund sein: Die **kognitive Psychologie** liefert meist Befunde, die für alle Menschen Gültigkeit beanspruchen können. Der grundlegende Bauplan unserer Gehirne ist derselbe. Insofern kann es nicht überraschen, dass die meisten Abweichungen vom Rationalmodell, welche die kognitive Psychologie dokumentiert hat, sich bei der großen Mehrzahl aller Menschen in derselben Weise wiederfinden. Sie sind also systematisch und damit in gewissem Maße vorhersagbar. Mit dieser Systematik ist das Fundament für ein relativ allgemeingültiges „quasi-rationales“ Verhaltensmodell gelegt.
- 487 Zunehmend wird der Begriff der Verhaltenswissenschaft extensiv verstanden. Manche Autoren beginnen etwa die Neurowissenschaften mit einzubeziehen. Von einem verbesserten Verständnis des Schaltplans unseres Gehirns versprechen sie sich tiefere Einsichten in die von der Psychologie entdeckten Effekte und ihre Zusammenhänge. Daneben versuchen andere, die Soziologie für ökonomische Fragen fruchtbar zu machen und dem *homo oeconomicus* gleichsam eine soziale Dimension jenseits der spieltheoretisch-strategischen Interaktion zu verleihen.<sup>6</sup> Wie groß der Nutzen dieser Erweiterungen ist, muss sich freilich noch zeigen.

---

<sup>6</sup> Vgl. etwa *Kahan*, Between Economics and Sociology: The new path of Deterrence, *Michigan Law Review* 95 (1997), 2477 ff.

## 2. Ökonomische Komponente

*Behavioral Law and Economics* ist nicht gleich *Law and Psychology* – genauso wenig wie die (bisherige) Verhaltensökonomik „nur“ Psychologie ist. Tatsächlich steckt noch jede Menge Ökonomik in dem neuen Ansatz. Zunächst einmal geben auch Verhaltensökonomien die Suche nach Gleichgewichten und effizienten Lösungen nicht grundsätzlich verloren, auch wenn sie zu anderen Ergebnissen kommen mögen. Selbiges gilt auch für die ökonomische Methodik: Methodologischer Individualismus,<sup>7</sup> mathematische Formalisierung der Annahmen und logische Deduktion von Ergebnissen – den Wert dieser Vorgehensweisen stellen nur wenige grundsätzlich in Frage. Richtig ist zwar, dass die traditionell beliebteste Methode der Verhaltensökonomik sicherlich die experimentelle ist; im Entscheidungslabor wird den Teilnehmern ein Problem vorgelegt, das sie bearbeiten müssen, ohne in ihren üblichen soziokulturellen Kontext eingebettet zu sein. Daneben bleiben aber auch alle anderen Methoden der Ökonomik zulässig:<sup>8</sup> Felddaten bleiben ebenso relevant wie Feldexperimente oder Computersimulationen.

Ferner erscheint die Rationalitätsannahme des Standardansatzes nach wie vor in den meisten Arbeiten der Maßstab zu sein, von welchem aus erst „Anomalien“ festgestellt werden können. Zwar stellen Vertreter beider Ansätze dies in Frage, doch können ihre Argumente nicht überzeugen. Orthodoxe Anhänger des Rationalmodells verweisen gern darauf, dass viele scheinbare Anomalien in Wahrheit zumindest „global“ rational seien. Selbst wenn dies zutrifft, übersieht dieser Einwand, dass der traditionelle Ansatz allein niemals in der Lage gewesen wäre, diese Effekte überhaupt zu identifizieren. Umgekehrt hoffen viele Verhaltensökonomien, irgendwann eine allgemeinere Verhaltenstheorie präsentieren zu können, als deren bloßer Spezialfall die rationale Wahl erschiene. Trotz einiger unbestreitbarer Fortschritte ist es, zumindest bis zum gegenwärtigen Zeitpunkt, bei dieser Hoffnung geblieben. Die Verhaltensökonomik lässt sich also nicht als kritische Fußnote zum Standardansatz abtun, ist aber auch nicht so eigenständig, wie es einige ihrer Vertreter gern hätten. *Mark Kelman* hat das Verhältnis sehr treffend als „rhetorisches Duett“ bezeichnet.<sup>9</sup> Prosaischer könnte man sagen, dass die Verhaltensökonomik eine (tiefgehende) Modifikation des Standardansatzes ist.

<sup>7</sup> Hierzu Rz. 63.

<sup>8</sup> Hierzu ausführlicher § 7.

<sup>9</sup> *Kelman*, Behavioral Economics as Part of a Rhetorical Duet: A Response to Jolls, Sunstein, and Thaler, Stanford Law Review 50 (1998), 1577 ff.

### 3. Juristische Komponente

- 490 Eine Einsicht ist in ihrer Wichtigkeit kaum zu überschätzen: Die kognitive Psychologie, welche die Verhaltensökonomik empirisch unterfüttert, ist keine normative Disziplin. Ihr geht es darum, zu verstehen, wie das menschliche Gehirn arbeitet. Ob das Resultat gut oder schlecht ist, interessiert sie nicht. Entsprechend kann auch die von der kognitiven Psychologie inspirierte Verhaltensökonomik nur eine positive Theorie menschlichen Verhaltens liefern. Ihren Erkenntnissen ist freilich bei der normativen Analyse Rechnung zu tragen. Für sich genommen haben sie jedoch keinen normativen Charakter. Verhaltensökonomien identifizieren Abweichungen von den Postulaten des Rationalmodells und setzen diesem ein eigenes positives Modell entgegen. Ob die Devianzen indes gut oder schlecht sind, ist eine davon verschiedene Wertungsfrage.
- 491 Ein Fach, das sich mit Wertungsfragen blendend auskennt, ist die Rechtswissenschaft. Insofern spielen Juristen bei der Anwendung verhaltensökonomischer Einsichten eine wichtige Rolle. Tatsächlich ist, wie wir sehen werden, ihr Wertungsspielraum gegenüber der traditionellen ökonomischen Analyse sogar erweitert. Letztere bildet nämlich ein in sich geschlossenes System: Autonomes Marktverhalten und effiziente Ergebnisse sind gleichsam zwei Seiten einer Medaille. Staatliche Intervention ist nur dann angezeigt, wenn hohe Transaktionskosten den freien Austausch auf dem Markt beeinträchtigen.<sup>10</sup> Tatsächlich aber – so eine Einsicht der Verhaltensökonomik – können viele kognitive Phänomene ganz ähnlich wie Transaktionskosten wirken. Sie alle durch staatliche Regulierung zu neutralisieren, ist weder wünschenswert noch praktisch machbar. Dem Rechtswissenschaftler stellt sich folglich ein Selektionsproblem: Welche Effekte sind so schädlich, dass sie nach paternalistischer Regulierung verlangen? Welche sind neutral oder gar positiv? Und wie ist zu verfahren, wenn ein Effekt ambivalent ist, also positive wie negative Effekte für ein oder mehrere Individuen haben kann?
- 492 Juristen müssen in diesem neuen Wertungsdickicht die normative Axt anlegen. Vielfach – dies ist eine weitere These von *Behavioral Law and Economics* – haben sie das bereits getan. Die klassische ökonomische Analyse des Rechts behauptete oft, dass das *Common Law* bereits implizit nach größtmöglicher Effizienz strebe.<sup>11</sup> Ähnlich besagt die deskriptive Analyse der Vertreter von *Behavioral Law and Economics*, dass diverse Rechtsinstitute und Institutionen als intuitive Reaktion auf die von ihnen

<sup>10</sup> Vgl. dazu oben Rz. 149 ff.

<sup>11</sup> Posner, *The Economics of Justice*, 1983, S. 4.

präzisierten Einsichten über die Beschränktheit menschlicher Rationalität verstanden werden müssten.<sup>12</sup> Im Folgenden werden nun einige prominente Einsichten der Verhaltensökonomik skizziert und rechtliche Anwendungsbeispiele diskutiert. Diese Beispiele entstammen überwiegend aus dem straf(prozess)rechtlichen Bereich. Hierauf ist die Relevanz von *Behavioral Law and Economics* aber nicht begrenzt. Die einschlägige Literatur hat mindestens ebenso viele Anstöße für das Zivilrecht und das öffentliche Recht geliefert.<sup>13</sup>

### III. Einzelne Einsichten der Verhaltensökonomik und ihre Bedeutung für das Recht

Jolls, Sunstein und Thaler sprechen in ihrem grundlegenden Aufsatz zur Methodik von *Behavioral Law and Economics* von den „three bounds“ – den drei Beschränkungen –, welche das neue psychologisch informierte Verhaltensmodell vom traditionellen Rationalmodell unterscheiden. Sie nennen diese (1) beschränktes Eigeninteresse (*bounded self-interest*), (2) beschränkte Rationalität (*bounded rationality*) sowie (3) beschränkte Willensstärke (*bounded willpower*).<sup>14</sup> Andere Autoren haben diese – freilich nicht zwingende – Trias aufgegriffen, und auch die nachfolgende Darstellung wird sich an ihr orientieren. 493

#### 1. Begrenztes Eigeninteresse

Francis Edgeworth erhob das Eigeninteresse in einem berühmten Zitat zum „ersten Prinzip der Ökonomie“. Damit ist gemeint, dass der *homo oeconomicus* seinen Vorteil innerhalb der bestehenden Restriktionen optimal zu verwirklichen sucht.<sup>15</sup> Die Definition lässt indes im Vagen, wie dieser Vorteil beschaffen sein mag. Die Ökonomen machen hierzu bewusst keine materiellen Aussagen. Auch wenn sie sich in ihren Modellen häufig 494

<sup>12</sup> Jolls/Sunstein/Thaler, A Behavioral Approach to Law and Economics, Stanford Law Review 50 (1998), 1471 (1508 ff.).

<sup>13</sup> Allgemeine Überblicke geben Jolls/Sunstein/Thaler, ebd.; Korobkin/Ulen, Law and Behavioral Science: Removing the Rationality Assumption from Law and Economics, California Law Review 88 (2000), 1053 ff.; Englerth, Behavioral Law and Economics: Eine kritische Einführung, in: Engel *et al.* (Hg.), Recht und Verhalten, 2007, 60 ff.

<sup>14</sup> Jolls/Sunstein/Thaler, A Behavioral Approach to Law and Economics, Stanford Law Review 50 (1998), 1471 (1476).

<sup>15</sup> Zum Modell des *homo oeconomicus* bereits Rz. 69 ff.

auf den materiellen bzw. monetären Vorteil konzentrieren, ist diese Einschränkung doch keineswegs zwingend. Dass Menschen auch daran interessiert sein können, Respekt oder Ansehen zu erlangen oder Missbilligung zu meiden, hat nie jemand ernstlich bestritten. Der Dieb stiehlt, weil er die Diebesbeute höher schätzt als ein ruhiges Gewissen, und die Nonne betet, weil ihr eine gute Beziehung zu ihrem Gott den größten Nutzen stiftet. Beide maximieren ihren **Erwartungsnutzen** (*expected utility*) gemäß der in der Entscheidungstheorie üblichen Formel:

$$EU = p_1 U(x_1) + p_2 U(x_2) + \dots + p_n U(x_n)$$

495  $U(x_1)$  bis  $U(x_n)$  bezeichnen hierbei den Einzelnutzen bestimmter unsicherer Ereignisse, die jeweils mit den Wahrscheinlichkeiten  $p_1$  bis  $p_n$  eintreten. Dieser Nutzen bestimmt sich nicht objektiv, sondern subjektiv, also gemäß der individuellen Nutzenfunktion  $U$  des Entscheiders, die zudem auch seine Risikoeinstellung ausdrückt. Ein und dasselbe Ergebnis  $x$  kann also von

496 zwei Menschen ganz unterschiedlich bewertet werden. Steht dann aber nicht jedes beobachtete Verhalten im Einklang mit der Rationaltheorie? Dass sie angesichts einer solchen Beliebigkeit nicht in den Teufelskreis fehlender Falsifizierbarkeit gerät, liegt an zweierlei Dingen: Erstens gibt eine Reihe von sogenannten **Rationalitätsaxiomen** den möglichen Präferenzen einen strikten Rahmen vor. Jede Entscheidung muss nicht nur für sich als rationale Wahl interpretierbar sein, sondern darf auch nicht im Widerspruch mit den in anderen Entscheidungen bekundeten Präferenzen treten. Zweitens machen Ökonomen konventionell eine weitere Einschränkung: Danach sucht der *homo oeconomicus* tatsächlich nur seinen eigenen Vorteil. Der Nutzen anderer hat folglich keinen Einfluss auf seinen Nutzen. **Soziale Präferenzen** sprengen zwar nicht grundsätzlich den Rahmen der Rationaltheorie, werden aber implizit regelmäßig ausgeschlossen. Der *homo oeconomicus* ist folglich weder altruistisch noch missgünstig. Er interessiert sich wenig für seine Mitmenschen. Dass die meisten Menschen in Wirklichkeit nicht ganz so abgeklärt sind, verrät einem bereits die Alltagserfahrung. Psychologen und Ökonomen haben in einer Reihe von Experimenten diese Erfahrung zu präzisieren versucht.<sup>16</sup>

497 Das bekannteste wird als **Ultimatum-Spiel** bezeichnet; wir haben es in § 2 ausführlich beschrieben.<sup>17</sup> Zur Erinnerung: Spieler 1 erhält die Aufgabe, einen bestimmten Geldbetrag zwischen sich und einem zweiten Spieler aufzuteilen. Spieler 2 kann das Angebot annehmen oder ablehnen.

<sup>16</sup> Hierzu Rz. 207 ff., 252 ff.

<sup>17</sup> Vgl. dazu oben Rz. 82.

Nimmt er an, wird die von Spieler 1 vorgeschlagene Aufteilung vorgenommen. Lehnt er aber ab, gehen beide Spieler leer aus. Ebenso einfach ist die Voraussage, welche die Rationaltheorie für diese Situation tätigen würde. Ein rationaler Spieler 1 würde den geringstmöglichen Betrag anbieten, um das Maximum für sich zu behalten. Ein rationaler Spieler 2 müsste dieses Angebot annehmen, da auch der kleinste Betrag besser ist als überhaupt nichts. Tatsächlich lässt sich das von *Rational Choice* vorhergesagte Verhalten im Ultimatum-Spiel jedoch nahezu nie beobachten. Angebote unter einem Drittel der Gesamtsumme werden von Spieler 2 regelmäßig zurückgewiesen – sie verletzen damit das Eigeninteresse-Postulat. Die meisten Teilnehmer in der Rolle von Spieler 1 scheinen dessen Unwillen, sich mit wenig abspesen zu lassen, bereits zu antizipieren und zeigen sich großzügiger. Eine Fortentwicklung ist das **Diktator-Spiel**: In diesem Spiel entscheiden die Sender einseitig über die Aufteilung des Geldbetrags; die Empfänger erhalten von der Gesamtsumme, was die Sender festlegen. Anders als im Ultimatum-Spiel haben sie nicht die Möglichkeit, das ihnen unterbreitete Angebot abzulehnen. Gleichwohl kommt es auch hier höchst selten zu Angeboten des Minimalbetrags. Anstatt den größtmöglichen Anteil für sich zu reservieren, geben die meisten Spieler in der Anbieterrolle einen beträchtlichen Teil an die Empfänger ab (auch wenn die Angebote hinter denen im Ultimatumspiel etwas zurückbleiben). Damit verhalten sich hier auch die Sender im Widerspruch zur Rationaltheorie.

Was sich aus den beiden Experimenten folgern lässt, ist, dass Menschen 498 sehr häufig nicht nur ihrem eng verstandenen Eigeninteresse folgen, sondern offenbar auch durch **Fairnessnormen** motiviert werden. Deren Herkunft mag unterschiedlich sein. Manche – wie das **Prinzip der Gegenseitigkeit** – mögen evolutionär adaptiv gewesen sein.<sup>18</sup> Dafür sprechen die Befunde aus einem von *Robert Axelrod* organisierten Turnier: In Form einer Simulation ließen Forscher Programme gegeneinander antreten, die jeweils unterschiedliche Strategien verkörperten, die zueinander in Konkurrenz traten. Dabei setzte sich letztlich das nur vierzeilige Programm des Mathematikers *Anatol Rapaport* durch, das auf der einfachen Strategie „Wie du mir, so ich dir“ (*tit for tat*) beruhte. Ganz ähnlich verhalten sich reale Personen im Rahmen des Ultimatum-Spiels.

Andere Fairnessnormen mögen kulturell determiniert sein und sich an 499 einer „Referenztransaktion“ orientieren. Dafür spricht unter anderem die Manipulierbarkeit der Ergebnisse im Diktator- sowie im Ultimatumspiel.

<sup>18</sup> Vgl. *Axelrod*, Die Evolution der Kooperation, 2009; *Gintis*, Game Theory Evolving, 2000.

Suggeriert man etwa Teilnehmer 1, er habe seine Anbieterrolle durch Leistung (z.B. aufgrund der richtigen Beantwortung einer schwierigen Frage) erworben, so fallen seine Offerten deutlich niedriger aus als ohne diese Suggestion. Insgesamt lässt sich festhalten, dass Fairness die klassische Vorstellung des Eigeninteresses unterminiert. Die Realität scheint erheblich komplexer zu sein. Fairnessnormen wirken bisweilen wie eine Steuer auf die Verwirklichung unseres Vorteils. Der Ökonom *Matthew Rabin* hat als Erster versucht, diese Erkenntnis zu formalisieren und in ein ökonomisches Modell zu übertragen.<sup>19</sup> Viele andere sind ihm mittlerweile mit eigenen Modellen gefolgt.

500 Das Gesagte impliziert nicht, dass die Standardannahme eigeninteressierten Verhaltens immer verfehlt wäre. Die zitierten Spiele wurden meist ohne Wiederholung (*one shot*) und unter der Voraussetzung völliger Anonymität der Teilnehmer gespielt. Im Wirtschaftsleben dürften solche Voraussetzungen relativ selten sein. Auch der Volksmund weiß: „Man begegnet sich immer zweimal“. Gleichwohl gibt es auch wirtschaftlich relevante Verhaltensweisen, die sich mittels der Annahme reinen Eigeninteresses nur schwer erklären lassen. So hinterlassen etwa die meisten Menschen Trinkgeld an Autobahnraststätten, die sie mit Sicherheit nie wieder besuchen werden. Wären die oben zitierten Befunde in ihrer Relevanz auf Alltagsanekdoten beschränkt, könnte man sie getrost vernachlässigen. Tatsächlich aber ist ihre Bedeutung für die Rechtswissenschaft nicht zu unterschätzen.<sup>20</sup> Zwei nur kurz angerissene Beispiele mögen dies illustrieren:

501 Dass Menschen nicht nur aus Furcht vor Strafe rechtstreu sind, wie es die ökonomische Standardtheorie behauptet, gilt den meisten (deutschen) Strafrechtlern und Kriminologen als selbstverständlich. Sie glauben an die Macht der **positiven Generalprävention**. Fairnesspräferenzen – insbesondere das Konzept der Gegenseitigkeit – können diese Behauptung empirisch unterfüttern. Solange Menschen glauben, dass sich auch die meisten anderen ihrer Mitmenschen an bestehende Normen halten, sind sie auch selbst dazu bereit, ohne dass es der staatlichen Knute bedürfte. Würde jede Gelegenheit Diebe machen, wäre die öffentliche Ordnung wohl kaum zu garantieren. Diese Logik funktioniert aber nicht nur in eine Richtung. So kann die sichtbare Erosion von Normen in einer Gemeinschaft auch zu dem Gefühl führen, der Ehrliche sei der Dumme. Das Eigeninteresse kann

<sup>19</sup> *Rabin*, Incorporating Fairness into Economics and Game Theory, *American Economic Review* 83 (1993), 1281 ff.

<sup>20</sup> Vgl. ausführlich hierzu *Magen*, Fairness, Eigennutz und die Rolle des Rechts. Eine Analyse auf Grundlage der Verhaltensökonomik, in: Engel *et al.* (Hg.), *Recht und Verhalten*, 2007, 261 ff.

dann wieder ungehemmt aufleben und der Verfallsprozess beschleunigt sich dadurch, dass immer mehr Menschen, die sonst zur bedingten Kooperation bereit wären, diese aufkündigen.

Bisweilen verbietet das Recht an sich effiziente<sup>21</sup> Transaktionen, ohne dass die ökonomische Theorie dies plausibel begründen könnte.<sup>22</sup> Manche interpretieren etwa das Verbot von Schwarzhandel unter Rekurs auf bestimmte Fairnessnormen:<sup>23</sup> Schwarzhandel entspreche letztlich einer Auktion und sei folglich hoch effizient. Auktionen gewährleisten nämlich, dass ein Gut in die Hände desjenigen gelangt, der es am höchsten schätzt – genau dies ist das vom Standardansatz empfohlene Ergebnis. Trotz ihrer Effizienz lehnen allerdings die meisten Menschen diese Art der Zuteilung für einige Güter ab. Stattdessen präferieren sie ein Allokationssystem, in dem etwa derjenige die begehrten Konzertkarten erhält, der am längsten für sie angestanden hat. Fairnessnormen sorgen in solchen Fällen dafür, dass eigentlich effiziente Tauschvorgänge verboten werden und ein alternativer, weniger effizienter Mechanismus sich durchsetzt.

502

## 2. Begrenzte Rationalität

Menschen sind keine Computer: Ihre kognitiven Fähigkeiten sind genauso begrenzt wie ihr Gedächtnis. Diese triviale Beobachtung wollen freilich auch *Rational-Choice*-Ökonomen nicht leugnen. Sie halten das Wissen um unsere **kognitive Beschränktheit** aber auf der Grundlage der **Als-ob-Annahme**<sup>24</sup> für unwesentlich; Menschen mögen keine Computer sein, aber wir können ihr Verhalten vereinfacht so modellieren, als ob ihr Gehirn ähnlich leistungsstark sei wie ein Hochleistungsrechner. Zudem erachten viele rationale Nutzenmaximierung als evolutionär adaptiv – diejenigen, die grob gegen die Rationalitätspostulate verstießen, würden bald von rationaleren Akteuren aus dem Markt gedrängt.<sup>25</sup>

503

Verhaltensökonomien finden keines dieser Argumente überzeugend. Gestützt auf diverse empirische Befunde argumentieren sie, dass im Falle der Rationalitätspostulate unrealistische Annahmen tatsächlich auch zu falschen Voraussagen führten und die Theorie damit an ihren eigenen Maß-

504

<sup>21</sup> Hierzu Rz. 79 f.

<sup>22</sup> Vgl. dazu oben Rz. 47 ff.; vgl. ferner oben Rz. 257.

<sup>23</sup> *Jolls/Sunstein/Thaler*, A Behavioral Approach to Law and Economics, Stanford Law Review 50 (1998), 1471 (1510 ff.).

<sup>24</sup> Dazu oben Rz. 71.

<sup>25</sup> *Alchian*, Uncertainty, Evolution, and Economic Theory, Journal of Political Economy 58 (1950), 211 ff.



staben scheitere. Selbst im Falle von Unternehmen und korporativen Akteuren führe Nicht-Rationalität nachweislich nicht zur sofortigen Insolvenz. Vielmehr sei der Markt zu jeder gegebenen Zeit voller ineffizient wirtschaftender Firmen. Noch viel mehr gelte dieser Einwand für Individuen. Wer nicht-rationale Konsumentenentscheidungen trifft, mag auf niedrigerem Niveau leben, kann dies aber durchaus sehr lange tun.

- 505 Daher fordern Verhaltensökonominnen eine Revision der Annahmen der ökonomischen Verhaltenstheorie. Statt unbegrenzte Rationalität zu unterstellen, sollten Ökonomen von der Psychologie lernen und bestimmte systematische Abweichungen von ihren Postulaten ausdrücklich in ihre Modelle einarbeiten. Die Abweichungen, die in einer Vielzahl von Experimenten belegt wurden, lassen sich grob mittels der aus der empirischen Entscheidungstheorie bekannten Kategorien „Urteilsbildung“ (*judgment*) und „Entscheidung“ (*choice*) einteilen. Urteilsbildung betrifft die kognitiven Vorgänge, die zu einer Vorstellung von Wahrscheinlichkeiten führen. Entscheidung betrifft die Auswahl zwischen verschiedenen Optionen auf dieser Grundlage.

#### a. Urteilsbildung

- 506 Die ökonomische Standardtheorie erhebt Menschen zu intuitiven Statistikern. Der *homo oeconomicus* sammelt zunächst die optimale Menge an Informationen. Kommen neue Informationen hinzu, aktualisiert er sein Wahrscheinlichkeitsurteil gemäß dem **Bayes-Theorem**, d.h. er überführt eine vorhandene *A-priori*-Wahrscheinlichkeit unter Einhaltung bestimmter statistischer Prinzipien in *A-posteriori*-Verteilungen. Diese Theorie bietet durchaus Raum für eine psychologische Unterfütterung, da sie nicht angibt, wie die *A-priori*-Wahrscheinlichkeiten bestimmt werden. Die Standard-Ökonomik lässt diesen Raum allerdings meist unausgefüllt. Zudem sind mit dem Bayes-Theorem einige Annahmen verbunden, deren empirische Validität gering ist. Die Verhaltensökonomik kann in diesem Bereich einen korrigierenden Beitrag leisten. Damit ist indes wie schon im Falle der Fairnessnormen keineswegs zwingend die Aufgabe der mathematisch-formalen Modellierung zugunsten eines rein verbalen *Ad-hoc*-Ansatzes verbunden. Vielmehr scheinen sich sogenannte quasi-bayesianische Modelle durchzusetzen, welche einige der im Folgenden beschriebenen Effekte in den Standardrahmen zu integrieren suchen.<sup>26</sup>

<sup>26</sup> Ein Beispiel für ein solches Modell ist *Rabin/Schrag*, First Impressions Matter: A Model of Confirmatory Bias, *Quarterly Journal of Economics* 114 (1999), 37 ff.

aa. **Heuristiken.** Das wohl bekannteste Beispiel aus dem Urteilsbereich sind die zahlreichen, von Psychologen seit Langem dokumentierten Heuristiken – kognitive Daumenregeln, die uns helfen, uns trotz begrenzter kognitiver Kapazitäten in einer hochkomplexen Welt zurecht zu finden. Die Kehrseite dieser *mental shortcuts* ist jedoch, dass sie in bestimmten Situationen zu Fehlurteilen führen. 507

Ein prominentes Beispiel dafür stellt die **Verfügbarkeitsheuristik** (*availability heuristic*) dar.<sup>27</sup> Daniel Kahneman und Amos Tversky hatten Teilnehmer eines Experiments gebeten zu schätzen, wie viele Wörter eines Textabschnitts auf die Silbe „ing“ endeten. Dabei erhielten sie regelmäßig deutlich höhere Schätzungen, als wenn sie nach der Häufigkeit von Worten fragten, deren zweitletzter Buchstabe ein „n“ ist. Dies ist widersprüchlich, denn der zweite Fall ist ein Unterfall des ersten und damit logisch wahrscheinlicher. Den Grund, weshalb die meisten Teilnehmer dies anders einschätzten, benennt die Verfügbarkeitsheuristik. Die meisten (englischsprachigen) Menschen haben eine Regel dafür im Gedächtnis, wann Wörter auf „ing“ enden (Partizip Präsens und Gerundium). Eine vergleichbare Regel für Worte, deren zweitletzter Buchstabe „n“ ist, war ihnen nicht verfügbar. Die Verfügbarkeitsheuristik führt also allgemein formuliert dazu, dass Menschen aktuell verfügbare Information zulasten der *A-priori*-Wahrscheinlichkeit überbewerten. 508

Eine ganz ähnliche Verzerrung bewirkt die ebenfalls von Kahneman und Tversky dokumentierte **Repräsentationsheuristik** (*representativeness heuristic*).<sup>28</sup> Ihr liegt die Beobachtung zugrunde, dass Menschen oft Ereignisse in Kategorien einordnen. Bei der Bildung eines Wahrscheinlichkeitsurteils lassen sie sich entsprechend maßgeblich von der Ähnlichkeit zu einer bekannten Kategorie leiten, vernachlässigen aber die Wahrscheinlichkeit, mit der die Kategorie selbst auftritt. Ein Beispiel: Den Teilnehmern eines Experiments wurden Informationen über eine Frau mit dem Namen Linda präsentiert, die stereotypisch mit einer Feministin assoziiert werden konnten. Anschließend wurden sie gefragt, was wahrscheinlicher sei: (1) Linda ist eine Bankangestellte oder (2) Linda ist eine Bankangestellte, die sich in der Frauenrechtsbewegung engagiert. Die Mehrheit der Teilnehmer entschied sich für (2), obgleich diese Antwort wiederum nur ein logischer Unterfall von (1) ist und daher nicht wahrscheinlicher sein 509

<sup>27</sup> Grundlegend Kahneman/Tversky, Subjective probability: A judgment of representativeness, *Cognitive Psychology* 3 (1972), 430 ff.

<sup>28</sup> Vgl. Kahneman/Slovic/Tversky, *Judgement under Uncertainty: Heuristics and Biases*, 1982, S. 23 ff.

kann. Warum sollten sich Juristen mit Heuristiken beschäftigen? Hierauf gibt es viele Antworten. Zwei seien exemplarisch angeführt:

- 510 Die ökonomische Standardtheorie des Straf- und Ordnungswidrigkeitenrechts geht davon aus, dass Rechtsbrecher sich (neben der Strafhöhe auch) durch eine Erhöhung der Entdeckungswahrscheinlichkeit abschrecken lassen. Dies setzt natürlich voraus, dass sie sich ein realistisches Urteil darüber bilden, wie wahrscheinlich es ist, gefasst zu werden. Offenbar hängt ihr Wahrscheinlichkeitsurteil aber nicht nur von der tatsächlichen Wahrscheinlichkeit ab, sondern wird auch durch die Verfügbarkeitsheuristik beeinflusst. Kennen Kriminelle jemanden, der inhaftiert worden ist, oder ist ihnen ein solcher Fall aus den Medien bekannt, wird das ihr Urteil unverhältnismäßig stark beeinflussen. Eine sichtbare Art der Vollstreckung hat folglich Vorteile gegenüber einer unauffälligen. Die Polizei tut auch nicht schlecht daran, Erfolge in bestimmten Bereichen an die Medien zu vermelden. Nicht umsonst gehen nach der medienwirksamen Verhaftung eines einzelnen Prominenten wegen Steuerhinterziehung regelmäßig zahlreiche Selbstanzeigen ein.<sup>29</sup>
- 511 Auch Richter sind für Denken in Repräsentativität anfällig. Dies kann zu gefährlichen Fehlurteilen führen, wie jeder Strafverteidiger bestätigen kann. *Mark Schweizer* gibt hierfür ein anschauliches Beispiel aus dem Sexualstrafrecht.<sup>30</sup> Oftmals wird aus dem Vorliegen bestimmter, für einen Missbrauch typischer Symptome geschlossen, dass ein Kind tatsächlich missbraucht worden sei. Ein solcher Schluss ignoriert aber die *A-priori*-Wahrscheinlichkeit dieser Symptome. Allein aus der Tatsache, dass etwas typisch ist, folgt keinesfalls, dass es auch beweiskräftig ist. Vielmehr muss zusätzlich bekannt sein, wie oft die in Rede stehenden Symptome auch bei nicht betroffenen Kindern auftreten. Ist dies nicht bekannt, ist ihre Typizität kein *schlechter* Beweis, sondern überhaupt keiner. Experten im Bereich Kindesmissbrauch kennen sich naturgemäß mit missbrauchten Kindern besser aus als mit unversehrten. Ihr Urteil ist daher ergänzungsbedürftig um das Wissen unbeeinflusster Quellen.
- 512 **bb. Rückschaufehler.** Ähnlich wie die Verfügbarkeitsheuristik zur Überschätzung der Wahrscheinlichkeit aktuell präsenter Ereignisse führt, ver-

<sup>29</sup> Dass einem solchen Vorgehen normative Grenzen gesetzt sind, ist keine Frage. Besonders grausame Bestrafungen würden sicherlich zu einer hohen „Verfügbarkeit“ des Bestrafungsrisikos führen, wären aber ethisch nicht zu rechtfertigen. Unterhalb dieser Schwelle – im Bereich normaler Informationstätigkeit – mag aber durchaus Spielraum bestehen.

<sup>30</sup> *Schweizer*, Kognitive Täuschungen vor Gericht, 2005, S. 133 f.

leitet der sogenannte **Rückschaufehler** (*hindsight bias*) Menschen dazu, Ereignisse, die bereits stattgefunden haben, für wahrscheinlicher zu halten als alternative Verläufe („Das hätte man doch kommen sehen müssen.“). In einer Studie von *Baruch Fischhoff* etwa erhielten fünf Gruppen von Teilnehmern einen Text über den Konflikt zwischen der nepalesischen Ethnie der Gurkas und Briten im 19. Jahrhundert.<sup>31</sup> Die Texte glichen sich, was die Schilderung der Vorgeschichte anbetraf, waren aber mit vier unterschiedlichen Enden versehen. Der fünften Gruppe wurde das Ende vollständig vorenthalten. Nach der Lektüre wurden die Teilnehmer gefragt, welcher Gang der Ereignisse der wahrscheinlichste sei. Dabei zeigten sie sich systematisch zugunsten desjenigen Ausgangs eingenommen, der ihnen als der tatsächliche präsentiert worden war. Lediglich in der Kontrollgruppe konnte eine derartige Voreingenommenheit nicht festgestellt werden. Der Rückschaufehler hängt möglicherweise eng mit der Verfügbarkeitsheuristik zusammen. Ereignisse, die uns als wahr präsentiert werden, sind leichter vorstellbar und werden deshalb in ihrer Wahrscheinlichkeit überschätzt.

Juristen sind gut beraten, dieses Phänomen zur Kenntnis zu nehmen. Wann immer sie über die **Verletzung eines Sorgfaltsmaßstabs** zu entscheiden haben, tun sie dies *ex post*. Sie wissen also, dass es zum Schadensfall gekommen ist: Die Kapitalanlage ist verloren, der Patient tot oder eine Chemiefabrik explodiert. Dem Gericht stellt sich die Frage der Vorausschbarkeit des tragischen Ausgangs. Tatsächlich wird diese nur selten verneint. Diverse empirische Studien aus dem straf- und zivilrechtlichen Bereich belegen, dass Geschworene und Richter vielmehr regelmäßig dem Rückschaufehler erliegen.<sup>32</sup> Dies führt nicht nur zu ungerechtfertigten Verurteilungen, sondern schafft auch Anreize zu ineffizienter Überregulierung in Zukunft. Daher liegt es nahe, den Rückschaufehler möglichst weitgehend zu neutralisieren. Vieles weist darauf hin, dass der verstärkte Einsatz von Gruppendeliberation – wie in Richterkollegien oder unter Geschworenen – dazu einen Beitrag leisten kann. Auch die möglichst präzise Ausformulierung von Kodizes und Sorgfaltsmaßstäben *ex ante* erscheint hilfreich. Schließlich ist zu erwägen, ob in bestimmten Bereichen nicht eine Garantiehaftung zu gerechteren Ergebnissen führen könnte als die Verschuldenshaftung.

513

<sup>31</sup> *Fischhoff*, Hindsight is not equal to foresight: The effect of outcome knowledge on judgment under uncertainty, *Journal of Experimental Psychology* 1 (1975), 288 ff.

<sup>32</sup> *Guthrie/Rachlinski/Wistrich*, Inside the Judicial Mind, *Cornell Law Review* 86 (2001), 777 ff.; *Hastie/Schkade/Payne*, Juror judgement in civil cases: Hindsight effects on judgment of liability for punitive damages, *Law and Human Behavior* 23 (1999), 597 ff.

- 514 cc. **Überzogener Optimismus und selbstwertdienliche Verzerrung.** Das Bayes-Theorem setzt voraus, dass Menschen die Wahrscheinlichkeit des Eintritts eines Ereignisses von dem damit einhergehenden Nutzen trennen können. Tatsächlich gelingt dies aber den meisten Menschen nicht konsequent. Vielmehr erliegen sie verschiedenen Formen von Wunschdenken. Selbst wenn sie die statistische Wahrscheinlichkeit bestimmter Ereignisse abstrakt richtig einschätzen, beziehen sie diese Einschätzung nicht auf sich selbst. Ein klassisches Beispiel: Die meisten Menschen wären wohl in der Lage die Scheidungsrate in Deutschland auf etwa 30 Prozent zu beziffern. Fragt man aber Heiratswillige nach der Wahrscheinlichkeit, dass ihre eigene Ehe scheitern werde, so antworten sie erwartungsgemäß mit „null“.<sup>33</sup> Was im Beispiel trivial klingt, lässt sich auf diverse Gebiete ausweiten. In einer Studie von *Neil Weinstein* zeigten sich die meisten der befragten Studierenden überzeugt, dass ihnen in ihrem Leben positive Ereignisse häufiger als dem Durchschnitt ihrer Kommilitonen zuteilwerden und dass sie von Schicksalsschlägen eher verschont blieben.<sup>34</sup>
- 515 Übermäßiger Optimismus mag mit einer Selbstüberschätzung der eigenen Fähigkeiten zusammenhängen. Menschen glauben etwa, ihr Unfallrisiko sei unterdurchschnittlich, weil sie überdurchschnittlich gute Autofahrer seien. Ein spezielles Beispiel für diese **systematische Selbstüberschätzung** der eigenen Fähigkeiten liefert der sogenannte *overconfidence bias*, d.h. das übersteigerte Vertrauen in das eigene Urteilsvermögen oder Wissen. Dieser Effekt lässt sich leicht demonstrieren. Üblicherweise stellt man dazu Versuchspersonen eine schwierige Frage und bittet sie anzugeben, wie sicher sie sich ihrer Antwort sind. Selbst wenn die Kandidaten angeben, „100 %“ sicher zu sein, sind ihre Antworten regelmäßig nur zu 85–90 % richtig. Je schwieriger die Fragen, desto deutlicher tritt der Effekt außerdem zu Tage.<sup>35</sup>
- 516 Eng mit den zuvor genannten Phänomenen verwandt ist der sogenannte *self-serving bias* (**selbstwertdienliche Verzerrung**), der ebenfalls das reine *bayesianische* Urteil trüben kann.<sup>36</sup> Man versteht darunter die menschliche Tendenz, Erfolge auf die eigene Leistung zurückzuführen und Miss-

<sup>33</sup> Vgl. zur amerikanischen Situation *Baker/Emery*, When every relationship is above average: Perceptions and expectations of divorce at the time of marriage, *Law and Human Behavior* 17 (1993), 439 ff.

<sup>34</sup> *Weinstein*, Unrealistic optimism about future life events, *Journal of Personality and Social Psychology* 39 (1980), 806 ff.

<sup>35</sup> *Deffenbacher*, Eyewitness accuracy and confidence, *Law and Human Behavior* 4 (1980), 243 ff.

<sup>36</sup> Dazu *Miller/Ross*, Self-serving biases in the attribution of causality: Fact or fiction?, *Psychological Bulletin* 82 (1975), 213 ff.

erfolg externen Faktoren anzulasten. Im Alltag lassen sich viele Beispiele dafür beobachten: Studenten, die bei einem Examen gut abgeschnitten haben, sehen dies als Beleg ihrer Leistungsfähigkeit, während diejenigen, die nicht bestanden haben, dies der unfairen Aufgabenstellung zuschreiben. Auch der *self-serving bias* ermöglicht folglich dem Individuum zu glauben, dass es in für seine Selbstachtung wichtigen Bereichen überdurchschnittlich begabt sei.

Insgesamt ist die Terminologie im Hinblick auf die hier diskutierten Phänomene äußerst uneinheitlich. Dies ist allerdings nur das Symptom einer tiefergehenden, sachlichen Unsicherheit. In der Anamnese – d. h. in der Feststellung, dass übermäßiger Optimismus und Urteilsverzerrungen reale Phänomene darstellen – sind sich die meisten Wissenschaftler einig. Bislang erfolgt aber die Diagnose – die Verknüpfung der Beobachtung mit einer Ursache – noch recht wenig systematisch. Welche Ursachen den Beobachtungen zugrunde liegen, ist nicht immer eindeutig. Tatsächlich wird es häufig verschiedene Erklärungen für einen Effekt geben. Der *self-serving bias* etwa wird häufig motivational erklärt: Menschen haben ein Interesse daran, ihr Selbstbild zu erhalten, was ihre Wahrnehmung in diese Richtung hin verzerrt. Dies rechtfertigt die Einordnung des *self-serving bias* in die Kategorie Wunschdenken. Andere behaupten dagegen, der *bias* sei das Resultat strategischer Aufplusterung. Schließlich wird auch die Funktionsweise unseres Gedächtnisses zur Erklärung bemüht. Internale Gründe für Erfolge seien uns „verfügbarer“ als externale. Letztere Erklärung würde den Effekt eher in die Nähe der Verfügbarkeitsheuristik rücken.

Umstritten ist schließlich auch die gebotene Therapie. Wunschdenken und positives Denken haben nicht nur eine schlechte Seite. Dies belegt am eindrucksvollsten jene Gruppe von Individuen, die als einzige ihre Lebensrisiken korrekt einschätzt: die chronisch Depressiven. Dagegen zeichnen sich viele erfolgreiche Unternehmer durch ein besonders ausgeprägtes Vertrauen in die eigenen Fähigkeiten aus. Trotz aller Streitigkeiten im Detail dürfte klar sein, dass tiefgreifende Verzerrungen der menschlichen Risikowahrnehmung, wie die genannten, nicht einfach ignoriert werden können, sondern systematisiert und weiter ergründet werden müssen. Ähnlich klar ist, dass der an einem Verhaltensmodell interessierte Rechtswissenschaftler bereits heute nicht um die Beschäftigung mit ihnen umhin kommt. Einige Beispiele:

Dass übermäßiger Optimismus und Selbstüberschätzung die Abschreckung von Straftätern erschweren können, ist leicht ersichtlich. Wenn der Täter die Wahrscheinlichkeit einer Bestrafung zwar kennt, sie aber nicht auf sich bezieht, unterminiert dies die Logik der Standardtheorie. Auf eine

höhere Kriminalitätsrate muss es dennoch nicht zwingend hinauslaufen. Möglicherweise wird der sich selbst überschätzende Straftäter nämlich auch unvorsichtiger und dadurch leichter geschnappt.<sup>37</sup> Schwieriger als diese deskriptive Beobachtung gestaltet sich der normative Umgang mit übermäßigem Optimismus. Wann etwa schlägt er in Fahrlässigkeit um und ist bestrafungswürdig? Viele große Vermögen wurden unter Inkaufnahme „unvernünftiger“ Risiken erwirtschaftet. Viel (vielleicht noch mehr) Geld wurde bei solchen Manövern aber auch verloren. Wie das Recht etwa im Bankenbereich Sorgfaltsmaßstäbe angesichts dieser normativen Ambivalenz bestimmen sollte, ist zweifelhaft und erfordert eine Diskussion, die selbstverständlich auch die Problematik des *hindsight bias* einbezieht.

## b. Entscheidung

520 Eine wichtige Annahme der ökonomischen Standardtheorie und folglich der ökonomischen Analyse des Rechts ist, dass Menschen über ein **stabiles System von Präferenzen** verfügen, welches bestimmten Axiomen gehorcht.<sup>38</sup> Die Stabilitätsannahme ist weniger empirisch fundiert denn methodisch hilfreich. Soll eine positive Theorie falsifizierbar bleiben, muss ein Parameter als konstant behandelt werden. Ökonomen behaupten nicht, dass Präferenzen sich nie ändern können, nehmen aber an, dass sich solche Änderungen langsamer vollziehen als Änderungen der Rahmenbedingungen (Restriktionen), unter denen die Entscheidung stattfindet. Während dies häufig richtig sein mag, hat die Verhaltensökonomik in den letzten Jahrzehnten verschiedene Situationen identifiziert, in denen die Annahme systematisch verletzt wird. Präferenzen – so lautet das Fazit – sind nicht immer vorgegebene Konstanten, sondern werden durch Entscheidungskontext und -verfahren beeinflusst.

521 aa. **Ankereffekt.** Ein Beispiel hierfür ist der **Ankereffekt** (*anchoring*). Dieser wurde ursprünglich in einem Urteilkontext entdeckt. Dass der

<sup>37</sup> Auch dies ist nicht zwingend. Wenn potentielle Opfer ebenfalls überoptimistisch sind und annehmen, ihr Viktimisierungsrisiko sei unterdurchschnittlich, werden sie weniger effiziente Schutzmaßnahmen (Tresore, Wegfahrsperren etc.) installieren und dadurch dem Kriminellen sein Handwerk erleichtern. Viel spricht aber dafür, dass jedenfalls gegenwärtig eine solche Neigung in der Bevölkerung nicht existiert. Dafür mag die Verfügbarkeitsheuristik ausschlaggebend sein. Ausgelöst durch zahlreiche Medienberichte überschätzen die meisten Menschen etwa das Risiko ermordet zu werden gegenüber anderen Todesursachen (Rauchen, Diabetes etc.) erheblich.

<sup>38</sup> Hierzu Rz. 66.



*homo oeconomicus* die optimale Menge an Information in Betracht zieht, bedeutet im Umkehrschluss, dass er irrelevante Tatsachen außer Acht lässt. Wie *Daniel Kahneman* und *Amos Tversky* in einem vielzitierten Experiment nachgewiesen haben, lassen sich reale Menschen jedoch durchaus von kognitivem Störfeuer beeindrucken.<sup>39</sup> Sie baten die Teilnehmer zu schätzen, welcher Anteil der UNO-Mitgliedsstaaten in Afrika liegt. Vor der Beantwortung wurde ein Glücksrad gedreht, welches die Experimentatoren so manipuliert hatten, dass es immer entweder bei 10 oder bei 65 anhielt. Die Probanden wurden dann gebeten anzugeben, ob der Prozentsatz afrikanischer UNO-Mitglieder über oder unter diesem „zufällig“ bestimmten Wert liege. Erst danach wurden sie nach dem exakten Anteil gefragt. Ob das Glücksrad nun 10 oder 65 anzeigte, sollte selbstverständlich keinen Einfluss auf die spätere Beantwortung der Frage haben. Tatsächlich hatte es dies aber sehr wohl. Hielt das Glücksrad bei 10 an, schätzten die Teilnehmer den Prozentsatz afrikanischer UNO-Mitgliedsstaaten durchschnittlich auf 25 Prozent. Stoppte es hingegen bei 65, so lag die Durchschnittsschätzung bei 45 Prozent. Der zufällig ermittelte Wert wirkte also als Anker für die spätere Schätzung.

Dass der Ankereffekt nicht nur im Urteilsbereich auftritt, sondern in den Bereich der Entscheidung übergreift, ist mittlerweile hinlänglich bekannt. Nicht nur Häufigkeits- oder Wahrscheinlichkeitsurteile, sondern auch der Wert, der einem Gut beigemessen wird, scheinen dafür anfällig zu sein. Entsprechend wird der Begriff inzwischen für sämtliche unwillentliche Anpassungen eines numerischen Urteils an einen (willkürlich) vorgegebenen Vergleichswert gebraucht. Eine solche Anpassung widerspricht freilich der Vorstellung von stabilen Präferenzen. Kompatibel damit ist eher das Bild der *ad hoc* formierten, kontextabhängigen Präferenz.

Dass die Existenz des Ankereffekts auch Juristen beschäftigen sollte, zeigen diverse empirische Studien, von denen zwei exemplarisch skizziert seien: *Birte English* und *Thomas Mussweiler* baten deutsche Richterinnen mit einer Berufserfahrung von durchschnittlich 15 Jahren, aufgrund eines kurzen Sachverhalts ein Strafmaß für einen Vergewaltiger festzulegen.<sup>40</sup> Die verteilten Schilderungen unterschieden sich nur hinsichtlich der Forderung des „Staatsanwaltes“, welchen der Sachverhalt als Informatikstudenten zu erkennen gab. In einem Fall forderte dieser 34 Monate Freiheitsstrafe, im anderen nur zwölf Monate. *English* und *Mussweiler* beobachte-

<sup>39</sup> *Tversky/Kahneman*, Judgement under uncertainty: Heuristics and Biases, Science 185 (1974), 1124 ff.

<sup>40</sup> *English/Mussweiler*, Sentencing under Uncertainty: Anchoring Effects in the Courtroom, Journal of Applied Social Psychology 31 (2001), 1535 ff.



ten, dass sich die RichterInnen in voller Kenntnis seiner Rechtsunkundigkeit von dem Antrag des „Staatsanwaltes“ beeinflussen ließen. Forderte er 34 Monate Haft, verurteilten sie im Durchschnitt zu knapp 36 Monaten. Forderte er nur zwölf, so lautete das durchschnittliche Urteil auf 28 Monate. In verschiedenen anderen Untersuchungen wurden ähnliche Effekte beobachtet.

- 524 Im deliktsrechtlichen Bereich haben *Gretchen Chapman* und *Gary Bornstein* den Ankereffekt untersucht.<sup>41</sup> Mit „The more you ask for, the more you get“ bringt schon der Titel ihren Befund zum Ausdruck. Die zugesprochenen Schadenersatzsummen in ihrem Experiment wurden teilweise dramatisch von der Höhe der Anfangsforderung beeinflusst. Es liegt auf der Hand, dass diese Zufallsanfälligkeit nicht unserer Vorstellung eines fairen Verfahrens entspricht. Wie der Einfluss des Ankereffekts aber gemildert werden kann, ist weniger klar. In diversen Studien ist er als sehr robust nachgewiesen worden. Selbst Fachkenntnisse, Erfahrung oder die explizite Warnung vor seinem Einfluss vermögen ihn nicht ganz zu neutralisieren. Ganz hoffnungslos ist die Situation gleichwohl nicht: Es hat sich erwiesen, dass die Kenntnis ungefährer Orientierungswerte den Ankereffekt abschwächt. Das lässt hoffen, dass eine gefestigte Rechtspraxis in häufig verhandelten Konstellationen den Effekt eindämmt. Im Deliktsrecht sorgen hierfür bereits die zahlreichen (freilich unverbindlichen) Schmerzensgeldtabellen.
- 525 Teilweise mag sogar eine Einschränkung des richterlichen Entscheidungsspielraumes geboten sein, wie sie im amerikanischen Recht die sogenannten *Federal Sentencing Guidelines* bewirken. Diese Richtlinien sollten die Strafzumessung bundesweit harmonisieren. In der Rechtssache *United States v. Boker*<sup>42</sup> erklärte der US-amerikanische Supreme Court dieses Instrument in seiner ursprünglichen Form jedoch wegen Verstoßes gegen das Recht auf ein Jury-Verfahren für verfassungswidrig. Entsprechend gelten die *Guidelines* heute nicht mehr als bindend (dienen aber weiterhin als Orientierungswert). Vor dem Hintergrund des Ankereffektes kann man die Entscheidung des Obersten Gerichtshofs nicht begrüßen. Wo Richtwerte keine Orientierung geben, ist ein guter Anwalt dennoch nicht komplett hilflos. Er kann den Richter dazu bringen, sich ganz konkret Argumente vorzustellen, die gegen den „Anker“ sprechen. Dass der

<sup>41</sup> *Chapman/Bornstein*, The More You Ask for, the More You Get: Anchoring in Personal Injury Verdicts, *Applied Cognitive Psychology* 10 (1996), 519 ff.

<sup>42</sup> US Supreme Court, *United States v. Boker*, 543 U.S. 220.

als Erster Plädierende dennoch einen Vorteil behält, ist aber wohl kaum zu verhindern.

**bb. Aversion gegen Extreme.** *Rational Choice* postuliert, dass sich die Wahl, die ein Individuum zwischen zwei Optionen trifft, nicht dadurch ändern darf, dass eine dritte, nicht gewählte Option hinzutritt. Tatsächlich scheint dies in der Realität aber häufig der Fall zu sein. Dieser gut dokumentierte Effekt ist als **Aversion gegen Extreme** (*extremeness aversion*) bekannt. Verkäufer machen ihn sich gern zunutze. Nehmen wir an, ein Konsument wolle eine Hi-Fi-Anlage erwerben. Vor die Wahl gestellt zwischen einem preiswerten, aber weniger soundstarken Gerät für € 100 und einer hochwertigeren Anlage für € 200 tendiert er zu dem günstigen Angebot. Ein kluger Verkäufer zeigt ihm nun eine dritte, extrem hochwertige Anlage für € 800. Der Käufer wird diese wahrscheinlich nicht erwerben. Die Aversion gegen Extreme mag aber dazu führen, dass ihm die Anlage für € 200 plötzlich als guter Kompromiss erscheint. Menschen lassen sich häufig auf solche Kompromisse ein und bevorzugen instinktiv mittlere Alternativen. Dies mag häufig nicht unvernünftig sein, belegt aber wiederum die Kontextabhängigkeit unserer Präferenzordnung. 526

Welche Relevanz hat dieser Befund nun für die Rechtswissenschaft? 527 Eine Untersuchung von *Mark Kelman*, *Amos Tversky* und *Yuval Rottenstreich* weist nach, dass die Aversion gegen Extreme auch im juristischen Umfeld Auswirkungen hat.<sup>43</sup> Die Teilnehmer des Experiments sollten aufgrund desselben Sachverhaltes über die Verurteilung wegen eines Tötungsdelikts befinden. Die rechtlichen Erwägungen wurden jedoch so manipuliert, dass die eine Versuchsgruppe zwischen „qualifiziertem Mord“, „Mord“ und „vorsätzlicher Tötung“ entscheiden musste. Die andere Gruppe stand vor der Wahl zwischen „Mord“, „vorsätzlicher Tötung“ und „fahrlässiger Tötung“. In beiden Gruppen wurde am häufigsten die mittlere Option gewählt. Gruppe A verurteilte folglich wegen Mordes, während Gruppe B aufgrund desselben Sachverhaltes nur eine vorsätzliche Tötung annehmen wollte.

Vor dem Hintergrund der Aversion gegen Extreme muss sich ferner der Gesetzgeber bewusst sein, dass er mit der Einführung einer neuen Qualifikation oder eines Milderungstatbestandes auch die Präferenzen für den Grundtatbestand beeinflusst, wie *Mark Schweizer* in einer Studie mit 528

<sup>43</sup> *Kelman/Tversky/Rottenstreich*, Context-Dependence in Legal Decision Making, *Journal of Legal Studies* 96 (1996), 287 ff.

Richtern in der Schweiz nachweist.<sup>44</sup> Diese entschieden sich deutlich häufiger, statt einer Gefängnisstrafe die schwerere „ordentliche Verwahrung“ anzuordnen, wenn ihnen zugleich noch die dritte Option der „lebenslangen Verwahrung“ offenstand.

529 cc. **Prospect Theory.** Das *Coase-Theorem*<sup>45</sup>, welches als theoretische Grundsteinlegung der ökonomischen Analyse des Rechts gelten darf, enthält eine schlichte Kernaussage: In einer Welt mit niedrigen Transaktionskosten gelangt ein Gut unabhängig von der Anfangsverteilung durch die Kräfte des Marktes in die Hände desjenigen, der ihm den höchsten Wert beimisst. Diese Person wird das Gut nämlich von dem Eigentümer erwerben, der dafür eine Gegenleistung erhält. Dieses Argument beruht freilich auf der Prämisse, dass Individuen eine feststehende Präferenzordnung hinsichtlich der einzelnen Güter haben. Andernfalls wären interpersonale Vergleiche nicht möglich. Verfahren, Kontext und Reihenfolge der Entscheidung sollten keine Rolle spielen. Beispielhaft: Wer vor dem Kauf einen VW Golf einem Honda Civic vorzieht, müsste auch bereit sein, einen bereits in seinem Besitz befindlichen Honda gegen den Volkswagen einzutauschen. Die Entdeckung des sogenannten **Besitzeffektes** (*endowment effect*) hat diese Prämisse jedoch in Zweifel gezogen.

530 Häufig zitiert wird in diesem Zusammenhang ein Experiment, welches mit Studenten der *Cornell University* durchgeführt wurde.<sup>46</sup> Darin wurde ein Markt simuliert, auf dem *Cornell*-Kaffeetassen gehandelt wurden. Die eine Hälfte der teilnehmenden Studenten erhielt eine solche Tasse, während die andere Hälfte mit jeweils \$ 6 Kaufkraft ausgestattet wurde. Diejenigen Studenten, die im Besitz einer Tasse waren, wurden gebeten zu spezifizieren, für welchen Betrag sie diese verkaufen würden. Die Mitglieder der anderen Gruppe sollten den für sie maximal akzeptablen Kaufpreis beziffern. Anschließend berechneten die Experimentatoren den Markträumungspreis und vollzogen die zu diesem Preis möglichen Transaktionen. Das *Coase-Theorem* legt die Voraussage nahe, dass ungefähr die Hälfte aller Tassen den Eigentümer wechseln sollte. Die Transaktionskosten lagen nahe null und die Ursprungsverteilung war zufällig. Die Realität sah jedoch anders aus. Tatsächlich konnten nur wenige Transaktionen vollzogen werden. Die Minimalverkaufspreise der Tassen-Besitzer lagen im Durchschnitt etwa doppelt so hoch wie die maximale Zahlungsbereit-

<sup>44</sup> Schweizer, Kognitive Täuschungen vor Gericht, 2005, S. 256 ff.

<sup>45</sup> Hierzu Rz. 150 ff. sowie Rz. 250 ff.

<sup>46</sup> Kahneman/Knetsch/Thaler, Experimental Tests of the Endowment Effect and the Coase Theorem, *Journal of Political Economy* 98 (1990), 1325 ff.

schaft der potentiellen Käufer. Die Experimentatoren folgerten, dass Menschen eine Sache, die bereits in ihrem Besitz ist, mehr schätzen als dies zuvor der Fall war. Zahlreiche weitere Experimente haben die Existenz des Besitzeffektes erhärtet und andere Ursachen für sein Auftreten – etwa strategisches Handeln oder Reichtumseffekte – ausgeschlossen. Vielmehr scheint der Besitz allein die höhere Wertigkeit einer Sache für Menschen zu begründen. Damit ist auch die Feststellung kompatibel, wonach der Besitzeffekt umso stärker auftritt, je länger ein Mensch eine Sache schon besessen hat.

Der Besitzeffekt ist wohl lediglich Teil eines größeren Mosaiks. In ihm, 531  
so lässt sich aufgrund anderer Experimente vermuten, äußert sich die generelle menschliche Neigung, Verluste stärker zu gewichten als Gewinne (**Verlustaversion**). Vor die Wahl gestellt zwischen einem sicheren Gewinn von € 240 und einer Lotterie, die mit 25 % Wahrscheinlichkeit € 1000 und mit 75 % nichts einbringt, entscheidet sich eine sehr große Mehrheit der Menschen für die sichere Option. Anders wenn es um Verluste geht: Ein Spiel, das mit 75 % Wahrscheinlichkeit zum Verlust von € 1000 führt, mit 25 % aber einen Verlust komplett vermeidet, wird einem sicheren Verlust von € 750 regelmäßig vorgezogen. Geht es um Gewinne, verhalten sich die meisten Menschen also risikoscheu und bevorzugen Sicherheit. Geht es dagegen um die Vermeidung von Verlusten, sind sie bereit, mehr aufs Spiel zu setzen. Dieses Ergebnis lässt sich nicht mit der traditionellen Erwartungsnutzentheorie in Einklang bringen. Diese lässt zwar individuell unterschiedliche Risikoneigungen zu. Dass sich die Risikoeinstellung eines Individuums aber daran orientiert, ob über Verluste oder Gewinne entschieden wird, ist mit ihr unvereinbar. Gleichwohl wurde in zahlreichen Experimenten diese Beobachtung dokumentiert. Menschen scheinen Verluste im Durchschnitt etwa doppelt so stark zu gewichten wie Gewinne in gleicher Höhe. Dies ist nicht nur im Labor nachgewiesen worden, sondern hilft auch, bestimmte im Feld beobachtete Phänomene zu erklären.

Auch die Verlustaversion selbst scheint aber auch nur ein etwas größerer 532  
Ausschnitt des Mosaiks zu sein. Zusammen mit dem Besitzeffekt lässt sie sich als Facette eines Effektes begreifen, der gelegentlich als *status quo bias* bezeichnet wird. Damit ist eine starke Präferenz für den Ist-Zustand gemeint. Menschen ändern diesen nur, wenn sie mit starken Anreizen konfrontiert werden. Dass Menschen sich in ihrem Entscheidungsverhalten weniger an absoluten Vermögensständen orientieren, wie es die Erwartungsnutzentheorie annimmt, sondern Veränderungen ausgehend von einem Referenzpunkt – häufig dem Status quo – bewerten, gehört zu den Grundannahmen des bedeutendsten positiven Alternativmodells zur Er-

wartungsnutzentheorie: Der von *Daniel Kahneman* und *Amos Tversky* entwickelten **Prospect Theory**,<sup>47</sup> die die zitierten empirischen Befunde integriert.

533 Die *Prospect Theory* modelliert Entscheidungsverhalten mittels zweier zentraler Komponenten: einer Wertfunktion und einer Funktion, welche die objektiven Wahrscheinlichkeiten gewichtet. Fasst man die Grundaussagen der Theorie zusammen, ergibt sich folgendes Bild:

(1) Die **Wertfunktion** ist S-förmig und weist einen Knick am Referenzpunkt auf. Von diesem aus verläuft sie im Gewinnbereich konkav, im Verlustbereich hingegen konvex. Dies bedeutet, dass bei Entscheidungen zwischen Optionen, welche relativ zum Referenzpunkt als Verluste erscheinen, die meisten Menschen risikofreudig agieren. Erscheinen die Optionen hingegen als Gewinne gegenüber dem Referenzpunkt, verhalten sie sich risikoavers.

(2) Diese Risikopräferenzen kehren sich jedoch infolge des Einflusses der **Gewichtungsfunktion** in denjenigen Fällen um, in denen es um Verluste oder Gewinne von geringer Wahrscheinlichkeit geht. Hier entscheiden die meisten risikofreudig, wenn es um Gewinne geht, aber risikoavers im Hinblick auf Verlustoptionen. Die Gewichtungsfunktion verläuft an den Enden nämlich sehr steil, was bedeutet, dass kleinen Wahrscheinlichkeiten unverhältnismäßig viel Gewicht eingeräumt wird.

(3) Bei Wahrscheinlichkeiten von 0,3–0,4 stimmen subjektives Empfinden und tatsächliche Wahrscheinlichkeit am besten überein. Kleinere Wahrscheinlichkeiten werden über-, größere tendenziell unterschätzt.

534 Entscheidend für die praktische Nützlichkeit der *prospect theory* ist natürlich, ob sich der Referenzpunkt bestimmen lässt. Dies wird nicht immer ex ante möglich sein. Im Regelfall dürfte sich der Referenzpunkt jedoch aus der Natur der Sache ergeben. Üblicherweise wird er mit dem Ist-Zustand übereinstimmen, wie der *status quo bias* nahelegt. In manchen Fällen kommt aber auch ein bestimmter (und bestimmbarer) Soll-Zustand in Betracht, etwa wenn es um Ertragsziele oder die Erfüllung eines Tagessolls geht.

535 **dd. Framing.** Die Tatsache, dass es einen „objektiven“ Referenzpunkt nicht gibt, schafft die Grundlage für Manipulation. Eine gezielte Beeinflussung des Referenzpunktes wird als **framing** bezeichnet. Die bekanntes-

<sup>47</sup> Grundlegend *Kahneman/Tversky*, Prospect Theory: An Analysis of Decision under the Risk, *Econometrica* 47 (1979), 263 ff.; zu Weiterentwicklungen *dies.*, Advances in Prospect Theory: Cumulative Representation of Uncertainty, *Journal of Risk and Uncertainty* 5 (1992), 297 ff.

te Illustration dieses Problems ist das sogenannte *Asian-Disease-Szenario*.<sup>48</sup> *Kahneman* und *Tversky* baten zwei Gruppen von Testpersonen sich vorzustellen, sie müssten angesichts einer Seuche, die das Leben von 600 Menschen bedroht, zwischen zwei möglichen Rettungsplänen wählen. Der ersten Gruppe wurden die Wahlmöglichkeiten folgendermaßen dargestellt: Plan A rettet mit Sicherheit 200 Menschen. Plan B dagegen mit einer Wahrscheinlichkeit von 1/3 alle 600 bedrohten Bürger, mit der Restwahrscheinlichkeit aber niemanden. Der zweiten Gruppe wurde das Dilemma präsentiert als Entscheidung zwischen Plan C, der den sicheren Tod von 400 Menschen zur Folge hat, und Plan D, bei dem eine 1/3-Chance besteht, dass niemand stirbt, während mit 2/3 Wahrscheinlichkeit alle 600 Menschen ihr Leben verlieren. Plan A und Plan C sowie Plan B und D entsprechen sich inhaltlich. Verschieden ist nur die Art ihrer Darstellung. Plan A und B haben einen positiven *frame*, stellen also die Wahl als eine zwischen Gewinnoptionen (gerettete Menschen) dar, während Plan C und D einen negativen *frame* aufweisen, bei dem zwischen Verlusten zu entscheiden ist. Dieser Unterschied hat gravierende Konsequenzen. Die meisten Menschen bevorzugten Plan A gegenüber B, aber D gegenüber C. Sie ziehen also die riskante Version im negativen Frame vor, jedoch die sichere bei der Wahl zwischen Gewinnen. Dies entspricht exakt den Prognosen der *prospect theory*.

Die Beobachtung, dass sich die Risikoneigung und damit die Entscheidung von Menschen durch die Darstellung des Entscheidungsproblems beeinflussen lassen, hat auch Bedeutung für das Recht. Ein anschauliches Beispiel liefert das Problem der **Steuerhinterziehung**. Die Entscheidung, ob man Steuern hinterziehen will, entspricht der Wahl zwischen einer sicheren Option und einer riskanten Option, deren Ausgang besser (Steuerersparnis) oder schlechter (Nachzahlung, Bestrafung, etc.) sein kann als bei der sicheren. Wie die Entscheidung ausfällt, hängt letztlich von der individuellen Risikoneigung ab. Diese wiederum aber ist beeinflussbar – etwa durch das Verfahren der Steuererhebung. Direkte Steuern wie die deutsche oder US-amerikanische Einkommenssteuer werden an der Quelle erhoben. Dadurch werden Manipulationen eher über das Instrument der Rückzahlung möglich. Dagegen wird – etwa in der Schweiz – die Steuer aus Mitteln bezahlt, die sich bereits auf dem Konto des Arbeitnehmers befinden. Aus Sicht des Standardansatzes sollte dieser Unterschied die Häufigkeit von Steuerdelikten nicht beeinflussen. Die *prospect theory* dagegen erlaubt

536

<sup>48</sup> Dazu *Tversky/Kahneman*, Judgement under Uncertainty: Heuristics and Biases, Science 185 (1974), 1124 ff.

eine abweichende Voraussage: Die Steuerrückzahlung dürfte regelmäßig als Gewinn wahrgenommen werden, da sie das Ist-Vermögen steigert. Die beim Steuerpflichtigen erhobene Steuer dagegen wird regelmäßig als Verlust wahrgenommen, weil sie den auf dem Konto vorhandenen Betrag verringert. Da Menschen sich im Gewinnbereich eher risikoavers verhalten, also die sichere Alternative gegenüber der riskanten bevorzugen, sollten in Ländern mit direkt an der Quelle erhobenen Steuern Steuerhinterziehungen seltener sein als etwa in der Schweiz. Diese Voraussage wird in der Realität bestätigt. Vorauszahlungen auf die Steuerschuld zu verlangen, scheint für den Staat also durchaus einen Vorteil zu haben.

537 Dass Menschen dem Status quo nicht-rational hohes Gewicht einräumen und sich in ihrem Entscheidungsverhalten durch die Darstellung des Entscheidungsproblems beeinflussen lassen, wirft ferner ein anderes Licht auf die Debatte um die Zulässigkeit **paternalistischer Regulierung**. Für Rationaltheoretiker steht fest, dass Menschen selbst am besten wissen, was sie glücklich macht. Aber selbst wenn dies ausnahmsweise nicht zutrifft, so argumentieren die meisten Paternalismuskritiker, darf man Menschen nicht zu ihrem Glück zwingen. Was aber wenn das Glück des Menschen dergestalt kontingent ist, dass die Art und Weise, wie man danach fragt, zu unterschiedlichen Antworten führen kann? Ein Beispiel mag das illustrieren:<sup>49</sup> In den US-Staaten Pennsylvania und New Jersey wurde ein identisches Basisversicherungspaket angeboten. Die Versicherungsnehmer hatten aber die Wahl zwischen einer teureren Variante mit Klagerecht und einer billigeren ohne dieses Recht. In New Jersey war die teurere Variante der Standard. Versicherungsnehmer mussten also eine bewusste Entscheidung treffen, um zu der billigeren Variante zu wechseln. In Pennsylvania verhielt es sich exakt umgekehrt. Aus rationaltheoretischer Sicht war zu erwarten, dass sich nach einiger Zeit in beiden Staaten ein ähnlicher Anteil von Versicherungsnehmern für das teure und das billige Versicherungspaket entscheiden würde. Tatsächlich aber behielt die große Mehrheit in beiden Staaten die jeweilige Standardoption bei. Sachliche Gründe für diesen Unterschied ließen sich nicht ausmachen.

538 Offensichtlich war es tatsächlich die bloße Tatsache, dass eine Option als „Standard“ (*default*) angeboten wurde, die sie als vorzugswürdig erscheinen ließ. Daraus haben namentlich *Cass Sunstein* und *Richard Thaler* gefolgert, dass Paternalismus unausweichlich sei.<sup>50</sup> Der Staat manipulierte

<sup>49</sup> Nach *Sunstein*, *Behavioral Law and Economics: A Progress Report*, *American Law & Economic Review*, 1 (1999), 115 (124).

<sup>50</sup> Vgl. *Sunstein/Thaler*, *Libertarian paternalism is not an oxymoron*, *University of Chicago Law Review* 70 (2003), 1159 ff.



schon mit seiner Entscheidung für einen bestimmten Standard das Entscheidungsverhalten der Bevölkerung, selbst wenn er ihnen die freie Wahl lasse, vom Standard abzuweichen. Diesem Faktum müsse der Staat Rechnung tragen und für den Bürger mitdenken. *Sunstein* und *Thaler* nennen ihre Idee „**libertären Paternalismus**“; wir werden uns das aus diesen Ideen abgeleitete Konzept des „**Nudging**“ gleich noch etwas genauer anschauen. Wollte man ihrem Argument folgen, hätte dies weitreichende Konsequenzen für diverse Rechtsgebiete von Verbraucherschutz bis zum Gesundheitsrecht.

### 3. Begrenzte Selbstdisziplin

Das oben bereits geschilderte „*Asian-Disease-Szenario*“ zeigt, dass die Vorstellung eines stabilen, kontextunabhängigen und geordneten Präferenzsystems nicht immer der Realität entspricht. Vielmehr kann es bisweilen zu einer Präferenzumkehr kommen. Im Falle des *Asian-Disease-Szenarios* geschieht dies aufgrund der unterschiedlichen Darstellungen (*frames*) des Entscheidungsproblems. Die Tatsache, dass Menschen ungleiche Risikopräferenzen für die Entscheidung zwischen Gewinnen und Verlusten haben, führte in diesem Beispiel zu einer normativ widersprüchlichen Wahl. Ein anderer Fall, in dem unsere Präferenzen zueinander in Widerspruch treten können, ist jedem Menschen aus der Alltagserfahrung bekannt. Die Rede ist von einer Präferenzumkehr im Zeitablauf. Wir nehmen uns heute vor, nie mehr zu rauchen, werden jedoch nächste Woche schon wieder rückfällig. Die meisten Raucher geben an, sie wollten „eigentlich“ aufhören. Was sich hinter dem Wort „eigentlich“ verbirgt, ist eine Kollision zwischen **langfristigen Präferenzen** (gesunder Lebenswandel) und **kurzfristigen Präferenzen** (Suchtbefriedigung). Oftmals obsiegen in diesem Konflikt letztere – „*the heat of the moment*“, so eine englische Redewendung, kann uns überwältigen.

Der *homo oeconomicus* in seiner reinsten Form kennt solche Probleme nicht. Das bedeutet nicht, dass nicht auch er sofortige Freuden gegenüber zukünftigen bevorzugen würde. Dies folgt aber einem streng ökonomischen Kalkül. Zukünftiger Nutzen wird diskontiert, denn die Zukunft ist unsicher und man weiß nicht, ob man sie erleben wird. € 100 heute oder in zehn Jahren sind ein großer Unterschied. Geld, das ich heute erhalte, kann bereits Zinsen abwerfen. In zehn Jahren mag die Inflation seinen Wert aufgefressen haben. Der *homo oeconomicus* mag eine hohe oder eine niedrige **Diskontierungsrate** – also eine starke Präferenz für die Zukunft oder die Gegenwart – haben, ohne dass dies inkompatibel mit der Ratio-



naltheorie wäre. Nur dürfen seine Präferenzen nicht widersprüchlich sein. Dies liegt an der speziellen Form der Diskontierungsfunktion, welche die Ökonomen ihm mitgegeben haben. Der *homo oeconomicus* diskontiert die Zukunft nämlich exponentiell. Der konstante Exponent aber schließt im Zeitablauf konfligierende Präferenzen aus. Wer sich also für die Schokolade am Abend entscheidet, darf dies am nächsten Morgen auf der Waage nicht bereuen.

541 Dass die Realität reichlich Belege für das Gegenteil bietet, ist auch Ökonomen nicht verborgen geblieben. Insbesondere die zahlreichen Beispiele vorausschauender Selbstbindung haben ihr Interesse geweckt. Viele Menschen, die abnehmen wollen, kaufen beispielsweise erst gar keine Schokolade, um nicht zu viel davon zu essen. Casinos bieten die Möglichkeit an, sich selbst auf Lebenszeit Hausverbot zu erteilen. Derartige Verhaltensweisen kann *Rational Choice* nicht sinnvoll deuten. Mit der Annahme exponentieller Diskontierung ist die Vorstellung, dass wir unseren Entscheidungsspielraum freiwillig einschränken, um uns vor unseren eigenen Präferenzen zu schützen, nicht kompatibel. Da sich aber nicht leugnen lässt, dass Menschen sich häufig wie Odysseus an den Mast binden lassen, hat die Ökonomik neue Wege der Modellierung ersonnen.<sup>51</sup> In diesen innovativen Modellen wird der Mensch nicht mehr mittels einer Präferenzordnung beschrieben, sondern erhält plötzlich mehrere Präferenzsysteme. Typischerweise stehen sich dabei ein kurzfristig und ein langfristig orientiertes Ich gegenüber.

542 *Behavioral Economics* führt den Grundgedanken dieser sogenannten *Multiple-selves*-Modelle konsequent weiter. Statt mehrerer „Ichs“ mit exponentieller Diskontierung wählen Verhaltensökonomien aber eine andere Modellierungstechnik. Diese ist als (quasi-)hyperbolische Abzinsung bekannt geworden.<sup>52</sup> Die Essenz des Ansatzes besteht in der Idee eines „Gegenwartsbias“ (*present bias*), der einer extremen Überbewertung sofortigen Konsums entspricht. Dieser Effekt kann zu Entscheidungen führen, die später bedauert werden, was wiederum deutlich macht, warum Selbstbindung manchmal sehr sinnvoll sein kann. Erfolgreiche Selbstbindung setzt freilich voraus, dass ein Individuum die Stärke seiner zukünftigen Bedürfnisse richtig einzuschätzen weiß. Ob dies der Fall ist, wird in der Verhaltensökonomik unterschiedlich bewertet; tatsächlich dürfte die Wahrheit situationsabhängig sein. Eine sorgfältige Rezeption psycholo-

<sup>51</sup> Das Bild stammt von *Elster*, Ulysses and the Sirens, *Studies in Rationality and Irrationality*, 1984.

<sup>52</sup> Grundlegend *Laibson*, Golden Eggs and Hyperbolic Discounting, *Quarterly Journal of Economics* 112 (1997), 434 ff.

gischer Erkenntnisse zum Thema Selbstkontrolle durch die Ökonomik ist geboten, um hierzu genauere Antworten geben zu können.

Für Juristen ist das Problem der Selbstkontrollmängel von großem Interesse, wie einige Beispiele illustrieren sollen: Teilweise wird damit der Idee des „Selbstpaternalismus“ eine normative Basis gegeben. Wenn Menschen Kollisionen zwischen gegenwärtigen und zukünftige Präferenzen voraussehen können, hilft dies zu begründen, warum Selbstausschlüsse aus Casinos und Ähnlichem vollstreckt und die Nichtdurchsetzung den Betreiber zum Schadensersatz verpflichten sollte. Auch das Betäubungsmittelrecht sollte nach Vorstellung einiger in dieser Weise reguliert werden.<sup>53</sup> Der Staat solle ein Konzessionssystem entwickeln, das Volljährigen nach Aufklärung über damit verbundene Risiken freien Zugang zu heute noch illegalen Drogen gewähre. Allerdings müsse dieses System auch die Möglichkeit der Festlegung einer persönlichen Höchstgrenze oder eines kompletten Selbstausschlusses vorsehen. Der Vorschlag ist nicht per se absurd. Angesichts des offenkundigen Scheiterns einer auf das Strafrecht bauenden Drogenpolitik ist die Suche nach alternativen Regulierungsmöglichkeiten naheliegend. Tatsächlich wird nämlich der Staat auch durch beliebig harte Strafen das Kosten-Nutzen-Kalkül vieler Abhängiger nicht mehr effektiv beeinflussen. Deren Langzeitpräferenzen stehen dem Drogenkonsum wahrscheinlich ohnehin entgegen. Indessen ist es das kurzfristig orientierte Ich, welches die strafrechtlich relevanten Entscheidungen zum Rückfall in die Sucht trifft. Wie allerdings verhindert werden soll, dass derjenige, der sich selbst den Zugang zu den legalen Drogenkanälen versperrt hat, wieder den Weg auf den Schwarzmarkt findet, wenn die Sucht ihn überkommt, beantwortet der genannte Vorschlag nur unzureichend.

Indem sie Selbstkontrollmängel ernst nimmt, trifft sich die Verhaltensökonomik ferner mit der modernen **kriminologischen Theorie**. Manche Kriminologen erachten fehlende Selbstkontrolle sogar als maßgebliche Determinante devianten Verhaltens.<sup>54</sup> Wenn Menschen zudem naiv ihren zukünftigen Bedürfnissen gegenüber sind, wird eine strafrechtliche Drohung oft wenig bewirken können. Möchte der Staat bestimmte Verhaltensweisen dennoch unterbinden, muss er stattdessen auf *Ex-ante*-Regulierung setzen. Er sollte, mit anderen Worten, nicht androhen, Menschen dafür zu bestrafen, dass sie die Kontrolle über sich verlieren. Wenn dies nämlich geschieht, sind sie für seine Strafdrohung bereits unerreichbar („Ich kannte mich selbst nicht.“). Er muss sie vielmehr ansprechen, bevor es zum

<sup>53</sup> Leitzel, Self Exclusion, SSRN Working Paper, 2008.

<sup>54</sup> Grundlegend Gottfredson/Hirschi, A General Theory of Crime, 1990.

Kontrollverlust kommt, und rechtliche Zäune um Situationen errichten, in denen Kontrollverluste häufig sind. Rechtsinstitute wie die strafrechtliche *actio libera in causa* richten an den Rechtsunterworfenen die normative Verhaltenserwartung, dass er sich selbst darum kümmern werde, nicht in einen solchen Zustand zu geraten. Dies setzt jedoch voraus, dass er seine zukünftige Neigung zur Straffälligkeit richtig einschätzen kann. Viel weist darauf hin, dass dies eine unrealistische Erwartung ist. Erfolgreicher wird der Staat sein, wenn die Sanktion bereits die Herbeiführung des Zustandes reguliert – etwa indem er etwa Gastwirte schärfer dazu anhält, erkennbar Volltrunkenen keinen Alkohol mehr auszuschenken.

#### IV. Nudging: Verhaltenswissenschaftliche Rezepturen für staatliche Steuerung?

- 545 Die Einsichten der Verhaltensökonomik sind aber nicht nur, wie in diesem Kapitel bisher gezeigt, punktuell für Juristen interessant, und die Systematisierung der Einsichten ist nicht nur aus einer wissenschaftlichen Perspektive vielversprechend: Auch aus rechtspolitischer Sicht verheißt eine systematische Analyse und Nutzung verhaltenswissenschaftlicher Einsichten effektivere Regulierung. Mit anderen Worten: Dass menschliches Verhalten einer begrenzten Rationalität und einer begrenzten Selbstdisziplin unterliegt und dass diese Verhaltensartefakte selbst systematisch auftreten und damit vorhersehbar sind, kann auch für die **Verhaltenssteuerung durch Recht** nutzbar gemacht werden.
- 546 Das Verhalten der Bürgerinnen und Bürger kann danach nicht nur durch Ge- und Verbote, durch Anreize und Sanktionen beeinflusst werden (wie uns das klassische ökonomische Paradigma lehrt), vielmehr können Menschen auch auf andere – vermeintlich „sanftere“ – Art und Weise zu bestimmten Verhaltensweisen veranlasst werden. Ferner können bestimmte Steuerungsmaßnahmen sinnvoll sein, wenn „vernünftige“ Entscheidungen nur aus Mangel an Weitsicht oder Selbstdisziplin nicht getroffen werden. Und schließlich lässt sich aus der Wohlfahrtsökonomik ableiten, dass Verhalten in solchen Fällen gesteuert werden sollte, in denen die damit erreichten Verhaltensweisen insgesamt gemeinwohlförderlich sind. Es geht also darum, Erkenntnisse der Verhaltenswissenschaften zu nutzen, um Verhalten staatlicherseits intelligent und wirkungsvoll zu steuern.
- 547 Diesem Ansatz haben sich in jüngerer Zeit verschiedene Regierungen verschrieben. Sein prominentester Vertreter, *Cass Sunstein*, hat den US-Präsidenten *Barack Obama* von 2009–2012 als Leiter des *Office of*

*Information and Regulatory Affairs* beraten. Die britische Regierung hat 2010 ein *Behavioral Insight Team* (inoffiziell bekannt als *Nudge Unit*) eingesetzt, um verhaltensökonomische Einsichten bei der Rechtsgestaltung einzusetzen. 2015 hat die auf den Koalitionsvertrag zurückgehende Arbeitsgruppe „Wirksam regieren“ im deutschen Bundeskanzleramt ihre Arbeit aufgenommen; sie hat den Auftrag, sich mit den Möglichkeiten verhaltenswissenschaftlich informierter Regulierung zu befassen.

## 1. Konzept

Hinter dem von *Richard Thaler* und *Cass Sunstein* geprägten Begriff *Nudging* (zu Deutsch etwa „schubsen“) steckt die Idee, dass menschliches Verhalten durch Mechanismen beeinflusst werden kann, die uns aus dem Marketing schon länger bekannt sind: Im Supermarkt wird etwa die teure Schokolade auf Augenhöhe präsentiert, um es uns als Käufern möglichst leicht zu machen, nach ihr zu greifen, während der Griff nach unten, zu günstigeren Produkten, etwas mehr Aufwand bedeutet. Die Marktbetreiber schubsen uns also durch die Art und Weise, wie sie das Regal befüllen und den gesamten Markt einrichten, dorthin, wo sie uns haben wollen (nämlich zu den teuren Produkten), und als Marketing-Maßnahme erscheinen uns diese „Psycho-Tricks“ vielleicht lästig, aber letztlich doch legitim. Ausgehend von solchen Beobachtungen definieren *Sunstein* und *Thaler* einen **Nudge** („Schubser“) als jede Maßnahme, die das Verhalten von Menschen auf vorhersehbare Weise verändert, ohne Handlungsalternativen zu verbieten oder die ökonomischen Anreize signifikant zu verändern. Um als *Nudge* gewertet werden zu können, muss der Betroffene der jeweiligen Intervention einfach ausweichen können: Das Obst auf Augenhöhe zu platzieren gilt demnach als *Nudge*, während ein Verbot von Fertiggerichten kein solcher ist.<sup>55</sup> 548

Diese Entscheidungsfreiheit, die bei den Rechtsunterworfenen verbleibt, ist der markante Unterschied zwischen *Nudging*-Maßnahmen und verbindlichen rechtlichen Regelungen: Vorab wird staatlicherseits entschieden, welches Verhalten „vernünftiger“ oder gemeinwohldienlicher ist, und die rechtliche Regulierung – die **Entscheidungsarchitektur** – sodann so ausgestaltet, dass es den Bürgern leichter gemacht wird, sich für genau diese Handlungsalternative zu entscheiden; wenn sie aber dennoch den anderen Weg gehen wollen, wird ihnen das weder verwehrt noch ökonomisch unattraktiv gemacht. Dadurch, so das Argument, greife der Staat 549

<sup>55</sup> *Thaler/Sunstein*, *Nudge: Wie man kluge Entscheidungen anstößt*, 2010, S. 13.

viel „sanfter“ in die Rechte der Betroffenen ein als durch klassische, sanktionierte Ge- und Verbote. *Nudging* eröffne den Rechtsunterworfenen die Möglichkeit, ihr Verhalten auch dort zu optimieren, wo ihre Willensschwäche ihnen im Weg stehe. Schließlich seien solche Entscheidungsarchitekturen letztlich unentrinnbar: Irgendwelche Güter müssten schließlich auf Augenhöhe platziert werden, und für irgendwelche *default*-Regeln<sup>56</sup> müsse sich das Recht entscheiden – warum also nicht jene nehmen, die dem wohlverstandenen, langfristigen Eigeninteresse und vielleicht sogar der gemeinschaftlichen Wohlfahrt am besten dienen?

## 2. Instrumente

- 550 Ein klassisches Beispiel für *Nudging* durch den Gesetzgeber liefert uns die Debatte um die Regelung zur Organspende: In Deutschland gilt eine sogenannte **Opt-in-Regel**: Wenn Sie Organspender sein möchten, müssen Sie dies ausdrücklich erklären, indem Sie einen Organspendeausweis ausfüllen. In vielen anderen Ländern, etwa in Frankreich, Italien und Österreich, gilt dagegen eine **Opt-out-Regel**: Hier sind Sie zunächst automatisch Spender, es sei denn Sie geben eine Erklärung ab, die dem ausdrücklich widerspricht. Die Varianten unterscheiden sich in der Frage, was die **Default-Regelung** ist. In Deutschland wird den Bürgern also Aufwand zugemutet, wenn sie Spender sein wollen, während ihnen das Nichtspenden leicht(er) gemacht wird. Wie im Supermarktbeispiel entscheiden sie jedoch in beiden Fällen frei, ob sie ihre Organe spenden oder nicht. Aus der Sicht der Gemeinschaft hat das dramatische Konsequenzen, denn in Ländern mit *Opt-out*-Regeln ist der Anteil der Organspender deutlich höher. Ebenfalls als *Nudge* gelten Ampeln auf Lebensmitteln, die darauf hinweisen, ob die entsprechenden Produkte unserem Körper mehr oder weniger guttun, oder Auszeichnungen zur Energieeffizienzklasse; sie sind zwar für die Hersteller zwingend, aber die Konsumenten können die Information, die ihre Entscheidung leiten soll, auch einfach ignorieren.
- 551 *Sunstein* nennt zehn Arten solcher *Nudges*, die aus der Steuerungsperspektive des Rechts wichtig seien:<sup>57</sup>
- **Standards** (*default rules* – etwa die „automatische“ Mitgliedschaft in Versicherungen, wie oben im Pennsylvania/New Jersey-Beispiel<sup>58</sup>),
  - **Vereinfachung** oder Komplexitätsreduzierung,

<sup>56</sup> Dazu oben Rz. 539.

<sup>57</sup> *Sunstein*, *Nudging: A Very Short Guide*, *Journal of Consumer Policy* 37 (2014), 583 ff.

<sup>58</sup> Dazu oben Rz. 537.

- den Einsatz **sozialer Normen** (bspw. eine Information darüber, wie sich die meisten anderen Menschen – vernünftig – verhalten)
- **Erleichterungen und Komfort** (bspw. die Präsentation von Obst in Augenhöhe)
- **Information und Aufklärung** (etwa zum Effektivzins eines Darlehens oder zur Energieeffizienz von Haushaltsgeräten)
- (**plastische**) **Warnungen** (z. B. auf Zigarettenschachteln)
- **Selbstbindungsstrategien** (etwa betreffend sportlicher Aktivität oder Spielsperren beim Glücksspiel)
- **Erinnerungen** (bspw. durch SMS oder E-Mails an fällige Rechnungen oder den anstehenden Vorsorgetermin<sup>59</sup>)
- **Abfrage von Handlungsintentionen** (z. B. hinsichtlich der Teilnahme an der Wahl oder Impfungen der eigenen Kinder)
- **Information über die Konsequenzen eigener Entscheidungen in der Vergangenheit** (etwa zu Strom- oder Krankheitskosten).

Regulierungskonzepte wie die Quellensteuer<sup>60</sup> oder die Standardeinstellung für Versicherungspolice in Pennsylvania und New Jersey<sup>61</sup> fügen sich also, wie es dieser Ansatz propagiert, nahtlos in die Systematik des *Nudging* ein.

### 3. Kritik

Die Vorstellung, dass Regierungen auf *Nudging* zurückgreifen, um Entscheidungen der Bürger zu beeinflussen, polarisiert – vor allem unter Juristen. Warum eigentlich, wenn es sich doch, verglichen mit klassischen Regulierungsinstrumenten, um „sanftere“ Maßnahmen handelt? Während eine Fraktion die sich gleichsam selbständig vollziehenden – und damit besonders mit Blick auf die Befolgungskontrolle kostengünstigen – verhaltenswissenschaftlichen Regulierungsinstrumente als faszinierenden und attraktiven Weg der Politikgestaltung betrachtet, beschwört die andere Fraktion die Gefahr eines überfürsorglichen Staates herauf, der mit „Psycho-Tricks“ seine Bürger manipuliert. Die Kritik hat zwei Stoßrichtungen: Zum einen erscheint der Gedanke einer Beeinflussung der Bürger durch „subtile“ oder schlicht manipulatorische Techniken mit der verfassungsrechtlich verbürgten Idee einer freien Willensbildung unvereinbar. Es

552

<sup>59</sup> Vgl. Altmann/Traxler, Nudges at the Dentist, European Economic Review 72 (2014), 19 ff.

<sup>60</sup> Dazu oben Rz. 536.

<sup>61</sup> Dazu oben Rz. 537.

geht dabei weniger um die Akzeptanz des „Ob“ eines staatlichen Eingriffs als vielmehr um das „Wie“, also die Art und Weise der Verhaltenssteuerung. Zum anderen stellen sich aus demokratietheoretischer Sicht Legitimationsfragen. So können gerade in den Vereinigten Staaten von Amerika viele *Nudges* exekutiv, also durch die Verwaltung vorgeschrieben werden; und auch wenn sie den Bürgern weitgehende Entscheidungsfreiheit belassen, greifen sie oftmals doch in die Rechte etwa der Anbieter von Gütern oder Dienstleistungen ein. Für den Hersteller eines energieintensiven Kühlschranks kann eine Pflicht zur Auszeichnung der Energieeffizienz schwerwiegende Nachteile im Markt mit sich bringen.

553 Der offensichtlichen Missbrauchsgefahr einer Beeinflussung von Menschen ins Auge blickend formulierte *Thaler* im Nachhinein Kriterien für „ethische“ *Nudges*: So sollen *Nudges* transparent und durch ein hinreichendes Allgemeinwohlinteresse gedeckt sein. Zudem müsse es einfach sein, sich gegen das nahegelegte Verhalten zu entscheiden. Allerdings wird auch mit diesen Einschränkungen das Problem nicht gelöst, dass wir als Individuen aus verhaltenswissenschaftlicher Sicht über mehrere Präferenzordnungen (*multiple selves*<sup>62</sup>) verfügen (können) – dieses Ich, das sich den Schokoladenkonsum oder den Casino-Gang verbietet, und jenes Ich, das der Versuchung dann doch nicht widerstehen möchte. Auch mit sanften Schubsern wird das vernünftigere Ich durch staatliche Intervention bevorzugt, obwohl es dafür freiheitsrechtlich – jedenfalls, wenn es keine gesellschaftlich relevanten Implikationen gibt – keine Rechtfertigung gibt.

554 Die Legitimationsfrage, die letztlich darauf abzielt, ob es für *Nudging*-Maßnahmen einer gesetzlichen Grundlage bedarf, ist jeweils vor dem Hintergrund einer konkreten verfassungsrechtlichen Ordnung zu beantworten. In Deutschland müssen die meisten Maßnahmen (ob die *Opt-out*-Regel bei der Organspende oder die Ampel-Kennzeichnung bei Lebensmitteln) der Wesentlichkeitstheorie des BVerfG folgend gesetzlich angeordnet werden und sind daher auch einer relativ breiten gesellschaftlichen Debatte ausgesetzt. Damit ist beiden Kritikpunkten begegnet: Gesetzgeberische *Nudges* weisen einerseits eine höhere Legitimation auf und bergen andererseits weniger die Gefahr, intransparent-manipulativ zu sein. Wo *Nudges* nicht vom Gesetzgeber verfügt werden (etwa bei Warnungen vor Lebensmitteln<sup>63</sup> oder vor Religionsgemeinschaften<sup>64</sup> durch die Bundesregierung) greift eine fein ausdifferenzierte Verfassungsrechtsdogmatik, die klare An-

<sup>62</sup> Dazu oben Rz. 541 f.

<sup>63</sup> BVerfGE 105, 252 (Glycolwarnung [2002]).

<sup>64</sup> BVerfGE 105, 279 (Osho-Bewegung [2002]).

forderungen an staatliches Informationshandeln bereithält; sie wurde schon entwickelt, lange bevor Konzepte wie das *Nudging* populär wurden.

Viele der Instrumente, die heute unter dem Begriff *Nudging* zusammengefasst werden, sind schon von deutschen Gerichten geprüft und an den überkommenen rechtlichen Maßstäben gemessen worden. Wenn so unterschiedliche Dinge wie GPS-Routenplanung, Energieeffizienzsiegel und *Default*-Regeln für Organspenden zum Gegenstand einer einheitlichen Debatte werden, dann gewinnt man aus rechtlicher Sicht den Eindruck, man habe es mit einem Gemischtwarenladen regulatorischer Ansätze zu tun; und da scheint es vernünftiger, jede Maßnahme einzeln und für sich zu diskutieren. Vor besondere Herausforderungen haben diese Instrumente das Recht bislang nicht gestellt. Es verwundert daher nicht, dass *Nudging* im hiesigen Rechtsraum aufgrund seiner Popularität in den USA und im Vereinigten Königreich zwar diskutiert wird, aber eben doch mit einer gewissen Gelassenheit (oder gar Gleichgültigkeit).

555

#### 4. Rhetorisches Mittel?

Bei der Bewertung des *Nudging* ist die mitunter grundlegend divergierende Einschätzung im deutschen und anglosächsischen Rechtssystem bemerkenswert. Während in Letzterem die Idee von *Nudging* seitens der Regierungen positiv aufgenommen wurde, schlossen sich kontinentaleuropäische Regierungen dem verhaltenswissenschaftlichen Ansatz später und zögerlicher an. Gleichzeitig wird die *Nudging*-Debatte im angelsächsischen Rechtsraum sehr kontrovers geführt, während namentlich in Deutschland Juristen so recht nichts Neues an den Instrumenten und ihrer rechtlichen Bewertung zu entdecken vermögen. Warum waren die angelsächsischen Gesellschaften, die sonst dem Staat so misstrauisch gegenüberstehen, Pioniere verhaltenswissenschaftlicher Politikgestaltung? Und woher kommt die kontinentaleuropäische Zurückhaltung gegenüber *Nudging*?

556

Die angelsächsische Debatte ist von einer tiefen Skepsis gegenüber staatlichen Interventionen gekennzeichnet, was Regulierung zu einem schwierigen Unterfangen macht – beinahe jede Intervention wird als Eingriff in individuelle Freiheiten und als **staatlicher Paternalismus** gesehen. Der Staat wird als eigenständiger Akteur betrachtet, vor dem Bürger sich im Zweifel in Acht nehmen und schützen müssen, damit er nicht zu viel Macht über sie erlangt. Steuern zu erhöhen (mit denen sich der Staat auf Kosten der Bürger finanziert) oder eine Helmpflicht für Fahrradfahrer einzuführen (und damit die Bürger dergestalt überfürsorglich zu bevormunden, dass sie sich selbst schützen müssen) ist praktisch unmöglich – zumal in

557



Systemen mit Mehrheitswahlrecht, in denen der Regierung eine starke Opposition mit Veto-Möglichkeiten gegenübersteht. *Nudges* und das Konzept eines „sanften“ oder **libertären Paternalismus** können somit als Ansatz verstanden werden, der Politik (und oftmals der Exekutive, die sanfte Eingriffe ohne Gesetzgebung vornehmen kann) Spielräume zu eröffnen, so dass die politischen Kräfte (wieder) eine gemeinsame Basis für Verhandlungen finden – oder als Beschwichtigungsstrategie, um staatliche Regulierung schmackhafter zu machen. In Kontinentaleuropa und vor allem im deutschen Verfassungsrechtsdiskurs wird der Staat nicht so sehr als unabhängige, aus sich selbst heraus bestehende Entität betrachtet, sondern als **Mechanismus kollektiver Selbstbindung**. Regulierung gilt damit nicht allein als Einschränkung individueller Rechte, sondern jedenfalls auch als zu rechtfertigender Ausgleich von Individualinteressen und Gemeinwohlbelangen. Damit geht eine hohe Bedeutung der Balance zwischen individuellen und kollektiven Rechten einher. Gelingt diese Balance, ist auch das Erhöhen von Steuern (etwa zur Finanzierung von öffentlichen Gütern) oder die Einführung der Helmpflicht<sup>65</sup> (etwa zur Reduzierung der Gemeinschaftskosten im Gesundheitssystem) kein Problem. Vor diesem Hintergrund verwundert es nicht, dass einerseits angelsächsische Regierungen das Konzept mit besonderem Eifer verfechten, damit aber auf harsche Kritik stoßen, während sich andererseits kontinentaleuropäische Regierungen dem *Nudging* nur verhalten zuwenden und dabei auf Juristen treffen, die jeweils auf die einzelne Intervention schauen und sie nüchtern an bewährten Rechtsgrundsätzen messen. Das nährt den Verdacht, dass das Konzept des *Nudging* am Ende nicht mehr ist als ein **rhetorisches Mittel**.

558 Man mag über die wissenschaftliche Unterscheidungskraft und den rechtspolitischen Nutzen des *Nudging*-Konzepts mit Recht unterschiedlicher Auffassung sein; es lediglich als Kampfmittel in einer polarisierten politischen Auseinandersetzung zu begreifen, wird der vielfältigen Bedeutung verhaltenswissenschaftlicher Einsichten bei der Gestaltung moderner Rechtsordnungen nicht gerecht. Die Zurückhaltung vor allem seitens deutscher Juristen dürfte daher schließlich auch mit unterschiedlichen Grundannahmen hinsichtlich der Stellung der Rechtswissenschaft (und der Juristenausbildung) zusammenhängen. Vor allem die US-amerikanische Rechtswissenschaft ist stark politikorientiert und richtet sich eher an den Gesetzgeber denn an die Rechtsanwender in Verwaltung und Justiz. In

<sup>65</sup> Vgl. dazu aber *Morell*, Die Rolle von Tatsachen bei der Bestimmung von «Obliegenheiten» im Sinne von § 254 BGB am Beispiel des Fahrradhelms, AcP 214 (2014), 387ff.

Mitteleuropa ist das typischerweise umgekehrt. Die Gesetzgebungswissenschaft fristet ein Schattendasein, und empirische Untersuchungen zu Verhaltensreaktionen werden unter Juristen kaum diskutiert. Im Gegensatz zu ihren US-Kollegen wenden sich deutsche Juristen, wenn sie sich äußern, in erster Linie an Gerichte und an die (nicht politische) Verwaltung; sie formulieren ihre politischen Argumente als Sache der „richtigen Auslegung“ des Gesetzes, die sie mit Hilfe von Hermeneutik und Dogmatik zum Ausdruck bringen. Weil sich diese Ausrichtung auch in der Juristenausbildung niederschlägt, sind viele deutsche Juristen schlecht gerüstet, Fragen der Verhaltensintervention zu diskutieren. Sie stützen sich meist auf Argumente des gesunden Menschenverstands, auf „anekdotische Evidenz“ statt auf empirische Daten. Verhaltenssteuerung ist bisher eher ein blinder Fleck. Die Debatte um den Einsatz von *Nudges* zeigt, dass die funktionale Perspektive auf Recht zur (möglichst effektiven) Verhaltenssteuerung an Bedeutung gewinnt und sich der Fokus stärker auf die empirische, verhaltensbezogene Dimension des Rechts richtet. Interdisziplinäre Bezüge zu Ökonomie und Psychologie werden häufiger, und in Forschung und Lehre zeigen Studiengänge wie *Law and Economics* und ökonomische Abteilungen in einer Reihe vormals rein rechtswissenschaftlich geprägter Institute den Wandel deutlich an.

## V. Offene Fragen

Der *homo oeconomicus* war vielen ein zu kantiger Zeitgenosse. Sein verhaltensökonomischer Wiedergänger wirkt dagegen sympathischer und vor allem menschlicher. Diese Menschlichkeit hat freilich ihren Preis: Je weiter sich die Ökonomik von ihrem traditionell sparsamen Verhaltensmodell entfernt und der komplexen Realität die Schleusen öffnet, desto größer wird das Risiko, dass sie von einer Flutwelle der Information überschwemmt wird. Eine Landkarte im Maßstab 1:1 ist wenig hilfreich. Ähnlich stehen auch Verhaltensökonomien vor der schwierigen Frage, wie viel Komplexität ein Modell menschlichen Verhaltens verträgt. Dieser Zielkonflikt zwischen Realitätsnähe und Sparsamkeit der Modellannahmen ist noch nicht gelöst. Er erscheint aber weniger gewaltig, wenn man sich vergegenwärtigt, dass es der ökonomischen Analyse des Rechts nicht um ein abstraktes Menschenbild, sondern um die Lösung konkreter und oft kleinteiliger Probleme geht. Diese Einsicht eröffnet die Möglichkeit problemspezifischer Mini-Verhaltenstheorien, die vom Standardansatz ausgehen, aber zusätzlich diejenigen Erträge der psychologischen For-

559

schung berücksichtigen, die im jeweiligen Kontext von besonders großer Relevanz sind.

- 560 Viele Vertreter der Verhaltensökonomik wollen in Zukunft jedoch höher hinaus. Was ihnen vorschwebt, ist ein *Grand Design* – eine neue eigenständige Verhaltenstheorie, die *Rational Choice* ablöst oder zu einem situationsspezifischen Unterfall degradiert. Man bestreitet nicht die grundsätzliche Nützlichkeit verhaltensökonomischer Forschung, wenn man derartige Ambitionen jedenfalls momentan für aussichtslos hält. Zu disparat erscheinen noch die Erkenntnisse. Selbst das Verhältnis einzelner ähnlicher Anomalien zueinander ist teilweise unklar und weniger von Verständnis als von (freilich gut begründeter) Spekulation getragen. Zudem sind die Grenzen ihres Auftretens alles andere als eindeutig. Teilweise sorgt eine geringfügige Umformulierung des Entscheidungsproblems dafür, dass ein Effekt komplett verschwindet und die Voraussagen der Rationaltheorie rehabilitiert werden.
- 561 Die Verhaltensökonomik hat einige Fortschritte in Richtung einer Systematisierung und Verknüpfung ihrer Beobachtungen gemacht. Trotzdem fehlt ihr noch größtenteils das theoretische Bindemittel. Es ist gut denkbar, dass ein verbessertes Verständnis der Funktionsweise unseres Gehirns in dieser Hinsicht irgendwann Abhilfe schaffen kann, indem es die Mechanismen „hinter“ vielen der beobachteten Effekte freilegt und Zusammenhänge sichtbar werden lässt. *Daniel Kahneman* hat kürzlich ein vielbeachtetes Modell vorgelegt, das auf der Unterscheidung zwischen einem langsamen, bewussten und einem intuitiven, automatischen Denkprozess beruht.<sup>66</sup> Dieser Ansatz entfernt sich weit vom klassischen Nutzenmaximierungsmodell, kann aber viele der beobachteten Anomalien erklären. Auch von der Verschmelzung von Ökonomik und Neurowissenschaft erhoffen sich manche mehr Klarheit.<sup>67</sup> Viele Ökonomen sprechen dem *Neuroeconomics*-Ansatz aber bislang noch seinen Wert als Erkenntnisquelle ab.
- 562 Ungeklärt ist ferner die normative Dimension der Verhaltensökonomik. Deren primäres Angriffsziel war zwar zweifellos die positive ökonomische Theorie menschlichen Verhaltens. Dabei ist es aber schon früh zu theoretischen Kollateralschäden gekommen. Die deskriptiven Einsichten der Verhaltensökonomien vertragen sich nicht immer mit der normativen Metrik der traditionellen Wohlfahrtsökonomie. Ein Beispiel hierfür liefert das

<sup>66</sup> *Kahneman*, Maps of Bounded Rationality: Psychology for Behavioral Economics, *American Economic Review* 93 (2003), 1449 ff.

<sup>67</sup> Vgl. *Camerer/Loewenstein/Prelec*, Neuroeconomics: How Neuroscience Can Inform Economics, *Journal of Economic Literature* 43 (2005), 9 ff.

**Coase-Theorem.**<sup>68</sup> Was bedeutet es, dass ein Gut in die Hände desjenigen gelangt, der es am meisten schätzt, wenn sich die Wertschätzung abhängig von den gegenwärtigen Besitzverhältnissen verändert, also das Sein das Bewusstsein bestimmt? Wenn jemand bereit ist, für ein Gut € 100 zu bezahlen, es aber nicht für unter € 150 verkaufen würde, während ein anderer dafür € 120 ausgeben würde, es sich aber für € 140 schon wieder abkaufen ließe, gerät die normative Analyse schnell aus den Fugen. Was soll nachgeahmt werden, wenn die Voraussetzungen des *Coase*-Theorems nicht gelten – Marktverhalten oder Marktergebnisse?

Ihre Ambivalenz macht es ferner sehr schwierig, die Auswirkungen vieler Effekte zu bewerten. Am Beispiel des unrealistischen Optimismus haben wir gesehen, dass ein *debiasing*, d. h. die Neutralisation des Effektes, nicht immer die beste Antwort sein kann. Und kaum jemand würde dem Menschen seine Fairnesspräferenzen aberziehen wollen. Wie man manche Effekte bewertet, hängt auch davon ab, worin man ihre Ursachen erblickt. Einige Forscher betrachten die Erkenntnisse der Verhaltensökonomien durch die Linse der **Evolutionstheorie**.<sup>69</sup> Manche scheinbar kognitiv verankerte „Fehler“ entpuppen sich dann als Ausdruck einer evolutionär adaptiven Strategie. Ob diese Strategie heute noch adaptiv ist oder ob sich die Rahmenbedingungen geändert haben, kann dann die normative Analyse beeinflussen. Die Verhaltensökonomik verhält sich gegenüber derartigen Deutungen bislang „agnostisch“.

Selbst wenn sich dies irgendwann einmal ändern sollte, wird die Evolutionsbiologie nicht alle normativen Fragen klären können. Schon heute scheint erkennbar, dass manche beobachteten Effekte aus dem Baukasten der Verhaltensökonomik zwar lokal (also im Einzelfall) nicht-rational, aber global – d. h. betrachtet auf die Lebensspanne – sehr vernünftig sein können. Solche Ambivalenz lässt sich nicht wegerklären. Heuristiken mögen uns in Einzelfällen zu falschen Einschätzungen verleiten, helfen uns aber, durch eine Welt zu navigieren, deren Komplexität uns ansonsten lähmen würde. Ein Notarzt – so ein bekanntes Beispiel – hat nicht die Möglichkeit, bei Einlieferung eines Schwerverletzten nochmals sein komplettes

<sup>68</sup> Dazu oben Rz. 258 ff.

<sup>69</sup> Vgl. *Cosmides/Tooby*, The Past Explains the Present: Emotional Adaptations and the Structure of Ancestral Environments, *Ethnology and Sociobiology* 11 (1990), 375 ff.; *Waldman*, Systematic Errors and the Theory of Natural Selection *American Economic Review* 84 (1994), 482 ff.; ferner *Jones*, Time-Shifted Rationality and the Law of Law's Leverage: Behavioral Economics Meets Behavioral Biology, *Northwestern University Law Review* 95 (2001), 1141 ff.

Medizinstudium zu rekapitulieren. Er muss vielmehr eine schnelle Entscheidung auf der Grundlage einer bewährten Daumenregel treffen.

- 565 Als einen weiteren Versuch, der Verhaltensökonomik ein normatives Fundament zu verschaffen, lässt sich ferner die sogenannte ökonomische **Glücksforschung** interpretieren.<sup>70</sup> Aus Sicht der neoklassischen ökonomischen Analyse maximieren Individuen konsequent ihren Nutzen, weswegen man sie möglichst mit Regulierungen verschonen sollte. Wenn ein Individuum jedoch konfligierende Präferenzen haben kann, wie die kognitive Psychologie nahelegt, unterminiert das diese Prämisse. Wenn wir wüssten, was Menschen tatsächlich langfristig glücklich macht, so hoffen einige, könnten wir ihre „realen“ Präferenzen identifizieren und wieder etwas normativen Boden unter den Füßen gewinnen. Bislang sind die Forschungen in diesem Bereich aber noch nicht weit genug vorangeschritten, um tatsächlich konkrete politische Maßnahmen darauf zu gründen. Zudem wirft die Glücksforschung philosophische Fragen auf – u. a. über die Legitimität des Versuchs, Menschen zu ihrem Glück zu zwingen –, die komplizierter sind als manche Vertreter der Forschungsrichtung wahrhaben wollen.

---

<sup>70</sup> Vgl. oben Rz. 70.

## Zu den Autoren

*Emanuel V. Towfigh*, geb. 1978, Inhaber des Lehrstuhls für Öffentliches Recht, Empirische Rechtsforschung und Rechtsökonomik an der EBS Law School und Professor für Rechtsökonomik an der EBS Business School, EBS Universität für Wirtschaft und Recht, Wiesbaden; Studium der Rechtswissenschaften und der Volkswirtschaftslehre in Münster und Nanjing; 2005 Promotion zum Dr. iur. in Münster; 2003–07 Wissenschaftlicher Mitarbeiter am Kommunalwissenschaftlichen Institut der Westfälischen Wilhelms-Universität Münster; 2007–15 Wissenschaftlicher Referent am Max-Planck-Institut zur Erforschung von Gemeinschaftsgütern, Bonn; 2011–12 Hauser Research Scholar und Global Research Fellow, New York University School of Law; 2012–13 Visiting Professor of Law, University of Virginia School of Law; 2014 Habilitation in Münster

*Niels Petersen*, geb. 1978, Inhaber des Lehrstuhls für Öffentliches Recht, Völker- und Europarecht sowie empirische Rechtsforschung an der Westfälischen Wilhelms-Universität Münster; Studium der Rechts- und Sozialwissenschaften in Münster, Genf und New York; 2004–06 Wissenschaftlicher Mitarbeiter am Max-Planck-Institut für ausländisches öffentliches Recht und Völkerrecht, Heidelberg; 2008 Promotion zum Dr. iur. in Frankfurt a. M.; 2007–14 Wissenschaftlicher Referent am Max-Planck-Institut zur Erforschung von Gemeinschaftsgütern, Bonn; 2009–10 M. A. in Quantitativen Methoden an der Columbia University; 2012–13 Hauser Research Scholar an der New York University School of Law; 2014 Habilitation in Bonn

*Markus Englerth*, geb. 1980, Rechtsanwalt; Studium der Rechtswissenschaften und der Volkswirtschaftslehre in Bonn und London; 2006–08 Wissenschaftlicher Mitarbeiter am Deutschen Bundestag; 2008–12 Wissenschaftlicher Referent am Max-Planck-Institut zur Erforschung von Gemeinschaftsgütern; 2010 Promotion zum Dr. iur. in Bonn

*Sebastian J. Goerg*, geb. 1980, Assistant Professor of Economics an der Florida State University, Tallahassee; Studium der Volkswirtschaftslehre in Bonn; 2007–08 Visiting Researcher am Antai College of Economics & Management, Shanghai Jiaotong University; 2010 Promotion zum Dr. rer. pol. in Bonn; 2009–12 Wissenschaftlicher Referent am Max-Planck-Institut zur Erforschung von Gemeinschaftsgütern, Bonn; 2011–12 Visiting Research Scholar an der University of Michigan School of Information

*Stefan Magen*, geb. 1966, Inhaber des Lehrstuhls für Öffentliches Recht, Rechtsphilosophie und Rechtsökonomik an der Ruhr-Universität Bochum; Studium der Philosophie und der Rechtswissenschaften in Frankfurt a. M. und Bochum; 1998–2001 Wissenschaftlicher Mitarbeiter am Bundesverfassungsgericht; 2003 Promotion zum Dr. iur. in Frankfurt a. M.; 2001–10 Wissenschaftlicher Referent am Max-Planck-Institut zur Erforschung von Gemeinschaftsgütern, Bonn; 2005 Visiting Scholar an der University of California at Berkeley School of Law; 2010 Habilitation in Bonn

*Alexander Morell*, geb. 1980, Wissenschaftlicher Referent am Max-Planck-Institut zur Erforschung von Gemeinschaftsgütern, Bonn; Studium der Rechtswissenschaften in Bonn und Paris; 2008 Visiting Scholar an der University of California at Berkeley School of Law; 2010 Promotion zum Dr. iur. in Bonn; 2015 Promotion zum Dr. rer. pol. in Jena

*Klaus Ulrich Schmolke*, geb. 1975, Inhaber des Lehrstuhls für Bürgerliches Recht, Handels-, Gesellschafts- und Wirtschaftsrecht an der Universität Erlangen-Nürnberg; Studium der Rechtswissenschaften und der Geschichte in Trier, Lausanne und Mainz; 2003 Promotion zum Dr. iur. in Mainz; 2006 Master of Laws (LL.M.), New York University School of Law; 2006–09 Wissenschaftlicher Angestellter am Institut für Handels- und Wirtschaftsrecht der Rheinischen Friedrich-Wilhelms-Universität Bonn; 2009–12 Wissenschaftlicher Referent am Max-Planck-Institut für ausländisches und internationales Privatrecht, Hamburg; 2012 Habilitation in Hamburg

## Glossar

**Abhängige Variable:** In einer empirische Studie werden in der Regel eine abhängige und eine oder mehrere unabhängige Variablen beobachtet. Die abhängige Variable ist dabei der Faktor, dessen Veränderung in Reaktion auf eine Änderung in der → unabhängigen Variable beobachtet werden soll. Mit anderen Worten untersucht man, ob Veränderungen der unabhängigen Variable Auswirkungen auf die abhängige Variable haben. (Rn. 387)

**Allokation:** Verteilung der verfügbaren Produktionsfaktoren auf die unterschiedlichen Verwendungsmöglichkeiten (Rn. 80).

**Arrows-Theorem:** Theorem des Ökonomen Kenneth Arrow, mit dem die Unzulänglichkeit herkömmlicher Abstimmungsverfahren bewiesen werden sollte. Arrow hat gezeigt, dass es, wenn man bestimmte Minimalanforderungen an ein Wahlverfahren stellt, keine Möglichkeit gibt, kumulierte individuelle → Präferenzen in eine kohärente kollektive Präferenz zu übersetzen. (Rn. 376)

**Bayes-Theorem:** Theorem in der Wahrscheinlichkeitstheorie, das es erlaubt, aus einer bestimmten Beobachtung Rückschlüsse auf die ursächlichen Faktoren zu ziehen, wenn die bedingten Wahrscheinlichkeiten bekannt sind. (Rn. 507)

**Bounded Rationality:** Eingeschränktes rationales Verhalten. Auf Grund von Einschränkungen der kognitiven Fähigkeiten wird die Suche nach einer optimalen Entscheidung frühzeitig gestoppt. Bei uneingeschränkter Rationalität wird solange nach einer Entscheidung gesucht, bis eine nutzenmaximierende Lösung gefunden wurde. Eingeschränkt rationales Verhalten ist nicht mit irrationalem Verhalten zu verwechseln. Der Begriff wurde durch den Nobelpreisträger *Herbert Simon* geprägt. (Rn. 492, 504)

**Coase-Theorem:** Theorem, das nach dem Nobelpreisträger *Ronald Coase* benannt ist und besagt, dass durch den Markt eine effiziente Allokation von Ressourcen erreicht werden kann, wenn Eigentumsrechte eindeutig definiert sind und es keine → Transaktionskosten gibt. Dieses Ergebnis ist unabhängig davon, wie die Güter ursprünglich verteilt gewesen sind. Die beiden Annahmen, die Coase macht, sind in der Praxis allerdings oft nicht gegeben. Zum einen ist es teilweise nicht möglich, Eigentumsrechte eindeutig zu definieren – wem gehört beispielsweise saubere Luft? Zum anderen sind wirtschaftliche Transaktionen in der Realität immer auch mit entsprechenden Kosten verbunden. Diese beiden Abweichungen von den Annahmen führen dazu, dass es in der Praxis dazu kommen kann, dass der Markt nicht immer eine effiziente Verteilung von Gütern bereit stellen kann. (Rn. 251)



**Diskontierungsrate:** Zinssatz, der den Wert von in der Zukunft liegenden Geldflüssen angibt. (Rn. 305)

**Diskrete Variable:** Variable, die entweder nur eine endliche Zahl von Werten annehmen kann, oder deren Werte zumindest zählbar sind. Diskrete Variablen haben also Sprünge. Beispielsweise kann die Bevölkerungszahl eines Ortes nur in ganzen Zahlen, nicht jedoch in halben oder viertel Einheiten gemessen werden. (Rn. 413)

**Effizienz:** In der Ökonomie bezeichnet Effizienz das Maß, anhand dessen bewertet wird, inwiefern Ressourcen und Güter optimal zugewiesen sind, um die gesellschaftliche Wohlfahrt zu maximieren. Bei der Beantwortung der Frage, was „optimal“ ist, helfen uns das → Pareto- und das → Kaldor-Hicks-Kriterium. (Rn. 79)

**Empirie:** Sammlung von Informationen, die auf gezielten Beobachtungen beruht. Empirische Forschung versucht also aufgrund von Beobachtungen, die entweder im Feld oder im Labor gewonnen werden, deskriptive oder erklärende Aussagen über bestimmte Phänomene zu machen (Rn. 59).

**Experiment:** Methode der Beobachtung, bei der die Beobachtungen in zwei oder mehrere Gruppen aufgeteilt werden. Zwischen diesen Gruppen gibt es dabei in der Regel nur einen systematischen Unterschied. Wenn sich die Beobachtungen in den einzelnen Gruppen unterscheiden, kann davon ausgegangen werden, dass dieser Unterschied in den Beobachtungen (→ abhängige Variable) auf den systematischen Unterschied in den Gruppen (→ unabhängige Variable) zurückzuführen ist. (Rn. 394)

**Externalitäten** → externe Effekte

**Externe Effekte:** Auswirkungen von Entscheidungen auf unbeteiligte Dritte. Da keine direkte Beziehung zwischen Verursacher und Betroffenen besteht, können diese nicht über Marktmechanismen kompensiert werden. Externe Effekte können sowohl positiv (dann externer Nutzen) als auch negativ (dann externe Kosten) sein. (Rn. 150, 209)

**Gefangenen-Dilemma:** Mathematisches Spiel aus der Spieltheorie, das ein soziales Dilemma modelliert, bei dem individuell-rationales Verhalten in ein kollektiv unerwünschtes Endergebnis mündet. Addiert man die Auszahlungen der beiden Spieler, erreichen diese das beste Ergebnis, wenn sie miteinander kooperieren. Individuell haben sie jedoch Anreize, nicht miteinander zu kooperieren. (Rn. 208)

**Gleichgewicht:** In der Spieltheorie liegt dann ein Nash-Gleichgewicht vor, wenn keiner der Akteure Anreize hat, sein Verhalten zu ändern. Das ist dann der Fall, wenn sich alle Akteure unter Berücksichtigung des Verhaltens aller anderen Akteure optimal verhalten. (Rn. 179) — In der Wirtschaftstheorie spricht man von einem Marktgleichgewicht, wenn die Menge des Angebots dem der nachgefragten Menge entspricht. (Rn. 138)

**Grenzkosten:** Beschreiben den Anstieg der Kosten durch die Produktion einer weiteren Einheit. (Rn. 130)

**Grenznutzen:** Zuwachs des Nutzens, welcher durch ein Gut generiert wird, wenn eine weitere Einheit eben dieses Gutes konsumiert wird. (Rn. 72)

- Heuristik:** Kognitive Daumenregel, die Komplexität reduziert und uns hilft, uns innerhalb einer hochkomplexen Welt zurechtzufinden. Sie werden von Menschen insbesondere dann zur Entscheidungsfindung oder Problemlösung genutzt, wenn entweder nicht alle notwendigen Informationen zur Verfügung stehen oder diese Informationen zu komplex sind, als dass sie verarbeitet werden könnten. (Rn. 508)
- Indifferenzkurve:** Mögliche Mengenkombination verschiedener (im einfachsten Fall zweier) Güter, welche für den Konsumenten/Haushalt denselben Nutzen stiften. Alle möglichen Kombinationen der Güter, welche auf der Indifferenzkurve liegen, werden als gleichwertig beurteilt. (Rn. 98)
- Informationsasymmetrie:** Form des Marktversagens, die dadurch entsteht, dass einer Vertragspartei bei Vertragsschluss mehr Informationen über den Vertragsgegenstand zur Verfügung stehen als der anderen. Wenn die Partei mit dem überlegenen Wissen diese Informationsasymmetrie zu ihren Gunsten ausnutzt, kommt es zu ineffizienten Verträgen, da die andere Partei den Vertrag unter vollständigen Informationen rationalerweise nicht geschlossen hätte. (Rn. 264)
- Internalisierungsproblem:** Situation, in der Individuen aufgrund von → externen Effekten nicht alle Kosten ihrer Handlungen berücksichtigen („internalisieren“), so dass aus wohlfahrtsökonomischer Sicht Fehlanreize entstehen (Rn. 161).
- Kaldor-Hicks-Effizienz:** Nach diesem Kriterium der → Effizienz wird eine Änderung in der Güterallokation schon dann als wohlfahrtssteigernd und damit effizienter angesehen, wenn aus den Gewinnen der bessergestellten Akteure die Verluste der Verlierer kompensiert werden können, und zusätzlich zumindest ein Akteur nach dieser Kompensation besser steht als im alten Zustand. (Rn. 85)
- Kardinalskala:** Skala, die Variablen in ein metrisches Verhältnis zueinander setzt. Dabei sind zwei Arten der Kardinalskala zu unterscheiden. Die Intervallskala misst den Abstand zwischen zwei Variablen, hat jedoch keinen natürlichen Nullpunkt, so dass Rechenoperationen wie die Multiplikation nicht durchgeführt werden können (Beispiel: Temperaturskala). Dagegen hat die Verhältnisskala einen natürlichen Nullpunkt, so dass die Variablen auch mithilfe einer Multiplikation oder Division miteinander ins Verhältnis gesetzt werden können. (Rn. 80, 413)
- Kausalität:** Beziehung zwischen Ursache und Wirkung. Eine Kausalbeziehung zwischen Faktor A und Faktor B liegt dann vor, wenn B durch A bewirkt worden ist, also in dieser Form nicht ohne das Auftreten von A zu beobachten gewesen wäre. (Rn. 388)
- Kollektivgut:** Auch öffentliches Gut genannt. Charakteristika eines Kollektivgutes sind die Nichtanwendbarkeit des Ausschlussprinzips sowie die Nichtrivalität im Konsum. Nicht-Ausschließbarkeit bedeutet, dass jemand nicht von der Nutzung des öffentlichen Gutes ausgeschlossen werden kann (z.B. saubere Luft). Mögliche Gründe können etwa technischer oder gesellschaftlicher Art sein. Da öffentliche Güter zeitgleich von mehreren Nutzern konsumiert werden können, spricht man von Nichtrivalität im Konsum. (Rn. 210)

- Konsumentenrente:** Die Differenz zwischen dem Marktpreis eines Gutes und dem Höchstpreis, den der Konsument für das Gut zu zahlen bereit wäre. Die Konsumentenrente ist ein Konzept der → Wohlfahrtsökonomik. (Rn. 118)
- Kontinuierliche Variable:** Variable, die einen beliebigen Wert annehmen kann, keine Sprünge hat und daher auf einer metrischen Skala messbar ist. Beispiele sind etwa Temperatur, Zeit oder Größe. (Rn. 413)
- Korrelation:** Statistischer Zusammenhang zwischen zwei Variablen. Zwei Faktoren sind miteinander korreliert, wenn wir bei der Veränderung einer der beiden Variablen gleichzeitig eine systematische Veränderung der anderen Variable beobachten können. (Rn. 451)
- Marktmacht:** Liegt vor, wenn Anbieter oder Nachfragen auf dem Markt durch fehlende Konkurrenz oder wegen fehlenden wesentlichen Wettbewerbs eine beherrschende Stellung einnehmen. Die stärkste Form der Marktmacht ist ein → Monopol (Rn. 282)
- Median:** Wert, durch den eine nach Größe geordnete Stichprobe in zwei Hälften geteilt wird. Genau die Hälfte der Stichprobe ist dabei größer als oder genauso groß wie der Median, die andere Hälfte kleiner als oder genauso groß wie der Median. (Rn. 423)
- Medianwähler-Theorem:** Theorem, das von dem Ökonomen Anthony Downs entwickelt wurde und versucht, die Konkurrenz von politischen Parteien um Wählerstimmen zu erklären. Downs beweist, dass politische Parteien sich, wenn man bestimmte Annahmen zugrunde legt, bei ihrem politischen Programm immer am Medianwähler orientieren, also an dem Wähler, von dem aus gesehen sich jeweils genau die Hälfte der Bevölkerung links und rechts im politischen Spektrum befindet. (Rn. 334)
- Methodologischer Individualismus:** Methodischer Grundstein der Ökonomik, demzufolge allein Handlungen von Individuen Gegenstand der wissenschaftlichen Analyse sind. Kollektiventscheidungen folgen danach nicht der Eigenlogik eines „Kollektivwillens“, sondern können auf das Zusammenwirken individueller Entscheidungsträger zurückgeführt und durch dieses erklärt werden. (Rn. 63)
- Mittelwert:** Durchschnitt aller Beobachtungen, der berechnet wird, indem man die Werte aller Beobachtungen addiert und dann durch die Zahl der Beobachtungen teilt. (Rn. 421)
- Modell:** Form der Beschreibung oder Erklärung von gesellschaftlichen Wirkungszusammenhängen durch Komplexitätsreduktion. In der Ökonomie werden Modelle oft in mathematischer Form beschrieben, dies muss in anderen Sozialwissenschaften aber nicht notwendigerweise der Fall sein. (Rn. 54)
- Modus:** Wert, der am häufigsten in einer (empirischen) Verteilung vorkommt. (Rn. 421)
- Monopol:** Liegt vor, wenn nur Anbieter (Angebotsmonopol) oder Nachfrager (Nachfragemonopol) vielen Nachfragern oder Anbietern oder nur Anbieter nur einem Nachfrager (zweiseitiges Monopol) gegenübersteht. (Rn. 282)

- Normalverteilung:** Eine Verteilung, welche die Form einer Glocke hat. Bei der Normalverteilung fallen  $\rightarrow$  Median,  $\rightarrow$  Modus und Erwartungswert ( $\rightarrow$  Mittelwert der Verteilung) zusammen. (Rn. 417)
- Ökonomik:** Methode der Wirtschaftswissenschaften. Die Ökonomik steht dabei im Gegensatz zur Ökonomie, die nicht Methode, sondern Untersuchungsgegenstand ist. (Rn. 3)
- Ökonometrie:** Teilgebiet der Wirtschaftswissenschaften, das versucht, wirtschaftstheoretische Modelle empirisch zu überprüfen. Ökonometrische Studien versuchen dieses Ziel vor allem durch die statistische Analyse von Felddaten zu erreichen. (Rn. 60)
- Opportunitätskosten:** Entgangener Nutzen oder entgangene Erträge, welche entstehen, wenn Handlungsalternativen nicht wahrgenommen werden. (Rn. 122)
- Ordinalskala:** Skala, die Rangordnung zwischen Variablen herstellt. Danach sind Aussagen möglich, dass  $a > b$ . Dagegen sind Rechenoperationen wie Addition oder Multiplikation mit ordinal geordneten Variablen nicht möglich. (Rn. 65, 413)
- Pareto-Effizienz:** Nach diesem Kriterium der  $\rightarrow$  Effizienz bedürfen Änderungen der Allokation von Gütern der Zustimmung aller (Einstimmigkeit), so dass jeder ein Vetorecht hat, das aber nur im Falle einer Schlechterstellung ausgeübt wird, nicht bei Indifferenz. Ein Zustand ist demnach effizienter als ein anderer, wenn durch ihn niemand schlechter, aber wenigstens einer besser gestellt wird. (Rn. 81)
- Population:** Grundgesamtheit aller statistischen Einheiten, auf die durch statistische Tests Rückschlüsse gezogen werden sollen. Wenn wir beispielsweise generelle Aussagen über das Verhalten von Richtern in Deutschland machen wollen, dann gehören zu der Population alle Richter, die jemals in Deutschland ein entsprechendes Amt bekleidet haben oder bekleiden werden. Die Stichprobe, von der aus Rückschlüsse auf die Merkmale der Population gezogen werden sollen, wird dabei wesentlich kleiner sein, da gar nicht alle Mitglieder der Population beobachtet werden können. (Rn. 386)
- Präferenz:** Präferenzen beschreiben innere Motive eines Akteurs, die unabhängig von den aktuellen Handlungsmöglichkeiten sind und die Wertvorstellungen des Individuums enthalten. In der ökonomischen Theorie wird in der Regel angenommen, dass die Präferenzen transitiv, ordinal und vollständig geordnet, dass sie interpersonal nicht vergleichbar und über die Zeit konstant sind. (Rn. 65)
- Prinzipal-Agenten-Beziehung:** Beziehung zwischen einem Prinzipal (Vorgesetzten) und einem Untergebenen (Agent). Diese Beziehung ist durch eine asymmetrische Verteilung von Informationen charakterisiert. (Rn. 294)
- Produzentenrente:** Die Differenz zwischen dem Marktpreis und dem niedrigsten Preis, zu dem ein Anbieter noch bereit wäre, ein Gut herzustellen. Wie die  $\rightarrow$  Konsumentenrente ist die Produzentenrente ein Konzept der  $\rightarrow$  Wohlfahrtsökonomik. (Rn. 134)
- Prospect Theory:** Theorie, welche die Abweichungen von (experimentell) beobachtetem Verhalten von dem erwarteten Verhaltens des *homo oeconomicus* zu erklären versucht. Hauptmerkmale der *Prospect Theory* sind die Beeinflussung

des aktuellen Verhaltens durch Referenzpunkte, die stärkere Gewichtung von Verlusten als von Gewinnen sowie die überproportionale Gewichtung von unwahrscheinlichen Ereignissen. Die *Prospect Theory* wurde durch den späteren Nobelpreisträger *Daniel Kahneman* und *Amos Tversky* eingeführt. (Rn. 533)

**Qualitative Variable:** Variable, die sich auf Kategorien bezieht und zwischen deren Werten keine logische Beziehung hergestellt werden kann (wird oft auch nominale oder kategoriale Variable genannt). Welche Kategorien dabei welcher Zahl zugeordnet werden, ist daher egal. Misst man etwa die Religion von Probanden ist egal, ob Muslime mit 1 und Christen mit 2 codiert werden oder umgekehrt. Anders als mit diskreten oder kontinuierlichen Variablen kann man mit qualitativen Variablen nicht rechnen. (Rn. 414)

**Randomisierung:** Methode, bei der Durchführung von Experimenten, durch die Testpersonen zufällig auf die unterschiedlichen Testgruppen aufgeteilt werden. (Rn. 394)

**Regression:** Analytisches Instrument, um gerichtete statistische Zusammenhänge zwischen einer  $\rightarrow$  abhängigen und einer oder mehreren  $\rightarrow$  unabhängigen Variablen bestimmen zu können. (Rn. 450)

**Signifikanz:** Wahrscheinlichkeit, mit der bei einem statistischen Test fälschlicherweise die Nullhypothese abgelehnt wird (Fehler 1. Art). Bei einem Signifikanzniveau von 5 % ist diese Fehlerwahrscheinlichkeit kleiner oder gleich 5 %. (Rn. 406 und Rn. 434)

**Stichprobe:** Beobachtungen, welche zufällig aus einer Grundgesamtheit ( $\rightarrow$  Population) gezogen werden. (Rn. 386)

**Sozialwissenschaften:** Wissenschaften, die sich mit der Beschreibung und Erklärung aller Aspekte menschlichen Zusammenlebens beschäftigen. Methodisch ist dabei zwischen positivistischer und interpretativer Sozialwissenschaft zu unterscheiden. Positivisten versuchen Aussagen über gesellschaftliche Gesetzmäßigkeiten zu machen, während Anhänger interpretativer Zugänge überwiegend annehmen, dass das Finden solcher Gesetzmäßigkeiten nicht unabhängig von der Vorprägung des Forschers möglich ist, objektive Gesetzmäßigkeiten also gar nicht gefunden werden können. Ihnen geht es vor allem darum, die Bedeutung bestimmter sozialer Phänomene zu analysieren. (Rn. 9)

**Standardfehler:** Ist, wie die  $\rightarrow$  Varianz, ein Streuungsmaß. Der Standardfehler nimmt mit dem Stichprobenumfang ab und mit der Varianz zu. (Rn. 431).

**Strategische Interdependenz:** Situationen, in denen das Resultat der Entscheidung eines Akteurs auch davon abhängt, wie sich der oder die anderen Akteure entscheiden (Rn. 162).

**Transaktionskosten:** Marktbenutzungskosten, welche das  $\rightarrow$  Coase Theorem ignoriert. Transaktionskosten entstehen durch die Benutzung von Märkten und beinhalten Kosten welche durch Verhandlungen (z.B. Anwälte) und Informationssuche entstehen. (Rn. 150, Rn. 262 und Rn. 306)

**Transitivität:** Aussage über die Relation zwischen mindestens drei Variablen. Wenn  $a > b$  und  $b > c$ , dann muss  $a > c$ , damit die Bedingung der Transitivität erfüllt ist. (Rn. 65)

**Unabhängige Variable:** In einer empirische Studie werden in der Regel eine abhängige und eine oder mehrere unabhängige Variablen beobachtet. Die unabhängige Variable ist dabei der Faktor, der verändert wird, um zu beobachten, welche Auswirkungen diese Veränderung auf die Entwicklung der abhängigen Variable hat. (Rn. 387)

**Utilitarismus:** Philosophie, die davon ausgeht, dass der moralische Wert einer Handlung allein aufgrund ihrer Konsequenzen bestimmt werden kann. Ziel ist dabei, das aggregierte Glück aller Individuen zu maximieren. (Rn. 48)

**Validität:** Generalisierbarkeit von Beobachtungen. Dabei lassen sich drei Arten der Validität unterscheiden. Die statistische Validität zeigt, ob eine Beobachtung möglicherweise auf Zufall beruht oder wir einen systematischen Effekt beobachten. Die interne Validität besagt, ob eine Beobachtung kohärent ist. Die externe Validität sagt schließlich etwas darüber aus, ob die Beobachtungen auch auf andere Kontexte übertragbar sind. (Rn. 406)

**Varianz:** Ein Maß, welches die Streuung einer Verteilung oder einer → Stichprobe wiedergibt. (Rn. 446)

**Verteilungsgerechtigkeit:** Meint die Gerechtigkeit von Verteilungsregeln (Regelgerechtigkeit) und deren Ergebnissen (Ergebnisgerechtigkeit). Die Ergebnisgerechtigkeit bezeichnet ein normatives Konzept, das solche gesellschaftlichen Gegebenheiten als gerecht bezeichnet, in denen allen Mitgliedern der Gesellschaft der gesellschaftliche Nutzen in grundsätzlich gleichem Maße zukommt. Die Regelgerechtigkeit bezieht sich demgegenüber auf solche zu den Ergebnissen führenden Regeln. (Rn. 313)

**Wohlfahrtsökonomik:** Beschäftigt sich mit der Frage, auf welche Weise gesamtgesellschaftliche (also nicht nur individuelle) Wohlfahrtssteigerungen oder soziale Optima erzielt werden können. Ihr Kernbegriff ist die → Effizienz. (Rn. 79)



# Sachwortverzeichnis

Verweise beziehen sich auf die Randziffern.

- Abschreckung 41, 519  
Abweichungsdiagramm 187, 193  
Agenda-Setter 377  
Agenda-Verfahren 374 ff., 388  
Agenturvertrag s. Prinzipal-Agenten-Beziehung  
Agenturkosten s. Prinzipal-Agenten-Beziehung  
Alignment of interests, s. Interessensangleichung  
Allmendegüter 166  
Allokation 88 ff., 150,  
– Ressourcenallokation 262  
Alternativhypothese 442 ff.  
Altruismus 179  
anchoring, s. Ankereffekt  
Angebotskurve 140 ff.  
Ankereffekt 521 ff.  
Annahmen 58 ff., 61 ff., 85 ff., 111, 217,  
Arbeitskosten 136  
Arrows-Theorem 371, 384 ff.  
Asian-Disease-Szenario 535,  
Auslegung 14 ff.  
– teleologische Auslegung 19 ff.  
Ausschließbarkeit 221 ff.  
availability heuristic, s. Verfügbarkeitsheuristik  
Aversion gegen Extreme 526 ff.  
  
backward induction, s. Rückwärtsinduktion  
Basisspiel 231 ff., 255 f.  
Bayes-Theorem 506, 514  
Behavioral Economics 481, 542  
Behavioral Law and Economics 483 ff., 492  
Besitzeffekt 529 ff.  
Bestimmtheitsmaß 471, 476,  
  
Borda-Verfahren 380 ff., 388  
Borda-Wert 380 ff.  
Bounded Rationality s. Rationalität, beschränkte  
Brennpunkte 209, 89 f.  
Budget, s. Konsummöglichkeiten  
Budgetmaximierung der Bürokraten 360 ff.  
Bürokraten 337 ff., 360 ff.,  
  
Ceteris-paribus-Annahme 105  
Chicken game, s. Feigling-Spiel  
Clubgüter, s. Klubgüter  
Coase-Theorem 161 ff., 261 f., 295, 529, 562,  
common knowledge, s. gemeinsames Wissen  
Condorcet-Paradox 378 ff.  
Condorcet-Verfahren 378 ff.  
Coordination device 267  
  
Datenmessung 409 ff.  
deadweight loss 149 f.  
Defektion 219, 226 ff.  
Demokratietheorie, kompetitive 347  
dichte Gruppen 254  
Dienstverträge 304  
Diktator-Spiel 497  
Diskontrate 232 f.,  
Diskordinierungsspiel 196  
Dominanz 184 ff.  
– schwache Dominanz 189  
– starke Dominanz 189  
– Pareto-Dominanz 217  
  
Economies of scale, s. Skalenvorteile  
effizientes Vertragsdesign 259  
Effizienz 34 f., 47, 87 ff., 164 ff. ,



- Effizienzhypothese 262 ff.
- Pareto-Effizienz 89 ff., 150, 198,
- Kaldor-Hicks-Effizienz 200 ff.
- Eigennutztheorem 61, 69 ff., 336,
- Einkommen 395
- Einkommenseffekt 118
- Einmal-Spiel, simultanes 178
- Einstimmigkeit 89
- Emissionshandelsrecht 20 ff.
- Empirie 59
- endowment effect, s. Besitzeffekt
- Entdeckungswahrscheinlichkeit 510
- Entscheidungsknoten 235 ff.,
- Entscheidungstheorie 120, 330, 505
- Entscheidungsfehler, systematische 297
- Erwartungsnutzen 74 f., 217, 494
- Erwartungsnutzentheorie 531 ff.,
- Evolutionsbiologie 564
- Experimentaldaten 412 f.
- Experimentalökonomie 59
- Extensivform 175 f., 183, 235 f.
- Externalitäten 161 f., 220 f., 225, 272
- externe Effekte 29, 161 ff.
- extremeness aversion, s. Aversion gegen
- Extreme
- Ex-post-Opportunismus, Problem des 317 ff.
  
- Fairnessnormen 498 f., 502
- Falke-Taube-Spiel 212 ff.
- Feigling-Spiel 213,
- Felddaten 412
- Feldexperiment 413, 488
- Feldstudien 412
- First mover advantage 238
- Fixkosten 136 ff., 146
- Focal points, s. Brennpunkte
- Folk-Theorem 232 f.
- Framing 82 f., 535 ff.
- F-Statistik 471
  
- Gefangenendilemma 188 ff., 203, 219 ff., 232 ff.
- Gegenseitigkeitsprinzip 498, 501
- gemeinsames Wissen 246 f.
- Gemeinwohlziele 198 ff.
- Generalprävention 41, 501
- Gesamtkosten 140
- Gesamtspiel 231 ff., 257
- Gewichtungsfunktion 533
- Gewinn 178 ff.
- Globalhaushalt 369
- Gleichgewichtsauswahl 194
- Gleichgewichtskonzepte, s. Lösungskonzepte
- Gleichgewicht
- in dominanten Strategien 184 ff.
- in gemischten Strategien 195 ff.
- in reinen Strategien 195 ff.
- Gleichgewichtspfad 242
- Gleichverteilung 425
- Glücksforschung 565
- Grenzertrag 141 f., 155
- Grenzkosten 140 ff., 169
- Grenznutzen 72 f., 76, 151
- Große Zahlen, Gesetz der 415
- Grundrechtsdogmatik 24
- Gruppenentscheidung 372
- Güter, öffentliche 224 f.
  
- Harberberger Dreieck, s. deadweight loss
- Harmoniespiele 204 ff.
- Haushalt 333, 340, 360 ff.,
- Heuristiken 298, 507 ff., 564
- hindsight bias, s. Rückschaufehler
- Hirschjagd, 210, 215 ff., 232
- Histogramm 423 ff.
- Holdout-Problem 163
- homo oeconomicus 53, 69 ff., 80, 494 ff., 506, 540, 559
- hypothetische Zahlungsbereitschaft 95
  
- Indifferenz 89, 106 ff., 116 ff., 349
- informale Institutionen 254 ff.
- Information
- imperfekte Information 244
- perfekte Information 244
- unvollständige Information 273 ff.
- Informationsasymmetrien 274, 281, 302 ff., 362
- Informationsmenge 244, 249
- Interessensangleichung 309 ff.
- Invarianzhypothese 262 f.
- inverted order paradox 382
- Isomorphismus 418
  
- Kaldor-Hicks-Kriterium 93 ff., 200, 229 f.

- Kampf der Geschlechter 194, 203, 210 ff.  
 Kapitalgesellschaft 321 ff.  
 Karteldilemma 177 ff., 220  
 Kausalität 396 ff., 458, 463,  
 Kennzahlen 419, 428  
 Klubgüter 222 ff.  
 kognitive Beschränktheit 503  
 kognitive Psychologie 481, 486, 490  
 kollektives Handeln 220  
 Kollektivgüter 221  
 Kollusion 180, 188, 257,  
 Kompensation 49, 93 f.,  
 Komplement 109 f.  
 Konfliktspiele 206 f., 211, 228  
 Konsumentenrente 125 ff., 146, 150 f.  
 Konsumentenwohlstandsstandard 126  
 Konsummöglichkeiten 112, 127 ff.  
 Konsumrivalität, s. Rivalität  
 Konventionen 255 f.  
 Kooperationsspiele 218 ff.  
 Koordinationsspiele 208 ff., 255  
 Korrelation 395 ff., 458 ff., 477  
 Korrelationskoeffizient 460  
 Kosten 139 ff., 161 kriminologische  
 Theorie 544  
  
 Labor 413  
 Lageparameter 429  
 Law & Economics 4, 80,  
 Lobbyismus 355  
 Lösungskonzepte 183, 191, 198  
  
 Markt  
 – Markteintrittsspiel 248  
 – Marktgleichgewicht 365  
 – Marktpreis 220  
 – Marktmacht 292  
 – Markträumung 530  
 – Marktversagen 153, 221, 292, 330  
 – perfekter Wettbewerbsmarkt 156  
 Mathematik 57, 112  
 maximale Zahlungsbereitschaft 530  
 Maximin-Auszahlung 233,  
 Maximin-Prinzip 214, 217,  
 Mechanismus-Design 173, 202  
 Median 431 ff.  
 Medianwähler-Theorem 342  
 Mehrheitsentscheidung 373 f.  
 Mehrheitswahlrecht 350, 557  
  
 Mehrparteiensystem 350  
 Merkmal 420 ff.  
 methodologischer Individualismus 63  
 Mikroökonomie 260  
 Mittelwert 429 ff., 452 f.  
 Modell 54 ff., 66 ff.  
 Monitoring-Problem, s. Überwachungs-  
 problem  
 Monopol 154 ff., 292 ff.  
 Moral hazard 302 f., 322  
  
 Nachfragekurve 120 ff.  
 Nash-Gleichgewicht 183, 191 ff., 213,  
 239 ff.  
 „Natur“ 109  
 Neue Institutionenökonomik, 3 f.  
 Neue politische Ökonomie 328 ff.  
 Nonpaternalismus 89  
 Normalform 175 ff.  
 Normalverteilung 425 f.  
 Normkonkretisierung 31 ff.  
 Nudging 545 ff.  
 Nullhypothese 442 ff.  
 Nullsummenspiel 207, 212  
 Nutzen 8, 69 ff., 286 ff., 332 ff., 494 ff.  
 Nutzenmaximierung 114 ff.  
  
 Öffentliche Güter 224, 329  
 Ökonometrie 59 f., 466, 478  
 Ökonomisches Paradigma 61  
 one shot games, s. Einmal-Spiele  
 Operationalisierung 411  
 Opportunitätskosten 132 f., 317  
 Optimierungsaufgabe 77  
 Ostrogorski-Paradox 390 f.  
 overconfidence bias 515 .  
  
 Pareto-Optimalität 49, 88 ff., 199, 273  
 Paternalismus 299, 537 ff., 557  
 politischer Prozess 332 ff.  
 Pooling-Vertrag, s. Durchschnittsvertrag  
 Population 394 , 420, 441  
 Präferenzen 48, 64 ff., 99, 179 ff., 198,  
 217, 228, 327 ff., 343, 372 ff., 496,  
 520, 539  
 Preis 97, 118 ff., 148 ff., 160, 252,  
 Preisdifferenzierung, perfekte 160  
 Preisnehmer 148, 154 f., 158  
 Prinzipal-Agenten-Beziehung 304 ff.  
 private Güter 166, 221

- probabilistische Aussage 13
- Produktivität 280 ff.
- Produktion
  - Produktionseffizienz 91
  - Produktionskosten 132
- Produzentenrente 146, 150 f.
- Prospect Theory 529 ff.,
- Public Choice Theory 331 ff.
- Punktwolke 459
  
- quantitatives Merkmal 421
- Quasi-Experiment 406
- Quasi-Rationalität 482
  
- Randomisierung 402
- rationale Ignoranz 335, 351
- Rationalitätsannahme 77 ff., 296, 336, 489
- Rationalität, beschränkte 296
- Rechtsdogmatik 17 ff.
- Rechtspaternalismus 299
- Rechtssetzung 40
- Regression 458, 466 ff., 471 ff.
- rent seeking 333, 355
- repeated games, s. wiederholte Spiele
- Repräsentationsheuristik 509
- Restriktionen 67 f.
- Richter 511 ff.
- Risikoaversion 76
- Risiko-Dominanz 217
- Rivalität 221 ff.
- Ressourcenallokation 262
- Ressourcenknappheit 61, 64
- Robustheit 58
- Rückschaufehler 512 ff.
- Rückwärtsinduktion 234, 243 ff.
  
- Satisficing s. Rationalität, beschränkte
- Schadensersatz 311
- Schatten der Zukunft 232
- Scheinrelationen 401
- Schleier des Nichtwissens 50
- Schwarzhandel 502
- Screening 283 ff.
- Selbstselektion 285
- Selbstüberschätzung 515, 519
- Selektion, adverse 274 ff.
- Self-serving bias 516 f.
- Separating-Verträge s. Selbstselektion
  
- Shortsightedness effect, s. politische Kurzsichtigkeit
- Signaling 277 ff.
- Signifikanz, statistische 413 ff., 442 ff.
- Signifikanzniveau 444 ff.
- Skaleneffekte 157
- Social Choice Theory 331, 372
- Sorgfaltsmaßstab 513
- soziale Normen 254 ff.
- Spieldefinition 175
- Spielmatrix 181
- Spieltheorie 170 ff.
- Staatsaufgaben 329
- Staatsversagen 330
- Stage game, s. Basisspiel
- Standardabweichung 438 ff.
- Standardfehler 439 f.
- Statistik 494 ff.
- Status Quo Bias 532 f.
- Stetigkeit 103
- Steuerungswirkung 4, 38
- Stichprobe 441 ff., 448
- Stichprobenvarianz 437
- Störfaktoren 402 ff.
- Strafmaß 41, 523
- Strategieprofil 183, 191, 227
- strategische Interdependenz 174
- Streuungsmaße, 436 ff.
- Substitution 109 ff.
- Substitutionseffekt 118
- Substitutionsrate 107 ff.
- Super game, s. Gesamtspiel
  
- Teilspielperfektion 239, 242
- Tendenz, zentrale 429
- Testverfahren 441 ff.
- Transaktionskosten 161 ff., 272
- Transitivität 101
- Tree pruning 243
- t-Statistik 472
- Turnier-Verfahren, s. Agenda-Verfahren
  
- Übernutzung 223 f.
- Überwachungsproblem 307 ff.
- Ultimatum-Spiel 82 ff., 497
- Unmöglichkeitstheorem, s. Arrows-Theorem
- Unsicherheit 78
- Unterdrückungseffekt 408
- Unterschiedshypothese 447

- unvollständige Information 246 f., 273, 292, 302
- Utilitarismus 8, 48, 53
- Validität 85, 413 ff.
- Variable 59, 85, 395 ff., 420 ff., 441, 459, 466 ff.,
- variable Kosten 136 ff.
- Varianz 448 ff.
- Verfügbarkeitsheuristik 508 ff.
- Verfügungsrecht 152
- Verhaltensmodell 69, 79 ff.
- Verhältnismäßigkeit 24
- Verhältnismäßigkeitswahlrecht 350
- Verhandlungsmacht 263, 292 f.
- Verhandlungsversagen 269
- Verlustaversion 287, 531 f.
- Verteilung 425
- Verteilungsform 448 ff.
- Verteilungsgerechtigkeit 34 f., 323
- Vertrag 260 ff.
  - Langzeitverträge 312
  - Vertragsökonomie 259
  - Vertragsrecht 262 ff.
- vollständige Information 78, 178, 247
- Vorabbindung durch Verträge 266
- Voting cycle, s. Condorcet-Paradox
- Voting paradox, s. Wahlparadox
- Wähler 334 ff.
- Wählerstimmen 332 ff.
- Wahlparadox 336
- Waste 367
- Wechselwirkung 397
- Weisungsrecht 37
- Wertfunktion 533
- wiederholte Spiele 231 ff.
- Wohlfahrtskriterien 198 ff.
- Wohlfahrtsverlust 150
- Zugfolge 178
- Zahlungsbereitschaft 71, 95, 125
- Zufallseffekt 400
- Zufallszug 247
- Zusammenhangshypothese 447