

Kabalak, Alihan

Working Paper

Kooperativität vs. Rationalität

Discussion Papers, No. 8/2011

Provided in Cooperation with:

Witten/Herdecke University, Faculty of Management and Economics

Suggested Citation: Kabalak, Alihan (2011) : Kooperativität vs. Rationalität, Discussion Papers, No. 8/2011, Universität Witten/Herdecke, Fakultät für Wirtschaftswissenschaft, Witten

This Version is available at:

<https://hdl.handle.net/10419/49947>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

discussion papers
Fakultät für Wirtschaftswissenschaft
Universität Witten/Herdecke

Neue Serie 2010 ff.
Nr. 8 / 2011
Kooperativität vs. Rationalität

Alihan Kabalack

discussion papers
Fakultät für Wirtschaftswissenschaft
Universität Witten/Herdecke
www.uni-wh.de/wirtschaft/discussion-papers

Adresse des Verfassers:

Alihan Kabalack
Max-Planck-Institut fuer Mathematik in den Naturwissenschaften
Inselstrasse 22 - 26
D-04103 Leipzig
e-mail: Kabalak@mis.mpg.de

Redakteure
für die Fakultät für Wirtschaftswissenschaft

Prof. Dr. Michèle Morner / Prof. Dr. Birger P. Priddat

Für den Inhalt der Papiere sind die jeweiligen Autoren verantwortlich.

Kooperativität vs. Rationalität

1. Zur Einleitung

Im Folgenden wird das Thema soziale Kooperation im Zusammenhang mit dem in den Gesellschaftswissenschaften allgegenwärtigen Thema Rationalität erörtert. Es hat sich als äußerst produktive Irritation erwiesen, dass es bislang merkwürdig schwergefallen ist, die intuitive Nähe von Kooperation bzw. Kooperativität, d.h. kooperativer Haltung, und Rationalität wissenschaftlich zu erfassen. Solange bestimmte Kooperationstypen nicht als rational ausgewiesen werden können, bleibt die Behandlung beider Konzepte lebendig. Man diskutiert, inwiefern eine unkooperative Rationalität wirklich in einem umfassenden Sinne ‚rational‘ sein kann. Daher werden vor allem jene Arten von allgemein vorteilhafter Kooperation auf ihre allgemeinen Merkmale hin überprüft, die bislang *nicht* als rational ausgewiesen werden können (Vgl. stellvertretend Binmore (2000) mit Gauthier (1982)). Auf diese Fälle werden auch wir uns beschränken. Das bestimmendes Merkmal einer jeden solchen Kooperation, der das gängige Rationalitätskonzept im Wege steht, ist dass mindestens einer der daran beteiligten Agenten einem Anreiz ausgesetzt ist, auf eine kooperative Handlung der anderen Agenten mit einer unkooperativen Handlung zum eigenen Vorteil und zum Nachteil dieser Anderen zu reagieren.

Was heute in den Wissenschaften, von der Ökonomik bis hin zur Informatik, meist unter ‚Rationalität‘ verstanden wird, ist ein Sprössling der ‚reinen Vernunft‘ oder – wenn es um rationale Handlungsentscheidungen geht – der ‚praktischen Vernunft‘ der Philosophie. Die grundsätzlichen Eigenschaften, die der Rationalität zugeschrieben werden, haben sich seit der Aufklärung kaum geändert, auch wenn sich das Vokabular verschoben hat (schon allein um den Unterschied von Wissenschaft und Philosophie zu markieren). Ohne Anspruch auf eine ‚kanonische‘ Formel, sei der Begriff ‚Rationalität‘ in dieser Arbeit so verwendet: Rational (praktisch vernünftig) ist ein Agent (Subjekt), der seine Urteile über seine und fremde Handlungen (Praxis) in Übereinstimmung mit allem fällt, was sich auf formal korrekte (logisch konsistente, folgerichtige) Art und Weise aus ihm gegebenen Informationen (seinem Wissen) ableiten lässt.

Da unsere Begriffsbestimmung etwas sperrig ausfällt, zwei Hinweise dazu, die im Folgenden eine Rolle spielen werden: 1. Wenn die Rationalität eines Agenten etwas rein privates oder ‚individuelles‘, d.h. etwas auf ihn allein bezogenes wäre, könnte man ‚Urteile‘ durch ‚Entscheidungen‘ ersetzen und ‚fremde Handlungen‘ streichen; tatsächlich ist Rationalität aber nicht individuell, sobald mehrere Agenten aufeinander treffen. 2. Ein Urteil, das in ‚Übereinstimmung mit‘ bestimmten Informationen, d.h. ohne Widerspruch zu ihnen, gefällt wird, muss nicht eindeutig aus diesen Informationen ableitbar sein; so bleibt Raum für Annahmen und andere Setzungen, die durch die Informationslage weder erzwungen noch ausgeschlossen sind.

Bleiben wir zunächst, ohne Punkt 1. zu vergessen, bei Entscheidungen hinsichtlich der eigenen Handlungen. Das Englische hält für einen Prozess rationaler Ableitung von Entscheidungen und anderen Urteilen aus gegebenen Informationen den Prozessbegriff '*reasoning*'¹ bereit, der zudem deutlich macht, dass ein solcher Prozess eine bei Bedarf explizit formulierbare *Begründung* (reason) für bestimmte Handlungen (bzw. Handlungsentscheidungen) und gegen alle nicht begründbaren und folglich irrationalen Alternativen liefert. Das reasoning zerlegt damit die Menge aller möglichen Handlungen (Alternativen) des Raisonierenden in zwei disjunkte Teilmengen, von denen sie eine bevorzugt: die Menge der begründbaren und damit rationalen Handlungen werden vor derjenigen der nicht begründbaren und damit irrationalen Handlungen besonders ausgezeichnet. Es ist bemerkenswert, dass damit nicht einzelne Handlungen aufgrund intrinsischer Eigenschaften das Prädikat ‚rational‘ oder ‚irrational‘ halten; sondern: jede Handlung wird *angesichts ihrer Beziehungen* zu all ihren Alternativen in eine der beiden Teilmengen oder ‚Klassen‘² eingeordnet.

Da sich die rationalen und irrationalen Teilmengen nicht überschneiden und da es kein drittes Abteil gibt, ergeben sich theoretische Probleme nicht im Verhältnis dieser beiden Klassen von Handlungen zueinander, sondern *innerhalb* der rationalen Klasse. Wenn das reasoning nämlich mehr als nur eine einzige Handlung als rational bestimmt, dann kann es keine rational begründete Entscheidung zwischen diesen gleichermaßen rationalen Handlungen geben. Das reasoning bleibt aus Sicht des Handelnden unvollendet: es schließt alle irrationalen Alternativen aus, unterscheidet aber nicht zwischen den restlichen (den rationalen) Alternativen. An diesem Punkt müssen die oben erwähnten Annahmen und Setzungen hinzukommen, um eine Handlungsentscheidung treffen zu können.

Ein berühmtes Sinnbild für das Entscheidungsproblem angesichts mehrerer rationaler Alternativen liefert ‚Buridans Esel‘: der Esel steht in der Mitte zwischen zwei identischen Heuballen. Um seinen Hunger zu stillen, müsste er sich für genau einen der Ballen entscheiden; denn in zwei Richtungen gleichzeitig kann er nicht laufen. Da die Heuballen aber gleichwertig sind, bleibt er in der Mitte stecken.³ Um nicht zu verhungern, muss der Esel aber eine Entscheidung für irgendeine Alternative innerhalb der rationalen Klasse treffen, also entweder zum ersten oder zum zweiten Ballen gehen – *dies* kann er rational aus seiner Präferenz ableiten, lieber nicht zu verhungern als zu verhungern.

¹ Der Einfachheit halber wird unterstellt, dass '*reasoning*' Rationalität impliziert; '*rationales reasoning*' wäre demnach eine Tautologie.

² Die mathematische ‚Menge‘ hat für tiefere Analysen zu wenig Struktur, da alle rationalen Handlungen durchaus eine Beziehung miteinander unterhalten, nämlich die Beziehung, gleichermaßen rational und damit Alternativen für einander zu sein. Das soll uns hier aber nicht weiter kümmern.

³ *Nicht*: weil der Esel von beiden gleichermaßen ‚angezogen‘ würde, sondern weil er keine rationale Überlegung anstellen kann, die zwischen den Ballen diskriminiert; es fehlt ein Antrieb.

M.a.W. muss er die dritte Option, in der Mitte stehen zu bleiben,⁴ rationalerweise ausschließen. Da er letztlich aber eine Entscheidung für eine bestimmte der beiden rationalen Alternativen und gegen die andere treffen muss, obwohl beide gleichermaßen rational wären, kann Rationalität bei diesem letzten Schritt nicht helfen. Die tragische Pointe der Geschichte ist, dass die Unfähigkeit zur Entscheidung innerhalb der rationalen Optionen de facto zur Vermeidung jeder dieser Optionen und stattdessen zur Realisation der irrationalen Option (Stehenbleiben) führt – es sei denn, der Esel wäre in der Lage, sich *rationalerweise willkürlich*⁵ auf eine der rationalen Optionen festzulegen.

2. Rationalität als selbstkonsistentes Urteilen

Um zu unserer Bestimmung von Rationalität zurückzukommen: Rationalität befähigt⁶ einen Agenten dazu, auf Basis beliebiger Informationen und Annahmen logisch konsistent zu denken und damit richtig zu urteilen (zu Schlussfolgern, korrekte Inferenzen anzufertigen)⁷. Stimmt das, ist Rationalität nicht nur für Handlungsentscheidungen relevant. Gemeinhin wird nämlich angenommen, dass das Denken eines Subjektes ohne diese Fähigkeit nicht begreifen könnte, was um es herum passiert. Dieser Punkt ist leicht zu akzeptieren, wenn man sich auf eine mittlerweile klassisch zu nennende Auffassung des subjektiven Denkens zu eigen macht: dass das Denken in irgendeiner Form die Realität ‚abbildet‘ oder ‚repräsentiert‘. Eine moderne, logisch-mathematisch ausgerichtete Abbildungstheorie des Mentalen vertritt z.B. (der junge) Wittgenstein (2010). Akzeptiert man eine Variante dieser Art Theorie, geht die Argumentation wie folgt: Nicht nur irgendwelche wahrnehmbaren Einzeldinge müssten abgebildet werden, sondern auch ihre *Beziehungen* zueinander, und auch die Beziehungen dieser Beziehungen usw. Solche Beziehungen (Relationen) aber lassen sich nicht direkt beobachten, sondern nur logisch ermitteln oder widerspruchsfrei postulieren. Ein aus der Philosophiegeschichte bekanntes Beispiel hierfür ist die Vorstellung einer physikalischen ‚Kausalität‘, die als Relation ‚zwischen‘ einer beobachtbaren Ursache und ihrer beobachtbaren Wirkung postuliert werden mag, aber selbst nicht beobachtet werden kann (u.a. dies leitet Kants Auseinandersetzung mit

⁴ Es wird angenommen, dass es die dritte Option gibt, in der Mitte stehen zu bleiben, um die Menge irrationaler Optionen nicht leer zu lassen. Wenn es diese Option nicht gibt, bleibt das Problem der Entscheidung zwischen den beiden rationalen Optionen aber unverändert erhalten.

⁵ Der Esel wird auch durch den modernen Probabilismus nicht gerettet, im Gegenteil. Selbst wenn der Esel in Wahrscheinlichkeitstheorie bewandert wäre, reichte es nicht, wenn er nach einer rationalen Wahrscheinlichkeitsverteilung für seine Handlungsalternativen suchte. Jede Verteilung (p : Ballen1, $1-p$: Ballen2) wäre genauso gut wie jede andere, denn *alle* Verteilungen mit $0 \leq p \leq 1$, die sich auf die beiden Ballen beschränken, haben denselben Erwartungswert. Nun müsste sich der Esel zwischen unendlich vielen gleichermaßen rationalen Verteilungen entscheiden, statt nur zwischen $p=1$ und $p=0$ zu entscheiden.

⁶ Wenn man dem Agenten wenig Freiheit lässt, dann kann man darüber hinaus auch annehmen, dass die Rationalität eines Agenten *sicherstellt*, dass er oder sie ausschließlich korrekte Inferenzen anstellt.

⁷ Im Folgenden werden urteilen, schlussfolgern und inferieren synonym verwendet.

Hume; vgl. Pearl (2010) für die moderne Kausalitätstheorie⁸). Darüber hinaus beziehen sich viele (wenn nicht alle⁹) Begriffe auf Relationen, obwohl sie wie Begriffe für Einzeldinge behandelt werden; darunter Begriffe wie ‚Gleichheit‘, ‚Alternative‘, ‚Implikation‘, ‚Unterschied‘, ‚Identität‘.

Generell wird daher dem rationalen Denken – nicht aber jeder beliebigen Art von Denken –, die Fähigkeit und die Aufgabe zugeschrieben, objektive Sachverhalte *mitsamt* ihrer Beziehungen zu erfassen. Im Idealfall entsteht so ein strukturell identisches (‚isomorphes‘) gedankliches (Ab-) Bild dieser Gesamtheit von objektiven Beziehung. Wenn, etwas vorsichtiger formuliert, nicht wenigstens eine strukturelle *Ähnlichkeit* (wenn schon keine Identität) zwischen

1. dem Urbild: der Ordnung der Dinge in der Welt (im Objektiven) und
2. seinem Bild: der Ordnung der Repräsentanten dieser Dinge (im Subjektiven)

herrsche, wäre das Bild ganz unbrauchbar; in der Folge wäre der zuständige Abbildungsprozess (Rationalität: objektive Urbild → subjektives Bild), d.h. das Denken selbst, unbrauchbar.

Eine ‚Parallelität‘ der Ordnungen von Welt und Vorstellung ist schon von Spinoza vertreten worden, und Abbildungstheorien haben insbesondere in der Bestimmung dessen Schwierigkeiten, was genau die ‚Objekte‘ im Bild (die Korrelate objektiver Sachverhalte in der Vorstellung, ‚Semantik‘) darstellen. Doch auch unabhängig davon, wie man zu solchen Abbildungstheorien steht, kann man die These akzeptieren, dass ein unsystematisches (nicht-rationales) subjektives Denken immer nur in zufälliger Beziehung zu den objektiven Sachverhalten stehen muss, falls diese Sachverhalte selbst eine Systematik aufweisen (wovon auszugehen ist). In der Folge stünden auch die Gedanken der Mitglieder eines Kollektivs von Subjekten in nur zufälliger Beziehung zueinander, wenn ihr Denken nicht dieselbe Rationalität aufweist. Eine allen Subjekten gemeinsame, einheitliche und einzigartige Rationalität würde hingegen sicherstellen, dass verschiedene Subjekte ihre Wahrnehmungen derselben Sachverhalte in derselben rationalen Weise verarbeiten, sodass sie zu denselben Schlussfolgerungen kommen. D.h. die Abstimmung des Denkens jedes einzelnen Subjekts mit den Sachverhalten in der Welt würde zu einer Abstimmung des Denkens verschiedener Subjekte in der derselben Welt führen, sodass sich diese Subjekte *über ihre gemeinsame Welt* verständigen könnten.

⁸ Die Problematik der Unbeobachtbarkeit einer kausalen Relation äußert sich in der gegenwärtigen Diskussion u.a. in der Frage des ‚Kontrafaktischen‘: die Kausalitätstheorie erörtert, inwiefern sich eine kausale Beziehung zwischen zwei tatsächlich gegebenen Phänomenen dadurch feststellen lässt, dass man die tatsächliche Situation mit nicht tatsächlichen (kontrafaktischen) Alternativsituationen vergleicht. Es ist unklar, ob und was man rational über Situationen aussagen kann, die *nicht* eintreten (vgl. Pearl 2010).

⁹ Es ist überhaupt fraglich, ob ein echter Unterschied zwischen Relationen und ‚Einzeldingen‘ durchgehalten werden kann. Man kann ein ‚Einzelding‘ oder ‚Objekt‘ O vollständig durch eine einzigartige, ihm eindeutig zuzuordnende Identitätsrelation $Id_O: O \mapsto O$ ersetzen. Vgl. z.B. Pumplün (1999) für die mathematische Kategorientheorie; dort wird versucht, die spezifischen Eigenschaften eines solchen O ausschließlich über dessen Relation zu anderen Objekten (mithin über Relationen von Identitätsrelationen) zu identifizieren.

Als man die Autonomie des subjektiven Denkens für gesichert hielt, man ihm also u.a. die Freiheit zuschrieb, auch irrational denken zu können, ließ sich hieraus leicht die Lehre ziehen: man muss sich anstrengen, rational zu denken, wenn man die Welt verstehen und von Anderen verstanden werden will. Es ist bemerkenswert, dass dies selbst schon ein rationales (praktisches) Urteil ist: wenn Rationalität der einzige Weg zum Verständnis der Welt und zur Verständigung mit Anderen ist, und *wenn* man dies anstrebt, dann ist es nur folgerichtig, Rationalität anzustreben. Es ist m.a.W. rational, rational zu sein, sofern man überhaupt etwas verstehen will; und generell ist es rational, ein rationales Urteil zu fällen. Diese *Reflexivität* des Rationalen – ein in jedem rationalen Urteil implizierter Selbstbezug, der eine widerspruchsfreie Selbstanwendbarkeit garantiert, – ist ein starker Hinweis darauf, dass das Phänomen nicht aus der Luft gegriffen ist. Diese auch *Selbstkonsistenz* genannte Eigenschaft der Rationalität ist Nachhall und Verfeinerung der berühmten cartesischen Einsicht, dass ‚das Denken‘ jenes Phänomen sei, das nicht denken kann, dass es nicht denkt (und also immer denken muss, dass es denkt)¹⁰.

Spezielle Aspekte der generellen Reflexivitätseigenschaft bzw. Selbstkonsistenz rationaler Urteile ziehen sich bis heute als roter Faden durch diverse Rationalitätskonzepte, sind also nicht nur philosophie- oder theoriegeschichtlich interessant. Ging es Cartesius noch darum, sich durch Erkenntnis der notwendigen Reflexivität seines Denkens seiner eigenen Existenz zu versichern, stand später (z.B. bei Kant unter dem Titel ‚praktische Vernunft‘) die Anwendbarkeit rationalen Inferierens auf Handlungsoptionen im Vordergrund. Auch eine Lesart des bekannten ‚kategorischen Imperativs‘ lässt sich als Selbstkonsistenzforderung und in diesem Sinne als Rationalitätsforderung interpretieren, obwohl der ‚Imperativ‘ vordergründig rein normativ auftritt: man soll nur solche Handlungen realisieren, die bei gleicher Sachlage auch von jedem anderen Agenten jederzeit Zeit realisiert werden könnten; das sind m.a.W. Handlungen, die in der Gesellschaft des Agenten zur Regel werden können, und die unter der Bedingung ihres regelmäßigen Auftretens immer noch eine ‚gute‘ Wahl wären. – Hiermit ist noch lange kein kooperatives Prinzip begründet. Es wird lediglich gefordert, dass ein Agent, der meint, dass eine bestimmte Handlung *h* rational sei, im Zuge dieses Urteils berücksichtigen muss, dass auch alle rationalen Anderen regelmäßig zu eben diesem Schluss kommen und kommen werden.

Im Handlungskontext geht es nicht nur darum, aus bekannten Wahrheiten auf damit notwendig einhergehende weitere Wahrheiten zu schließen, sondern auch darum, aus bekannten Wirkungen seiner Handlungsoptionen darauf zu schließen, was man tun müsste, um dieses oder jenes *wahr zu machen*. Freilich muss man schon festgelegt haben, *welche* der handelnd realisierbaren Wirkungen man wahrmachen möchte, bevor man sich an die rationale Lösung

¹⁰ Und in diesem Sinne: jenes Phänomen also, das dann und nur dann existiert (‚ist‘), wenn es denkt, dass es denkt (*sum cogitans*). Vgl. dies mit der FN oben zur Identifikation von Objekten mit ihren spezifischen Identitätsrelationen in der Kategorientheorie.

des so gestellten Entscheidungsproblems macht. Diese Festlegung wird in zeitgenössischen Modellen durch eine *Präferenzordnung* pro Agent besorgt, die der betroffene Agent dann als Information in seine rationalen Schlussfolgerungen einspeisen kann (s.u. mehr).

Auf dieser Linie spielt die Reflexivitätseigenschaft von Rationalität beim Zusammenhang der rationalen Urteile bzw. Schlussfolgerungen *mehrerer* aufeinander treffender handlungsfähiger Subjekte eine wesentliche Rolle. Nur solche Urteile sind rational, die mit sich selbst im Einklang stehen, sich also selbst für allgemeingültig und für allgemein ermittelbar erklären, und zwar egal wer sie gerade fällt. Daher sind sich allen gängigen Vorstellungen von Rationalität bei folgendem Prinzip einig: *Ein Urteil, eine Schlussfolgerung, die sich selbst widerspricht, sich selbst außer Kraft setzt oder sich sonstwie selbst leugnet, kann nicht rational sein.* Dies benennt eine (objektive) Eigenschaft aller rationalen Schlussfolgerungen, die von den sie anfertigenden Agenten unabhängig ist. Insbesondere muss daher ein Agent, der aus seinen Information und Annahmen den rationalen Schluss ϕ zieht, auch davon ausgehen, dass *jeder andere rationale Agent* bei gleicher Informationslage und unter denselben Annahmen zum selben Schluss ϕ kommen muss. Ein rationaler Agent darf m.a.W. nicht sowohl ϕ schlussfolgern als auch glauben (bzw. schlussfolgern), dass er allein dazu in der Lage ist, aus seinen Informationen und Annahmen ϕ zu inferieren. Dies hilft zum einen, irrationale Schlussfolgerungen auszusortieren; zum anderen ist damit klar, dass ein rationaler Agent sich darauf verlassen kann, dass alle anderen rationalen Agenten, die dieselben Informationen und Annahmen teilen, ebenfalls zur Gewissheit ϕ gelangen. Dies liefert das Fundament für die Möglichkeit einer ‚rationalen Interaktion‘ von mehreren Agenten; eine der aktuellen Fragen ist, ob dies auch für Kooperationen reicht.

Problematisch bleibt die Abhängigkeit folgerichtigen Denkens vom gegebenen ‚Material‘, d.h. den Informationen aus denen inferiert werden soll, was sonst noch alles gegeben, möglich und unmöglich ist. Ein Aspekt hiervon ist die Möglichkeit des Irrtums hinsichtlich des ‚Gegebenen‘: Eine korrekt angefertigte Schlussfolgerung aus einer falschen Information¹¹ führt wieder zu einer falschen Information; in gewissem Sinne mehrt sich so der Irrtum, während die korrekte Anwendung der rationalen Methode das Subjekt in falsche Sicherheit wiegen kann. Dass der Erfolg von Rationalität auf gegebene Informationen angewiesen ist, wiegt besonders schwer, wenn man ohne Abstufungen nur richtige von falschen Informationen unterscheidet: Soweit sich die Logik traditionell nicht um abgestufte Wahrheiten, sondern nur um die Opposition ‚wahr‘ (T, für weitere Inferenzen verwendbar) vs. ‚falsch‘ ($\perp$ = Nicht-T, nicht weiter verwendbar) kümmert, kann sie aus *unsicheren* Informationen nicht direkt auf etwas schließen. Da unsichere Informationen – sprich: Informationen über bloße *Wahrscheinlichkeiten* p ($0 < p < 1$) – weder gegebenen Wahrheiten (T=nicht- $\perp$)

¹¹ Falls man bei einem Irrtum, einer Fehlinformation überhaupt von einer falschen ‚Information‘ sprechen kann.

noch sicher falsch ($\perp$ = nicht-T) sind, lassen sie sich nicht in solch ein binäres System einordnen (vgl. zur Problematik z.B. Ramsey 1930).¹²

Versuche, dieses Problem zu überwinden, riefen schließlich die Verschmelzung von Mathematik und (formaler) Logik in der Wahrscheinlichkeitstheorie und Statistik auf den Plan (z.B. in Form von Carnaps (1952) ‚induktiver Logik‘). Dieser Pfad führt zur modernen Konzeption ‚individueller Rationalität‘ in der sog. Entscheidungstheorie (Savage 1954, Neumann / Morgenstern 1968), die nicht nur aus sicher gegebenen Informationen, sondern auch aus bloßen Wahrscheinlichkeitsangaben über alternative (einander ausschließende) Sachverhalte rationale Schlüsse ziehen kann. Natürlich bleiben solche Schlüsse unter dem Vorbehalt, dass sie von den zugrunde gelegten Wahrscheinlichkeitsangaben abhängen (auch qualitativ¹³).

Ein anderer problematischer Aspekt betrifft nur ‚praktische Urteile‘: Erst wenn die Rationalität auf gegebene subjektive Präferenzen (Bewertungen) angewandt wird, führt sie zu Schlussfolgerungen darüber, was das Subjekt *tun müsste und daher sollte*, um diese Präferenzen zu befriedigen. Hinsichtlich seiner eigenen Präferenzen kann ein Subjekt (Ego) selbst zwar nicht irren,¹⁴ aber *andere* Subjekte können sich hinsichtlich Egos Präferenzen irren. Und falls die Präferenzen eines Subjekts unabhängig davon festgelegt sind, was in der Welt sonst noch der Fall ist, können sie *nicht* sowohl auf rationale Weise als auch eindeutig aus Weltbeobachtungen abgeleitet werden. Folglich birgt die tatsächliche Präferenz eines Agenten für andere Agenten immer eine gewisse Unsicherheit, die sich nicht durch rationale Schlussfolgerungen auflösen lässt.

3. Präferenzen, Erwartungsnutzen

Im Folgenden werden einige konzeptionelle Entwicklungen der Forschung zur *individuellen* Rationalität in der oben erwähnten *Entscheidungstheorie* kurz beschrieben (hier im Wesentlichen: Neumann / Morgenstern 1961, Savage 1954, vgl. z.B. Schmidt 1995 für eine knappe Einführung). Es gibt gute Gründe, sich

¹² Daher muss ein binäres System unsicheres wie falsches behandeln, dann alles was nicht wahr ist, ist dort per Definition dem Wert ‚falsch‘ zuzuordnen. Da es dann nicht für Schlussfolgerungen herangezogen werden kann, wird unsicheres letztlich genauso wie völlig unbekanntes behandelt. Das Problem besteht darin, ob und wie man einen durchaus sicheren Schluss ziehen kann, der aus etwas unsicherem etwas anderes unsicheres ableitet.

¹³ Die gängigsten konkurrierenden Interpretationen sind grob: i) Eine rational verarbeitbare Wahrscheinlichkeitsangabe muss einer objektiven Wahrscheinlichkeit entsprechen; dann lassen sich daraus objektive Wahrheiten über inhärent unsichere (nicht-determinierte) Sachverhalte ableiten. ii) Eine Wahrscheinlichkeit gibt die Auftrittshäufigkeit (Frequenz) eines Phänomens unter bestimmten Bedingungen an, das zwar determiniert auftritt, dessen Auftritt aber durch einen uns im Detail unbekannten Mechanismus bestimmt wird; dann lassen sich hieraus Schätzung für zukünftige Frequenzen ableiten, sofern der Mechanismus erhalten bleibt. iii) Eine Wahrscheinlichkeitsangabe ist nur ein Ausdruck für die subjektive Gewissheit bzw. Unsicherheit eines Agenten über objektive Sachverhalte; dann lässt sich aus der subjektiven Wahrscheinlichkeit nur etwas über die damit logisch vereinbaren Gewissheiten und Unsicherheiten des Agenten ableiten (aber nichts über die Welt).

¹⁴ ‚Präferenzen‘ in einem strengen Sinne (reflexive, transitive Präferenzordnung), der die Anwendbarkeit von Rationalität schon sicherstellt und daher auch voraussetzt, dass die Präferenz ihrem Wirt bekannt ist. Daher kann man z.B. nicht einwenden: „Was, wenn sich ein Subjekt doch hinsichtlich seiner eigenen Präferenzen irrt?“ – Dann hätte es überhaupt keine Präferenz.

dieser Entwicklungen zu vergewissern, bevor es um die Interaktion mehrere Agenten geht. In diesem Bereich wurden konzeptionelle Traditionen begründet, an die die theoretischen Behandlungen von Rationalität im sozialen Kontext angeschlossen haben. Deren Kenntnis hilft daher dabei, die Schwierigkeiten einzuordnen, denen theoretische Versuche begegnen, Rationalität im Sozialen mit Kooperationen zu vereinbaren.

Präferenzordnung. Eine Präferenzordnung $\Pi <O, \succsim>$ über einer Menge von *outcomes* $O = \{o_1, o_2, \dots, o_r\}$ ist eine Menge von (reflexiven, transitiven)¹⁵ gerichteten zweiwertigen Relationen $\succsim$ zwischen Elementen von O . Zwischen je zwei *outcomes* o_x, o_y besteht mindestens eine der Relationen $o_x \succsim o_y$ oder $o_y \succsim o_x$. Diese *outcomes* bestehen aus Sachverhalten, die ein Agent handelnd herbeiführen kann, d.h. aus realisierbaren *Handlungsergebnissen*. Die Präferenzordnung Π eines Agenten legt fest, welche *relativen* Werte er den *outcomes* zuordnet (welche *outcomes* er also welchen vorzieht); solange insbesondere die Transitivitätsbedingung gewahrt bleibt (die für ‚Ordnung‘ sorgt), kann jeder rationale Agent jede beliebige Präferenz unterhalten. Eine jede in Π festgelegte Relation $o_1 \succsim o_2$ zwischen zwei *outcomes* o_1 und o_2 kann als „ o_1 wird o_2 nicht vorgezogen“ oder als „ o_1 ist nicht besser als o_2 “ gelesen werden.¹⁶ Gleichwertig sind zwei *outcomes* ($o_1 \sim o_2$) dann, wenn die Relation in beide Richtungen besteht, wenn also sowohl $o_1 \succsim o_2$ als auch $o_2 \succsim o_1$. Ein *outcome* o_2 ist besser als ein anderes o_1 (d.h. $o_1 \prec o_2$), wenn die Relation $o_1 \succsim o_2$ besteht, nicht aber die Relation $o_2 \succsim o_1$.

Nutzenfunktion. Solch eine rein qualitative Präferenzordnung Π (‚ordinale Präferenzen‘) wird in eine quantitative *Nutzenfunktion* U (‚kardinale Präferenzen‘) überführt, wenn die einzelnen *outcomes* in O zusätzlich Gewichte erhalten, d.h. absolute Maßzahlen für ihre Güte oder Wertigkeit. Jedes *outcome* o_x erhält dann einen absoluten Nutzwert $U(o_x)$, der in einer beliebigen (und daher interpretationsbedürftigen) Einheit ‚utils‘ angegeben werden kann.¹⁷ Durch diese Quantifizierung des ‚Wertes‘ eines Ergebnisses lassen sich nicht nur qualitative Relationen der Art „ist besser als“ oder „ist nicht besser als“ feststellen, sondern auch Abstände zwischen zwei Ergebnissen, $d(o_x, o_y) = U(o_x) - U(o_y)$, die sich zudem mit den Abständen anderer Paare von *outcomes* vergleichen und verrechnen (z.B. mitteln) lassen.

Welt als Mechanismus. Was hier ‚outcomes‘ heißt, wird mitunter ‚Alternativen‘ genannt; allerdings könnte dies zu Missverständnissen führen, da ein Agent i , der in einer Modellwelt handelt, nicht direkt eines der alternativen Ergebnis wählt, sondern eine seiner alternativen Handlungsoptionen $H^i = \{h_1^i, h_2^i, \dots, h_m^i\}$. Seine ‚Welt‘ sorgt dann dafür, dass jede seiner Handlungsoptionen in H^i zu einem Ergebnis in O führt. So kann die Welt W als *Mechanismus* interpretiert werden, der zwischen Handlungen und Ergebnissen

¹⁵ Angenommen, dass die gerichtete Relation R zwischen den *outcomes* o_x und o_y besteht, wird notiert als $o_x R o_y$. 1. Reflexivität: für alle o_x gilt: $o_x R o_x$. 2. Transitivität: wenn $o_x R o_y$ und $o_y R o_z$, dann auch $o_x R o_z$.

¹⁶ Alternativ kann $\sim$ als Relation dienen.

¹⁷ Solange immer wenn $o_x \succsim o_y$ gilt, auch $U(o_x) \geq U(o_y)$ gilt, ist Π mit U vereinbar.

vermittelt. M.a.W. bildet der Mechanismus W die Menge der Optionen H^i des Agenten i auf die möglichen Ergebnisse O ab: $W: H^i \mapsto O$.

Nutzen einer Handlung. Wenn nun eine Handlung h_a sicher zu einem bestimmten Ergebnis o_x führt, d.h. wenn $W(h_a)=o_x$, dann kann der ursprünglich dem Handlungsergebnis o_x zugeordnete Nutzwertwert an die zugehörige Handlungsoption h_a weitergereicht werden: $U(h_a) = U(o_x)$. Das Ergebnis (der Zweck) o_x ‚vererbt‘ der Handlung h_a (dem Mittel) seinen Wert, weil h_a sicher zu o_x führt. Diese Ausweitung des Konzepts ‚Nutzen U ‘ auf Handlungsoptionen setzt allerdings schon Rationalität voraus; denn die Übertragung der Nutzwerte von outcomes zu Handlungen verlangt vom Agenten i , einen rationalen Schluss aus seiner Nutzenfunktion U und der ihm bekannten Funktionsweise seiner Welt W zu ziehen. Hierauf gründet Webers (1990) frühe Kennzeichnung von ‚Zweckrationalität‘, wonach die Handlung und die Absicht eines rationalen Agenten in folgendem Sinne ‚übereinstimmen‘ müßten: Wenn Agent i weiß, dass eine bestimmte Handlungsoption h_a^* zu jenem Ergebnis o_x^* führt, das ihm in W den maximal möglichen Nutzen beschert, dann ist es für ihn zweckrational, eben diese beste mögliche Handlung h_a^* zu realisieren.¹⁸ In der Folge wird diese Handlung h_a^* mitunter als ‚rationale Handlung‘ bezeichnet, obwohl genau genommen die *Schlussfolgerung* rational ist, die von U und W zu h_a^* führt – aber nur *sofern* zusätzlich angenommen werden darf, die allgemeine (vor-rational festgelegte) Absicht des Agenten darin besteht, durch seine Handlungswahl einen möglichst hohen Nutzwert zu erzielen.

Erwartungsnutzen. Wichtig für den nächsten Schritt der Entscheidungstheorie ist, dass sich der Nutzwert $U(o_x)$ eines möglichen Ergebnisses o_x einer Handlung h_a auch mit einer Wahrscheinlichkeit $p_{x,a}$ seines tatsächlichen Eintritts wichten lässt. Voraussetzung ist, dass der Mechanismus W dafür sorgt, dass h_a nicht sicher, sondern eben nur mit Wahrscheinlichkeit $p_{x,a}$ zu o_x führt (und mit $1 - p_{x,a}$ zu anderen Ergebnissen). Ist dies dem Agenten bekannt, kann er den ‚unsicheren‘ bzw. den im Durchschnitt ‚zu erwartenden‘ Wert $p_{x,a} U(o_x)$ in rationale Schlussfolgerungen einbinden.

Dies führt zu dem erst im Laufe des 20. Jhs (Neumann / Morgenstern 1961) eingeführten ‚Erwartungsnutzen‘ $EU(h_a)$ einer Handlung h_a , die mit bestimmten Wahrscheinlichkeiten zu verschiedenen Ergebnissen führt. Der *Erwartungsnutzen einer Handlung* h_a , $EU(h_a)$, ist der mit Eintrittswahrscheinlichkeiten gewichtete Mittelwert der Nutzwerte aller

Ergebnisse $o_{x=1}, o_{x=2}, \text{etc.}$, zu denen h_a führen kann: $EU(h_a) = \sum_x p_{x,a} U(o_x)$, wobei $\sum_x p_{x,a} = 1$. Um den Nutzwert der Handlungsoption h_a voll zu erfassen, werden natürlich die Wahrscheinlichkeiten ausnahmslos all ihrer möglichen Ergebnisse benötigt, die sich zu 1 aufaddieren; d.h. man benötigt die

¹⁸ Es gilt also $U(h^*) = U(o^*) = \text{Max}_h U(h) = \text{Max}_o U(o)$.

Wahrscheinlichkeitsverteilung $p_a = (p_{x=1,a}, p_{x=2,a}, \dots)$ aller möglichen Ergebnisse von h_a .¹⁹

Angenommen eine Handlung h_a führt mit Wahrscheinlichkeit $p_{x,a}$ zum Ergebnis o_x und andernfalls mit $p_{y,a}=1-p_x$ zu einem anderen Ergebnis o_y . Dann lässt sich h_a auch so behandeln, *als würde* h_a eindeutig zu einem dritten Ergebnis o_z führen, welches einen durchschnittlichen Nutzen von $U(o_z) = p_{x,a} U(o_x) + p_{y,a} U(o_y) = EU(h_a)$ hat.²⁰ Auf diese Weise lässt sich eine Situation, in der eine Handlung h_a unsichere Konsequenzen (o_x, o_y) mit jeweils sicheren Nutzwerten $(U(o_x), U(o_y))$ hat, rational umformulieren: zu einer Situation, in der h_a eine sichere Konsequenz o_z hat, welcher aber nur ein unsicherer, d.h. durchschnittlicher, Nutzwert $U(o_z)$ zugeordnet ist. In beiden Interpretationen hat die Handlung h den Erwartungsnutzen $EU(h_a) = U(o_z)$, sodass sie in dieser Hinsicht ‚äquivalent‘ sind. Die Feststellung und Nutzbarmachung dieser Äquivalenz erfordert Rationalität.

Gemischte Handlung. Ein zweiter Aspekt des probabilistischen Zugangs ist, dass die Verrechenbarkeit von Eintrittswahrscheinlichkeiten mit absoluten Nutzwerten letztlich auch die Möglichkeit eröffnet, sich für eine ‚Mischung‘ verschiedener Handlungen zu entscheiden. Anstatt sich auf eine einzige Handlung festzulegen, kann ein Agent beschließen, verschiedene Handlungsoptionen mit verschiedenen Wahrscheinlichkeiten (die sich zu 1 addieren) zu realisieren. Wenn z.B. eine Handlungsoption h_a – entweder sicher oder, wie eben, rechnerisch – zu einem Ergebnis o_x führt und eine andere Handlung h_b zu einem anderen Ergebnis o_y , dann kann der Agent sich zu einem *gemischten Handlungsplan* h_c entscheiden, der diese Handlungen kombiniert.

Solch ein Handlungsplan h_c kann z.B. vorsehen, dass der Agent mit einer Wahrscheinlichkeit q_a Handlung h_a realisiert und alternativ mit der Wahrscheinlichkeit $q_b = 1 - q_a$ die andere Handlung h_b . Dann entspricht h_c einer planmäßigen Randomisierung von zwei ‚echten‘ oder ‚reinen‘ Handlungen, d.h. $h_c = q_a h_a + q_b h_b$. Somit lässt sich der Plan h_c rational so behandeln, *als wäre* er eine eigenständige Handlungsoption h_c , die eindeutig zu einem dritten Ergebnis o_z führt,²¹ welches einen mittleren Nutzen von $U(o_z) = q_x U(o_x) + q_y U(o_y)$ hat. Obwohl dieser Wert dann der ‚gemischten Handlung‘ h_c als der ihr spezifische Erwartungsnutzen zugeordnet werden kann, $EU(h_c) = U(o_z)$, ist zu beachten, dass der betroffene Agent in einer konkreten Einzelsituation durchaus eine der reinen Handlungen h_a oder h_b realisieren muss (was nicht unter allen Umständen unproblematisch ist). Um in einer konkreten Situation einen gemischten Plan h_c zu verfolgen, muss es der Agent irgendwie bewerkstelligen, seine in dieser Situation tatsächliche Handlung h_a oder h_b *zufällig*, aber unter Beachtung der Gewichte q_x und q_y auszuwählen. Sichtbar wird ein solches randomisiertes

¹⁹ Entweder wird p_a auf alle outcomes beschränkt, zu denen h_a mit positiver Wahrscheinlichkeit führt, oder die Verteilung wird für alle möglichen outcomes in O festgelegt, wobei keinem Ergebnis eine Wahrscheinlichkeit unter null zugeordnet werden darf.

²⁰ Hierin hallt die ursprüngliche Beschränkung der Logik auf zwei Wahrheitswerte nach.

²¹ $W(h_c) = o_z$

Vorgehen aber erst, wenn derselbe Agent in verschiedenen Instanzen desselben Situationstyps zwischen h_a oder h_b changiert.

Um festzustellen, dass der ‚ontische‘ Unterschied zwischen gemischten und reinen Handlungen bzw. Handlungsplänen hinsichtlich der Maximierung seines Erwartungsnutzens (unter günstigen Bedingungen) irrelevant ist, muss ein Agent rational sein. Stellt ein rationaler Agent i dies fest, weiß er, dass seine Handlungsoptionen H^i neben reinen Handlungsoptionen auch alle möglichen Mischungen dieser Optionen umfassen. Dies ist ein Beispiel, dafür was oben im zweiten Abschnitt allgemein beschrieben wurde. Ein Agent muss rational sein, um die *strukturellen Zusammenhänge* seiner Welt zu erschließen, die in den nicht direkt beobachtbaren, aber logisch konsistent inferierbaren Beziehungen der beobachtbaren Sachverhalte bestehen: dass Agent i neben seinen reinen Handlungsoptionen auch eine gemischte Option wie h_c zur Verfügung hat, kann er nur durch die rationale Ermittlung von $h_c = q_a h_a + q_b h_b$ und des jeweiligen Erwartungsnutzens $EU(h_c)$ erschließen. Dasselbe gilt für einen eventuellen Beobachter von i : nur ein rationaler Beobachter kann, wenn er i in verschiedenen Instanzen desselben Situationstyps beobachtet, schlussfolgern, dass i möglicherweise einem ganz bestimmten gemischten Plan h_c folgt, obwohl er in derselben Situation mal h_a und mal h_b realisiert.

Weltzustände. Eine weitere für unser Thema relevante Komplikation der Entscheidungstheorie betrifft die Möglichkeit, die Modellwelt W in sich zu differenzieren, in der sich ein Agent bewegt. Gemeinhin wird dazu eine Menge von ‚möglichen Weltzuständen‘ (*states*) oder Situationen $S = \{s_1, s_2, \dots, s_w\}$ angenommen. Dabei bleibt meist offen, was genau ein solcher ‚Zustand‘ sein soll – mögliche Zustände S einer Welt W oder viele alternative *mögliche Welten* S oder etwas ganz anderes. Fest steht, dass dies die Darstellung einer Modellwelt erlaubt, in der es vom jeweiligen Wert von S abhängt, welche möglichen Ergebnisse in O mit welchen Wahrscheinlichkeitsverteilungen den einzelnen Handlungsoptionen in H zuzuordnen sind. Anstelle einer direkten Vermittlung von Handlungen und Ergebnissen $W: H^i \mapsto O$ tritt ein Mechanismus, der alle möglichen Paare aus Handlungsoption und Weltzustand in jeweils ein Ergebnis übersetzt. Insgesamt überführt der Mechanismus W also das kartesische Produkt der Optionen und Zustände in outcomes, d.h. $W: H^i \times S \mapsto O$. Soweit dann einer bestimmten Handlungsoption h_a bei verschiedenen Weltzuständen s_u, s_v , etc. verschiedene Verteilungen $p_{a,u}, p_{a,v}$ etc. hinsichtlich ihrer Ergebnisse zugeordnet sind, hängt der Erwartungsnutzen dieser Handlung auch vom Weltzustand ab: $EU(h_a, s_u), EU(h_a, s_v)$, etc. Ist der Agent i hierüber informiert, kann er rational ermitteln, welche seiner Handlungsoptionen in den verschiedenen möglichen Weltzuständen den jeweils maximalen Erwartungsnutzen haben.

Reaktionsfunktion. Letztlich muss ein Agent i mit feststehender Nutzenfunktion U^i also denselben Handlungsoptionen verschiedene Erwartungsnutzenwerte zuordnen, nämlich je Weltzustand einen. Als Ergebnis dieser Überlegungen ermittelt der rationale Agent i eine (erwartungs-)

nutzenmaximierende *Reaktionsfunktion* RF^{i*} , die für jeden einzelnen möglichen Weltzustand festlegt, jeweils welche der (reinen oder gemischten) Handlungsoptionen aus H^i den maximalen Erwartungsnutzen birgt. Damit ist RF^{i*} eine Abbildung $RF^{i*}: S \mapsto H^i$ mit der besonderen Eigenschaft, dass für alle $s \in S$ gilt, dass $RF^{i*}(s) = \operatorname{argmax}_{h \in H^i} EU^i(h, s) = h^{i*,s}$.

Diese Art Reaktionsfunktion wird deswegen ‚rational‘ genannt (‚Bayesianische Rationalität‘), weil ein Agent i aus

1. seiner über alle outcomes O definierten Nutzenfunktion U^i ,
2. seinen Handlungsoptionen H^i ,
3. seiner Welt $W: H^i \times S \mapsto O$,
4. und seiner vor-rational festgelegten Absicht, seinen Erwartungsnutzen in *jedem einzelnen* der möglichen Weltzustände zu maximieren

genau diese Reaktionsfunktion RF^{i*} ableiten muss, sofern seine Ableitung logisch konsistent erfolgt. Um auf Grundlage dieser RF^{i*} konkret handeln zu können, muss der Agent i natürlich zumindest zu wissen glauben, welcher der möglichen Weltzustände aktuell herrscht.

Um es zu betonen: An der Ermittlung und Festlegung einer Reaktionsfunktion mit der beschriebenen besonderen Eigenschaft ist nicht die ‚Nutzenmaximierung‘ an sich rational, sondern die Art und Weise wie die pro Situation festgelegten Handlungen ermittelt werden, *wenn* in den Prämissen der rationalen Inferenz schon festgelegt ist, dass der Agent seinen Erwartungsnutzen in jeder einzelnen Situation maximiert *will*. Wenn der Agent etwas anderes wollte, müsste er rationalerweise einer anderen Reaktionsfunktion ermitteln und befolgen.

Gesetzt, Agent i will in jeder der möglichen Weltzustände S seinen Erwartungsnutzen maximieren. Dann kann eine solche Reaktionsfunktion $RF^{i*}: S \mapsto H^i$ als *lokal maximierende* Reaktionsfunktion betrachtet werden; denn die Reaktion $h^{i*,s} = RF^{i*}(s)$, die diese rationale Reaktionsvorschrift jedem einzelnen Zustand $s \in S$ zuordnet, ist im Einzelfall unabhängig davon, welche Reaktionen sie allen anderen Zuständen zuordnet. Anders formuliert, behandelt der Agent die verschiedenen Weltzustände s_u, s_v etc. wie *verschiedene Welten*, deren Zusammenhang (falls es einen gibt) für seine jeweiligen Handlungen h^{i*,s_u}, h^{i*,s_v} etc. in diesen Welten unerheblich ist. Sobald der Agent glaubt, ermittelt zu haben, welchen Zustand die Welt gerade hat, kann er die anderen Zustände ignorieren.

5. Interaktionen ohne Kooperationsmöglichkeit

Die eben beschriebene Konzentration auf lokale Maximierung ist berechtigt (eben: rational), solange ein einsamer Agent i innerhalb einer ihm vorgegebenen Welt W mit vorgegebenen Weltzustand s_u handelt, den er in keiner Weise durch

seine Wahl beeinflussen kann. Die aktuelle Handlungswahl des Agenten – egal, ob sie rational erfolgt oder nicht – steht in keiner systematischen Beziehung dazu, welcher Weltzustand ihm aktuell vorgegeben ist;²² folglich kann und muss ein Urteil des Agenten darüber, wie er rationalerweise handeln soll, die aktuelle Welt als gegeben behandeln.

Daher ist hier die Redeweise „*innerhalb* einer Welt handeln“ bzw. „*in* einer Welt im Zustand s_i handeln“ bedeutungsvoll, weil sie den hierarchischen Unterschied zwischen dem einzelnen Agenten und der ihn *umfassenden* Welt explizit macht. Jeder mögliche Weltzustand s_i gibt dem Agenten i vor, auf welche Weisen er handeln kann²³ und welche Ergebnisse diese Optionen mit welchen Wahrscheinlichkeiten zeitigen werden, sobald sie realisiert werden. Dem Agenten bleibt nichts anderes übrig, als sich zu beugen und eine der ihm aktuell durch seine Welt vorgegebenen Optionen auszuwählen und zu realisieren. Dieser Agent ‚reagiert‘ also zwar auf den aktuellen Zustand der Welt im beschriebenen Sinne; aber diese ‚Reaktion‘ besteht in einem *Agieren innerhalb* eben dieser Welt und nicht in einem *Interagieren mit* der Welt. Die Welt selbst ‚reagiert‘ nämlich nicht etwa auf die Handlungen des Agenten mit eigenen ‚Handlungen‘; sondern jede mögliche Handlung des Agenten ist ebenso Teil der Welt wie auch das durch die Welt festgelegte Ergebnis dieser Handlung.

Diese etwas spitzfindige Unterscheidung von Agent und Welt ist aufschlussreich, wenn es nun nicht um einen einsamen Agenten in seiner Welt geht, sondern um das Zusammentreffen mehrerer Agenten in derselben Welt. Die wesentlichen Grundlagen für diesen Übergang von rationaler Aktion (‚Entscheidungstheorie‘) zu rationaler Interaktion (‚Spieltheorie‘) wurden in zwei Schritten von Neumann und Morgenstern (1961)²⁴ und Nash (1951) gelegt. Anstelle einer Welt, die einem einzelnen Agenten verschiedene Handlungsoptionen und damit verbundene Handlungsergebnisse bereithält, tritt nun ein *Spiel*: eine Welt die mehreren Agenten Interaktionsoptionen und damit verbundene Interaktionsergebnisse bereitstellt. Eine einzelne Handlung allein führt in einem Spiel noch zu keinem Handlungsergebnis; stattdessen übersetzt eine Spiel-Welt die Handlungen aller Spieler *im Verbund* in ein Interaktionsergebnis für Alle.

Spiele. Das Konzept eines strategischen Spiels Γ lässt sich in wenigen Schritten aus der obigen Situation entwickeln, in der sich ein einsamer Agent i in einer Welt W bewegt. Dabei wird einiges davon, was oben beschrieben wurde, in modifizierter Form übernommen. Im Ausgangspunkt wird die probabilistische Komplikation ausgespart.

²² Wir ersparen uns die Erörterung der Komplikation, dass der *nächste* Weltzustand vom aktuellen Zustand und der aktuellen Handlung des Agenten abhängen kann. Wichtig ist hier der wesentliche Unterschied eines handelnden Agenten zu der Welt, innerhalb derer er handelt.

²³ Auch die Handlungsoptionen können mit s variieren, $H^i(s)$, was hier der Einfachheit halber unterschlagen wurde.

²⁴ Neumann (1928) zeichnet für die erste Formalisierung verantwortlich.

Ausgangspunkt: Agent i ist in einer Welt W mit möglichen Zuständen $S = \{s_1, s_2, \dots, s_w\}$, er hat in jedem dieser Zustände dieselben Handlungsoptionen $H^i = \{h_1^i, h_2^i, \dots, h_m^i\}$ und weist eine Nutzenfunktion U^i über alle möglichen outcomes seiner Handlungen auf. In jedem möglichen Weltzustand $s \in S$ führt jede der Optionen $h \in H^i$ eindeutig zu einem Ergebnis $o = W(h, s)$. D.h. alle möglichen Outcomes O , die die Welt W bereitstellt, ergeben sich aus allem möglichen Kombinationen von Handlungen und Weltzuständen: $W: H^i \times S \rightarrow O$. Da jedes outcome durch je eine Handlung ($h_1 \dots h_m$) und einen Zustand ($s_1 \dots s_w$) festgelegt ist, lässt es sich über zwei Indizes indentifizieren, d.h. $O = \{o_{1,1}, o_{1,2}, \dots, o_{2,1}, \dots, o_{m,w}\}$. Als rationaler Agent leitet i aus dieses Informationen seine Reaktionsfunktion $RF^{i*}: S \rightarrow H^i$ ab, die ihm für jeden möglichen Zustand $s \in S$ jene Handlungsoption $RF^{i*}(s) = h^{i*,s}$ liefert, die er gemäß U^i allen anderen Handlungsoptionen vorzieht, sobald er mit s konfrontiert ist.

Die Entscheidungssituation (bzw. die Welt W) lässt sich in einer Tabelle (Matrix) darstellen, in der alle möglichen Kombinationen aus Handlungen und Umweltzuständen eingetragen sind:

W	s_1	s_2	s_3	<i>Etc...</i>
h_1^i	$o_{1,1}$	$o_{1,2}$	$o_{1,3}$	...
h_2^i	$o_{2,1}$	$o_{2,2}$	$o_{2,3}$	...
h_3^i	$o_{3,1}$	$o_{3,2}$	$o_{3,3}$	...
<i>Etc ...</i>	...	...	...	...

Nun können die möglichen Outcomes durch die zugehörigen Nutzwerte gemäß der Nutzenfunktion U^i ersetzt werden. Beschränken wir uns auf je drei Optionen und Zustände:

W	s_1	s_2	s_3
h_1^i	$1 = U^i(o_{1,1})$	2	3
h_2^i	2	9	1
h_3^i	5	3	3

Die pro Weltzustand maximal erreichbaren Nutzwerte sind hervorgehoben. So schreibt die Reaktionsfunktion in Zustand s_1 die Handlung h_3^i vor: $RF^{i*}(s_1) = h^{i*,1} = h_3^i$. In Weltzustand s_3 hätte Buridans Esel ein Problem, da es keine einzige beste Aktion gibt. Es ist gleichgültig (gleichwertig), ob der Agent in s_3 seine Option h_1^i oder h_3^i oder aber eine gemischte Handlung auswählt, nach der diese beiden Optionen in beliebiger Mischung auftreten. Die Funktion $RF^{i*}(s_3)$ liefert hier mehrere gleichermaßen nutzenmaximierende Aktionen (darunter unendlich viele gemischte Handlungen). Wir nehmen an, dass Agent i vor des Esels Problem gefeit ist, und also trotz dieser Ambivalenz dazu fähig ist, eine der besten Optionen zu realisieren. Außerdem lässt sich ablesen, dass die Option h_3^i in jeder der möglichen Zustände mindestens genauso gut ist wie die Option h_1^i –

also kann i diese ‚dominierte‘ Option getrost ignorieren und sich auf die anderen beiden konzentrieren.

Um nun zu handeln, muss i eine Vorstellung davon haben, welcher Weltzustand aktuell herrscht. Weiß er gar nichts über den Zustand, muss er zur Not annehmen, dass alle drei Zustände mit gleicher Wahrscheinlichkeit aktuell sein könnten. Folglich wählt er rationalerweise jene Handlung, deren Erwartungswert angesichts dieser Unsicherheit maximal ist; das wäre hier h_2^i mit einer Wert von $EU^i(h_2^i)=12/3$.²⁵ Offensichtlich hängt dieser Erwartungswert von der *subjektiven* Einschätzung der Wahrscheinlichkeiten durch den Agenten ab (Savage 1954). Entscheidend ist, dass der tatsächliche Weltzustand dem Agenten vorgegeben und damit *unabhängig von seinen Präferenzen und seinen rationalen Schlussfolgerungen* festgelegt ist – folglich kann und muss der Agent die Angelegenheit mit dem Zustand als eine reine Frage von Wahrscheinlichkeit bzw. von Zufall behandeln. Insbesondere weist die Feststellung von RF^{i*} keinen systematischen Zusammenhang mit dem aktuellen und aktuell zu erwartenden Weltzustand auf. Kennt der Agent den aktuellen Zustand s genau, kann er getrost gemäß $RF^{i*}(s)$ seinen Nutzen in s maximieren und dabei alle outcomes samt Nutzwerten in allen anderen möglichen Weltzuständen ignorieren, d.h. er kann und muss lokal maximierend vorgehen.

Ein Gegenspieler: Zum Agenten i gesellt sich ein Agent j mit Handlungsoptionen $H^j = \{h_1^j, h_2^j, \dots, h_n^j\}$ und Nutzenfunktion U^j . Die Welt Γ ist nun so beschaffen, dass nicht isolierte Handlungen, sondern erst die möglichen Kombinationen der Handlungsoptionen beider Spieler zu outcomes führen: $\Gamma: H^i \times H^j \mapsto \mathcal{O}$. Die Welt selbst habe hingegen nur einen einzigen, eindeutig festgelegten und beiden Spielern bekannten Zustand, den wir ignorieren können. Beide Spieler müssen ihre Handlungen gleichzeitig und (nur) in diesem Sinne unabhängig voneinander festlegen; d.h. sie können einander Handlungen nicht beobachten, bevor es zu spät ist, die eigene Handlung zu revidieren.

Γ	h_1^j	h_2^j	h_3^j
h_1^i	$O_{1,1}$	$O_{1,2}$	$O_{1,3}$
h_2^i	$O_{2,1}$	$O_{2,2}$	$O_{2,3}$
h_3^i	$O_{3,1}$	$O_{3,2}$	$O_{3,3}$

Anstatt sich nun über den aktuellen Weltzustand Gedanken zu machen – eine Frage von Informiertheit oder Unsicherheit – muss sich Agent i überlegen, welche Handlung nun aktuell wohl vom Gegenspieler j zu erwarten ist. Obwohl die zur Beschreibung von $\tilde{A}$ aufgestellte Tabelle auf den ersten Blick ganz genauso aussieht wie die Darstellung der Welt W eben, handelt es sich um zwei gänzlich verschiedene Situationen. Denn angenommen, dass Agent j rational ist:

²⁵ Wir ersparen es uns, nun auch die Reaktionsfunktion an alle möglichen Unsicherheit über den Weltzustand anzupassen.

dann ist die Handlungswahl des j eine Frage seiner Rationalität und nicht etwa eine Frage von Wahrscheinlichkeiten oder Zufall.

Wenn i also aus irgend Gründen weiß oder glaubt, dass j rational ist, dann muss sich i fragen, welche Handlung aktuell für j rational wäre – und wenn sich so eine bestimmte rationale Handlung h^{j*} des j durch i *rational* ermitteln lässt, dann kann und muss sich i rationalerweise darauf verlassen, dass j tatsächlich h^{j*} realisiert (bzw. realisieren wird). *In der Folge kann der Agent i seinen eigenen Handlungen nicht wie in W einfach auf unabhängig von ihm Wahrscheinlichkeiten gegründete Erwartungsnutzen zuordnen.* Denn anders als in der Welt W ist die tatsächliche, durch die Handlung des j festgelegte ‚Situation‘, mit der i in Γ aktuell rechnen muss, also durchaus *mit einer Eigenschaft des i korreliert*: die Situation wird auf derselben rationalen Grundlage festgelegt wie die Handlung des i selbst, nämlich nach dem Nutzenmaximierungsgesichtspunkt.²⁶ Immer dann wenn es also für i rational ist, seinen eigenen Nutzen zu maximieren, ist es auch für j rational, dies angesichts U^j mittels h^{j*} ebenfalls zu tun. Und diese Rationalität des j legt am Ende fest, welche der in Γ prinzipiell bereitgehaltenen outcomes Agent i tatsächlich durch seine eigene Wahl realisieren kann; das sind all jene outcomes, die h^{j*} zulässt.

Im letzten Schritt muss dann i nur noch die Vorgabe seiner Reaktionsfunktion RF^{i*} angesichts des erwarteten h^{j*} realisieren, wobei diese Funktion nun nicht über Weltzustände, sondern über alle möglichen Handlungen des Gegenspielers j definiert sein muss: $RF^{i*}: H^j \rightarrow H^i$. Diese Funktion legt fest, mit welcher (oder welchen) eigenen Handlung(en) i seinen Erwartungsnutzen maximiert, sofern er eine bestimmte Handlung des j *erwartet* bzw. *annimmt*. Einerseits ist daher die Bezeichnung *Reaktionsfunktion* in Γ eher berechtigt als in W, da j für i ja nicht eine ihn umfassende Welt darstellt, sondern ein gleichberechtigtes Gegenüber, das in derselben Welt Γ agiert wie er. Andererseits sollte die Bezeichnung nicht darüber hinwegtäuschen, dass Agent i seine RF^{i*} schon zu Rate ziehen muss, *bevor* er sicher beobachten kann, wie j wirklich handelt.²⁷ Mithilfe dieser Funktion werden nicht Beobachtungen oder Wissen in Handlungen übersetzt, sondern bloße Erwartungen in eigene Antworten, sodass sie letztlich nur *Reaktionsbereitschaften* des i beinhaltet.

Obwohl diese Übertragung der nutzenmaximierenden Reaktionsfunktion RF^{i*} von W auf Γ wie ein harmloser, ganz folgerichtiger Schritt erscheinen mag, hat er später übrigens erhebliche Konsequenzen, die sich in interessanten Fällen als problematisch erweisen werden. Es wird sich zeigen, dass *das Zusammentreffen nutzenmaximierender Reaktionsfunktionen Kooperationen verhindert*.

Dass hier, im Vergleich zu W oben, etwas kompliziertere Überlegungen angestellt werden können, die auch die (mutmaßliche) Rationalität anderer Agenten berücksichtigen, macht aus dieser Welt Γ ein *strategisches Spiel* zwischen i und j (Neumann / Morgenstern 1961). Es wird auch deutlich, dass

²⁶ Die Welt W hingegen hatte und brauchte keine eigene ‚Nutzenfunktion‘ und konnte ihren Zustand folglich nicht ‚rational‘ festlegen.

²⁷ Deswegen erfordert neben der Ermittlung von RF^{i*} auch dessen Anwendung rationale Urteile.

hier der ‚soziale‘ Aspekt der oben beschriebenen Selbstkonsistenz von Rationalität eine Rolle spielt; denn um zu ermitteln, welche konkrete Handlung für j rational wäre, damit i darauf seine RF^i anwenden kann, muss i auch berücksichtigen, dass j die gleichen Überlegungen hinsichtlich i anstellt. D.h. i muss rational erwägen, was j rational erwägt. Nur wenn i und j gleichermaßen rational überlegen, können sie einander Überlegungen, Entscheidungen und Handlungen *vorhersehen* – und damit können sie auch vorhersehen, dass sie einander Überlegungen vorhersehen. Folglich verschafft ihnen ihre einheitliche Rationalität gewisse Informationen über einander, und zwar mitsamt der Information, dass sie über diese Informationen verfügen – *ohne* dass sie kommunizieren müssten. Dies begründet eine rationale strategische Interaktion. Freilich benötigen wir zunächst Angaben über die Nutzen von i und j, die den verschiedenen möglichen outcomes in der Welt Γ zugeordnet sind. Erst nach dieser Festlegung ist das strategisches Spiel zwischen i und j vollständig definiert: $\Gamma <\{i,j\}, H^i, H^j, U^i, U^j>$. Der Mechanismus Γ vermittelt dabei zwischen allen möglichen Interaktionen und outcomes, denen die Agenten (die Spieler) wiederum Nutzwerte zuordnen können. Folglich sind die Nutzenfunktionen auch über alle möglichen Konstellationen von Handlungen, d.h. über alle in Γ möglichen Interaktionen definiert. Daher können die Outcomes O in der Definition von Γ ausgespart werden. In der folgenden Matrix (der ‚Normalform‘ des Spiels) stehen die Nutzwerte für i jeweils vor (links von) denjenigen für j.

Γ	h^i_1	h^i_2	h^i_3
h^j_1	1, 9	2, 8	3 , 7
h^j_2	2, 8	9 , 1	1, 9
h^j_3	5 , 5	3, 7	3 , 7

Konstantsummenspiel. Die Nutzenfunktion des Agenten i ist in diesem Beispiel analog zu der obigen Situation in W festgelegt worden, um W und Γ leichter vergleichen zu können. Die Reaktionsfunktion RF^j des Agenten j sieht für jede der möglichen Handlungen des i nutzenmaximierende ‚Reaktionen‘ oder besser ‚Antworten‘ vor, die wieder fett hervorgehoben sind. Die für j eingetragenen Nutzenwerte zeigen, warum es sich bei diesem speziellen j im Wortsinne um einen *Gegenspieler* handelt: je höher i ein bestimmtes outcome bewertet, desto niedriger ist dessen Nutzenwert für j, und umgekehrt. Da die Funktionen U_i und U_j genau gegenläufig sind, handelt es sich hier um ein sog. Nullsummen- oder *Konstantsummenspiel*.²⁸ Dies kennzeichnet den ersten Spieltyp, der in der einflussreichen Arbeit von Neumann und Morgenstern (1961) formalisiert wurde.

²⁸ Bei jedem outcome addieren sich die Nutzwerte der Spieler zur selben Konstante, hier 10, sodass für alle outcomes gilt: $U^i(o_{x,y}) = 10 - U^j(o_{x,y})$.

Es ist bemerkenswert, dass diese spezielle Art Spiel einer Reihe von grundsätzlichen Unterschieden zwischen Spielen wie Γ und entscheidungstheoretischen Welten wie W die Schärfe nimmt. Die Spieltheorie und damit die erste Definition von Spielen überhaupt – die in anschließenden Arbeiten (Nash 1951, Selten 1965, Luce / Raiffa 1967, Harsanyi 1968, Aumann 1995) wesentlich übernommen wurde – begann also mit einer Erörterung von Spielen, *in denen Kooperation unmöglich ist*.

Wie wir gleich sehen werden, ist dieselbe Konstanzsummeneigenschaft, die für die Unmöglichkeit von Kooperation sorgt, dafür verantwortlich, dass das Interagieren *mit* anderen Agenten in solchen Spielen dem einsamen Agieren *in* einer Welt W noch recht ähnlich ist. In der Folge werden einige der oben erörterten sozialen Aspekte von rationalen Überlegungen in Spielsituationen in den Hintergrund gedrängt, bis die Konstanzsummenbedingung aufgehoben wird (s.u.). An deren Stelle treten im Konstanzsummenspiel Überlegungen, die aus der individuellen Situation in W übernommen werden, *darunter die a priori Annahme nutzenmaximierender Reaktionsfunktionen*. Insofern kann die Theorie dieser speziellen Spiele als Fortsetzung der Entscheidungstheorie angesehen werden.

Als Vorbild dienten Neumann und Morgenstern (1961) Gesellschaftsspiele wie Schach oder Dame, in denen es darum geht, dem eigenen Sieg möglichst nahe zu kommen, was dort immer dasselbe bedeutet wie, den Gegenspieler einer Niederlage möglichst nahe zu bringen. Näher am Sieg ist, wer einen höheren Nutzwert als der Gegner erzielt.²⁹ Folglich drückt die in RF^{i*} festgelegte Bereitschaft zur Nutzenmaximierung nichts anderes aus, als den Witz der Teilnahme an einem rein antagonistischen Spiel wie Schach: dem Sieg möglichst nahe zu kommen. Daher ist es in allen *Konstanzsummenspielen* rational, jede verfügbare Information dazu zu verwenden, den eigenen Nutzen ohne Rücksicht auf den Nutzen Anderer zu maximieren (und dabei denjenigen des Anderen zu minimieren) – genauso wie es beim Handeln in einer Welt W rational ist, jede verfügbare Information über die Welt W dazu zu verwenden, ausschließlich den eignen Nutzen zu maximieren. Für diese Parallele der beiden RF^{i*} ist es unerheblich, dass ein Agent i in W natürlich nicht so etwas versucht wie ‚die Welt zu besiegen‘, sondern sich schlicht deswegen auf seinen eigenen Nutzen konzentrieren muss, weil es keinen anderen Agenten gibt, dessen Nutzen i zusätzlich berücksichtigen könnte.

In dieser Hinsicht ist das rationale interagieren in einem Spiel wie Γ dem Agieren in einer Welt W sehr *ähnlich*, wenn man vom ontologischen Unterschied zwischen ‚agieren in‘ und ‚interagieren mit‘ abstrahiert: Gegeben 1. die Erwartung h^j in Γ oder die (Information oder) Erwartung s in W , und gegeben 2. dass i seinen Nutzen $U^i(h^i, h^j)$ oder seinen Nutzen $U^i(h^i, s)$ maximieren will, gibt es 3. keine andere rationale Schlussfolgerung als $RF^{i*}(h^j)$

²⁹ Je höher der im Konstanzsummenspiel erreichte eigene Nutzwert ist (d.h. je niedriger der Nutzwert des Gegners ist), desto besser ist diese spezielle Zielsetzung erfüllt. Die Nutzwerte der Spieler können auch relativer ‚Punktestand‘ gelesen werden.

oder $RF^{i*}(s)$ zu realisieren, da bei gegebenem h^j oder s der Nutzwert $U^i[RF^{i*}(h^j), h^j]$ oder $U^i[RF^{i*}(s), s]$ maximal ist. Daher könnte man die *Interpretation* vertreten, dass ein Agent in einem solchen Konstantsummenspiel Γ – fast – wie in einer Welt W handelt, die die Zustände H^j aufweisen kann. Die Ähnlichkeit endet dort, wo dem Agenten i klar sein muss, dass die aktuell anstehende Handlung h^j eines rationalen Gegenspielers auf ganz andere Weise festgelegt wird als der aktuelle Zustand s einer Welt W .

Die umgekehrte Interpretation ist unter Spieltheoretikern recht beliebt (Bernheim 1986, Aumann 1987, Brandenburger / Dekel 1987), aber weniger glücklich: dass ein Agent in einer Welt W so tun könne, *als ob er ein strategisches Spiel gegen die Natur W spiele*. Unabhängig davon, dass W über keine eigene Nutzenfunktion verfügt, sich folglich nicht rational für einen Zustand s entscheidet und dass W folglich die Nutzenfunktion U^1 und damit die mutmaßlich anstehende rationale Handlung des Agenten i stets ignoriert: das Wörtchen ‚gegen‘ schränkt die eben beschriebene Ähnlichkeit von W zu Recht auf Konstantsummenspiele ein. In anderen Spielen, s. u., lassen sich die Nutzwerte, die ein Spieler erzielen kann, nämlich nicht ohne Weiteres mit den Nutzwerten eines einsamen Agenten in W gleichsetzen. In der Folge spielt dort die Reaktionsfunktion eines Spielers eine wesentlich andere Rolle als in W und in Konstantsummenspielen, was bei möglichen Kooperationen eine fundamentale Rolle spielt.

Um zu betonen, was üblicherweise nicht thematisiert wird: Nur deswegen, weil es sich bei Γ um ein Konstantsummenspiel handelt, kann *ohne Einschränkung* die den je eigenen Nutzen lokal maximierende Reaktionsfunktion RF^{i*} aus der alten (einsamen) Entscheidungssituation in W im Prinzip unverändert³⁰ für das Entscheidungsproblem in Γ übernommen werden. Also:

- Agent i will einen möglichst hohen Nutzwert erzielen; und er weiß, dass dies im Konstantsummenspiel gleichbedeutend damit ist, seinem Gegner j einen möglichst niedrigen Nutzwert zu überlassen; Kooperationen sind unmöglich;
- Also kann i rationalerweise nicht erwarten, dass j ihm bei der Verfolgung dieses Ziels in irgend einer Form helfen will; und wenn j ebenfalls rational ist, kann i sich darauf verlassen, dass j versucht, dieses Ziel möglichst zu vereiteln (indem er versucht U^j zu maximieren und damit U^i zu minimieren);
- Also muss i , was immer j tun wird, mit einer nutzenmaximalen Antwort gemäß RF^{i*} aufwarten wollen, obwohl dies dem j sicher nicht gefällt; denn es kann keinen nutzenrelevanten Grund geben, dies aus Rücksicht auf j nicht zu tun. Dasselbe muss i von einem rationalen j erwarten.

Aufgrund dieser (oder äquivalenter) rationaler Überlegungen muss ein rationaler Agent i davon ausgehen, dass ein rationaler Gegenspieler in einem

³⁰ Abgesehen davon, dass S durch H^j ersetzt wird, weil der Unterschied dieser Mengen keine Rolle spielt.

Konstantsummenspiel – ebenso wie i selbst – auf Basis einer nutzenmaximierenden Reaktionsfunktion RF^{j*} entscheidet und handelt. Es ist wesentlich für diese Überlegung, dass eine Kooperation in Konstantsummenspielen schlicht unmöglich ist.

Maximin. Noch einmal das Spiel Γ zur Erinnerung:

Γ	h_1^j	h_2^j	h_3^j
h_1^i	1, 9	2, 8	3 , 7
h_2^i	2, 8	9 , 1	1, 9
h_3^i	5 , 5	3, 7	3 , 7

Um sich in solch einer Situation für eine eigene Handlung zu entscheiden, *muss* ein Agent irgendeine Annahme über die mutmaßliche Vorgehensweise seines Gegenspielers treffen – denn auf Grundlage der Kenntnis des Spiels Γ allein lässt sich keine rationale Entscheidung fällen. So weit ist die Situation analog zu dem Fall, dass Agent i in W handelt, aber nichts über den aktuellen Zustand s weiß. Zwei grundsätzliche Annahmen über den Agenten j kommen für i in Frage: dass j rational ist oder nicht.

Wenn Agent i seinen Gegner j *nicht* für rational hält, wird sich j in Γ aus subjektiver Sicht von i *zufällig* verhalten, d.h. auf für i unvorhersehbare Weise. In diesem Fall ist die nächstliegende Ersatzannahme, dass j jede seiner Handlungsoptionen mit gleicher Wahrscheinlichkeit ($1/3$) realisiert. Unterhält i dieses subjektive Wahrscheinlichkeitseinschätzung, hat seine Option h_2^i den maximalen Erwartungswert ($12/3=4$).³¹ Anders als in einer unbekannten Welt W ist solch eine Annahme in einem strategischen Spiel Γ *nie* rational, selbst wenn Agent i überhaupt nichts über j weiß. Hierin – in der Standardannahme unter vollkommener Unsicherheit – liegt im Kern der wesentliche Unterschied zwischen der individuellen Rationalität in W und der sozialen Rationalität in einem Spiel wie Γ : *Sobald ein anderer Agent j auftaucht, über den i gar nichts weiß, ist es rational, diesen j ebenfalls für rational zu halten.* Erst wenn irgendwelche Evidenzen gegen diese Annahme sprächen, sollte i rationalerweise von dieser Standardannahme abrücken. Solange dies nicht der Fall ist, würde eine vorgeblich ‚rationale Überlegung‘, die zu dem Schluss käme, dass der Andere nicht rational sei, eben jene allgemeine Rationalität in Frage stellen, aus der sie ihre eigene Legitimität beziehen will.

Also nimmt ein rationaler Agent i in Γ an, dass sein Gegner j rational ist, was nach der rationalen Überlegung etwas weiter oben impliziert, dass j die nutzenmaximierende Reaktionsfunktion RF^{j*} unterhält. Neumann und Morgenstern (1961) zeigen, dass sich unter dieser Annahme rational ableiten lässt, dass die Spieler in Γ die Interaktion (h_3^i, h_3^j) realisieren und somit die Nutzwerte $U^i(h_3^i, h_3^j)=3$ und $U^j(h_3^i, h_3^j)=7$ einfahren sollten. Anders als rationale

³¹ Außerdem könnte ein hinreichend ‚naiver‘ Agent i auf die Idee kommen, deswegen h_2^i zu realisieren, weil sie ihm mit etwas Glück den global höchsten Nutzenwert, $U^i(h_2^i, h_2^j)=9$, einbringen kann.

Überlegungen in W nehmen Schlussfolgerungen in Γ den Umweg über *kontrafaktische Situationen*, um dann darauf zu schließen, was tatsächlich zu tun ist. *Nach der Einführung von Wahrscheinlichkeitsverteilungen in logische Schlüsse, stellt die Einbeziehung von solchen was-wäre-wenn-Fragen den nächsten großen Schritt in der modernen Entwicklung der Rationalitätstheorie dar.*

Konkret muss jeder Spieler überlegen, was in Fällen passieren *würde*, in denen einer von beiden schon über die feststehende Handlungsentscheidung des jeweils anderen informiert *wäre* – obwohl solche Fälle in Wirklichkeit gar nicht auftreten. Aus Sicht von i stellt sich die Argumentation so dar:

1. Angenommen j wüsste, welche Handlung h^j der Agent i realisiert. Dann würde j rational mit $RF^j(h^i)$ antworten. Also könnte i in solch einer Situation maximal $U^i(h^i_3, h^j_3)=3$ einfahren, da die Antworten des j auf jede andere Handlung des i zu geringeren Nutzwerten führen würden. Folglich sollte i , wenn er vorlegen müsste, h^i_3 realisieren, und j würde mit $RF^j(h^i_3)=h^j_3$ antworten.
2. Angenommen i hätte den Informationsvorteil und wüsste, welche Handlung h^j der Agent j realisiert. Dann sollte i rational mit $RF^i(h^j)$ antworten. Da j dieser rationale Antwort rational vorhersieht wählt er – analog zu i im Punkt 1 – jene eigene Handlung h^j , die ihm angesichts $RF^i(h^j)$ den maximalen Nutzwert einbringt; dies ist die Handlung h^j_3 . Folglich würde i in dieser Situation $RF^i(h^j_3)=h^i_3$ realisieren. Wieder erhielte i $U^i(h^i_3, h^j_3)=3$.

Obwohl sich die Punkte 1 und 2 auf Situationen beziehen, die in Γ *nicht* auftreten – denn die Agenten fällen ihre Entscheidungen nicht nacheinander, sondern gleichzeitig –, lässt sich aus diesen kontrafaktischen Szenarien etwas darüber folgern, was in Γ tatsächlich rationalerweise passieren muss. Es wäre laut 1 und 2 gleichgültig, ob i zuerst entscheidet (d.h. vorlegen) und die vollinformierte Reaktion des j in Kauf nehmen müsste, oder ob i umgekehrt in der privilegierten Position wäre, seinerseits auf eine vorher festgelegte Handlung des j zu reagieren: in beiden Fällen würden die Agenten dieselbe Interaktion (h^i_3, h^j_3) realisieren. *Folglich müssten sie auch in dem mittleren Fall, dass sie in Γ gleichzeitig entscheiden, genau diese Interaktion realisieren.*

Wenn gemischte Handlungspläne zugelassen werden, gibt es in jedem Konstantsummenspiel mindestens eine rationale Interaktionsmöglichkeit der beschriebenen Art. Diese Lösung wird MaxiMin oder MiniMax-Lösung genannt, da die rationalen Handlungen in einem Konstantsummenspiel eine besondere Eigenschaft aufweisen, die aus der Unmöglichkeit von Kooperation folgt: die rationale Handlung h^i_3 des Agenten i ist genau jene Handlung, die ihm den maximal möglichen Mindestbetrag sichert; d.h. alle anderen Handlungsoptionen würden im jeweils schlechtesten Fall einen niedrigeren Nutzwert bergen. Insofern diese Handlung h^i_3 seine Alternative mit dem

maximalen Mindestnutzwert $U^i(h_3^i, h_3^j)=3$ darstellt,³² kann Agent i mit eben dieser MaxiMin-Handlung mindestens diesen Nutzwert *erzwingen*. Egal, was j tut, er kann nicht verhindern, dass i mindestens diesen Nutzwert erhält (daher ‚erzwingt‘ i diesen Wert). Da sich rationale Gegenspieler in Konstantsummenspiel gegenseitig nichts schenken, erhält i aber auch nicht mehr als eben diesen maximal erzwingbaren Nutzwert. Analoges gilt für die MaxiMin-Handlung h_3^j und den MaxiMin-Nutzwert $U^j(h_3^i, h_3^j)=7$ des Gegenspielers.³³ Wenn alle Teilnehmer eines Konstantsummenspiels rational sind, erhalten sie das und nur das, was sie erzwingen können.

Die Arbeit von Neumann und Morgenstern (1961) hat zwar die formal betriebene Spieltheorie begründet, sie krankt aber an ihrer Beschränkung auf Null- bzw. Konstantsummenspiele. Ob eine strategische Situation ein Konstantsummenspiel darstellt oder nicht, hängt von den *subjektiven* Nutzenzuweisungen der beteiligten Agenten an die outcomes der möglichen Interaktionen ab. Es reicht z.B. nicht, dass die Agenten irgendeinen objektiven outcome unter sich aufteilen müssen, sodass jedes Stück der ‚Beute‘, das Agent i erhält, dem anderen Agenten j vorenthalten wird – denn im Prinzip können die Agenten jedes outcome beliebig bewerten.³⁴ Überhaupt ist es äußerst fraglich, ob es überhaupt zulässig ist, von der *Summe* der individuell-subjektiven Nutzwerte U^i und U^j zu reden, wenn diese Werte etwas anderes darstellen sollen als relative Siegpunkte in einem Gesellschaftsspiel wie Schach oder Dame.

Gleich in der Folgearbeit von Nash (1951) wurde die Konstantsummenbedingung aufgehoben, um nach einem allgemeinen Rationalitätskonzept für beliebige Spiele zu suchen. Darin können die Spieler jedem einzelnen outcome unabhängig von den anderen outcomes – und unabhängig von der Bewertung anderer Spieler – jeden beliebigen Nutzwert zuordnen. Es stellt sich heraus, dass MaxiMin-Handlungen nur in Konstantsummenspielen systematisch rational ist. Dennoch entfaltet die Vorgängerarbeit immer noch einen erheblichen Einfluss, da alle anderen hier erörterten Modellannahmen – insbesondere die Annahme nutzenmaximierender Reaktionsfunktionen – im Wesentlichen beibehalten wurden.

6. Interaktionen mit Kooperationsmöglichkeit

Gleichgewicht. Im Beispiel oben ist das outcome, das der Interaktion (h_3^i, h_3^j) zugeordnet ist, besonders ausgezeichnet: Nur in diesem Ergebnis sollen die Agenten – laut RF^{i*} und RF^{j*} – auf einander Handlungen so ‚antworten‘, wie

³² $\text{Max}_{h^i} \text{Min}_{h^j} U^i = 3$ und $\text{Argmax}_{h^i} \text{Min}_{h^j} U^i = h_3^i$

³³ Außerdem erhält Agent i, wenn er seine MaxiMin-Handlung realisiert, immer einen höheren Nutzwert als den MaxiMin-Nutzwert, falls der Gegenspieler etwas irrationales tun sollte. Wegen dieser Eigenschaften wird die Orientierung am maximal erzwingbaren Nutzwert auch als ‚individuelle Rationalität‘ bezeichnet, obwohl wir gesehen haben, dass an ihrem Anfang die Annahme des Agenten steht, dass auch der andere rational ist.

³⁴ Z.B. könnten es beide am besten finden, dass etwas gleichmäßig aufgeteilt wird. Außerdem könnte ein Agent einem Anteil q der Gesamtbeute X einen höheren oder niedrigeren Nutzwert als $qU(x)$ zuweisen.

jeder von ihnen ohnehin schon agiert.³⁵ D.h. nur bei dieser Konstellation von Handlungen ist $h^i = RF^{i*}(h^j)$ und gleichzeitig $h^j = RF^{j*}(h^i)$,³⁶ sodass die Handlungen h^i_3 und h^j_3 gemäß U^i und U^j gegenseitige ‚beste Antworten‘ aufeinander darstellen. Solch eine Interaktion (h^i_3, h^j_3) wird *gleichgewichtig* genannt, weil sich diese Handlungen angesichts RF^{i*} und RF^{j*} quasi die Waage halten. Hinsichtlich Rationalität drückt diese Eigenschaft die *Selbstkonsistenz* von Überlegungen aus, die genau diese Interaktion betreffen. In gewissem Sinne ‚passen‘ nämlich die Handlungen h^i_3 und h^j_3 angesichts der beiden RF^* zu einander, sodass *auch die rationalen Überlegungen der beiden Agenten, die zur Realisation eben dieser Handlungen führen, einander bestätigen bzw. miteinander vereinbar sind.*

Aus Sicht des Agenten i stellt sich diese Konsistenz so dar: Agent i darf vorläufig annehmen, dass j die Handlung h^j_3 realisiert; dann hieraus ableiten, dass er selbst rationalerweise $h^i_3 = RF^{i*}(h^j_3)$ realisieren soll; dann ableiten, dass j als rationaler Agent diese rationale Schlussfolgerung vorhersieht und sich also auf h^i_3 vorbereitet; und hieraus schließlich folgern, dass der Mitspieler j gemäß RF^{j*} tatsächlich $h^j_3 = RF^{j*}(h^i_3)$ realisiert – was mit der ursprünglichen Annahme von h^j_3 in Einklang steht (konsistent ist). Hingegen führt bei allen anderen Option des j eine vorläufige Annahme des i , dass j sie realisieren werde, zu einem widersprüchlichen Schluss – nämlich dazu, dass j diese Optionen rationalerweise doch nicht realisieren sollte. Dasselbe lässt sich aus Sicht von j zeigen. *Folglich ist (h^i_3, h^j_3) die einzige Interaktion in Γ , die durch i rational gerechtfertigt werden kann, weil sie durch beide Spieler rational gerechtfertigt werden kann, sofern sie einander nur Rationalität zuschreiben.*

Auf dieser Grundlage, dass die rationalen Überlegungen *aller* Spieler mit einander im Einklang stehen müssen, entwickelt Nash (1951) das nach ihm benannte Gleichgewichtskonzept für beliebige Spiele. Das dem Nash-Gleichgewicht zugrunde liegende Kalkül hat sich seither als das Rationalitätskonzept für strategische Interaktionen schlechthin etabliert. Es läuft auf die Selbstkonsistenz von Überlegungen hinaus, die

- sowohl unter der *Annahme* angestellt werden, dass jeder Spieler eine angesichts seiner nutzenmaximierenden Reaktionsfunktion rationale Handlung realisiert und dies auch von jedem anderen Spieler annimmt
- als auch widerspruchsfrei zu der *Schlussfolgerung* gelangen, dass eben dieses Angenommene tatsächlich realisiert wird.

Es handelt sich nicht um einfache Zirkelschlüsse; denn zwischen diesen Annahmen und Schlussfolgerungen vermitteln die Nutzen- und Reaktionsfunktionen *aller* Spieler gleichzeitig. In der Folge kommen als rationale Interaktionen nur Gleichgewichte wie im Beispiel eben in Frage. Das in der Einleitung beschriebene *reasoning* über Handlungen wird nun auf die

³⁵ Wir lassen die Möglichkeit gemischter Handlungspläne beiseite.

³⁶ Mithin ist $h^i = RF^{i*}(RF^{j*}(h^i)) = RF^{i*}(RF^{j*}(RF^{i*}(RF^{j*}(h^i)))$ usw., analoges gilt für h^j .

kollektive Ebene von Interaktionen gehoben: Rationales Schließen, *reasoning*, in strategischen Spielen zerlegt die Menge aller in einem Spiel möglichen Interaktionen der Raisonierenden in (genau) zwei disjunkte Teilmengen, von denen sie die gleichgewichtigen als rational bevorzugt. Solche Gleichgewichte existieren in beliebigen Spielen, sofern auch gemischte Handlungspläne („gemischte Strategien“) zulässig sind (Nash 1951).

Lassen wir gemischte Handlungspläne wieder weg und bleiben wir bei einem strategischen Spiel $\Sigma = \langle \{i,j\}, H^i, H^j, U^i, U^j \rangle$ zwischen nur zwei Agenten i und j mit Handlungsoptionen H^i, H^j und Nutzenfunktionen U^i, U^j . Hieraus lassen sich die nutzenmaximierenden Funktionen RF^{i*} und RF^{j*} wie oben ableiten. Eine Interaktion (h^i, h^j) in Σ ist genau dann Nash-gleichgewichtig, und nach diesem Konzept rational, wenn $h^{i*} = RF^{i*}(h^j)$ und gleichzeitig $h^{j*} = RF^{j*}(h^i)$. In Konstantsummenspielen sind dies die Interaktionen, in denen alle Spieler ihre MaxiMin-Handlungen realisieren, in anderen Spielen nicht unbedingt.

Gleichgewichte werden auch Fixpunkte genannt; sie lassen sich nicht nur für nutzenmaximierende, sondern für beliebige Reaktionsfunktionen ermitteln. Wir können zur Vereinfachung eine allgemeine Abbildung bzw. Operation FIX definieren, die für jedes beliebige gegebene Spiel Σ und beliebige gegebene Reaktionsfunktionen RF^i, RF^j , etc. sämtlicher Spieler die *Menge* der gleichgewichtigen Interaktionen in Σ ermittelt. Für jede dieser durch FIX ermittelte Interaktion gilt, dass die gleichzeitige Anwendung der Reaktionsfunktionen aller Spieler wieder zu eben dieser Interaktion führt. Fassen wir die gleichzeitige Anwendung von $RF^i(h^j)$ und $RF^j(h^i)$ zu $RF(h^i, h^j)$ zusammen, so ergibt sich nur für diejenigen (h^i, h^j) , die FIX für Σ ermittelt, also für alle und nur für $(h^i, h^j) \in \text{FIX}(\Sigma, RF^i, RF^j)$, ein zirkulärer bzw. reflexiver Zusammenhang der Form: $(h^i, h^j) = RF(h^i, h^j) = RF(RF(h^i, h^j)) = RF(RF(RF(h^i, h^j))) \dots$ usw.

Im günstigsten Fall liefert FIX nur eine einzige Interaktion, kann aber auch null oder zwei und mehr Ergebnisse haben. Nash-Gleichgewichte sind jene speziellen Ergebnisse von FIX, die sich bei nutzenmaximierenden Reaktionsfunktionen ergeben (wobei gemischte Handlungspläne zugelassen sind). Z.B. liegt das Nash-Gleichgewicht im folgenden Spiel Σ bei (h_2^i, h_1^j) , es ist also $\text{FIX}(\Sigma, RF^{i*}, RF^{j*}) = \{(h_2^i, h_1^j)\}$, was zu dem einzigen rationalen outcome führt, dem die Nutzwerte (5,3) zugeordnet sind.

Σ	h_1^j	h_2^j
h_1^i	3, 1	1, 8
h_2^i	5, 3	2, 2

Die MaxiMin-Handlung h_2^j des Agenten j ist in Σ hingegen nicht rational, und der rationale Nutzwert $(U^j(h_2^i, h_1^j)=3)$ liegt *oberhalb* seines mittels h_2^j mindestens erzwingbaren MaxiMin-Nutzwerts $(U^j(h_2^i, h_2^j)=3)$. Mit der rationalen Handlung h_1^j wird zwar nicht sicher ausgeschlossen, dass j den global niedrigsten Nutzwert $U^j(h_1^i, h_1^j)=1$ einführt – doch unter der Annahme, dass

sein Mitspieler i rational ist und daher rationalerweise h_2^i realisiert, muss j diese Katastrophe gar nicht fürchten und also keine extra Vorkehrungen treffen. Schon die Tatsache, dass die Option h_2^i die Alternative h_1^i dominiert (also in jedem Fall mehr Nutzen birgt), könnte j davon überzeugen, dass i sicher h_2^i realisiert, worauf j rationalerweise mit $RF^{i*}(h_2^i) = h_1^i$ zu antworten hätte. Da die Konstellation zudem ein Nash-Gleichgewicht darstellt, weil gleichzeitig $RF^{i*}(h_1^i) = h_2^i$, kann ein rationaler Agent i ohne Widerspruch annehmen und schlussfolgern, dass j tatsächlich h_1^i realisiert – sofern j rational ist.

Ein trotz Rationalität unlösbares Problem ergibt sich für den Agenten i in Σ allerdings dann, wenn j *irrationalerweise* die Handlung h_2^i meint realisieren zu sollen – etwa, weil dies seine MaxiMin-Handlung ist oder vielleicht, weil j darauf hofft, den sehr hohen Nutzwert 8 einfahren zu können, etc. In dem Fall werden die irrationale Interaktion (h_2^i, h_2^j) und damit die Nutzwerte (2,2) realisiert. Im Vergleich zu den rationalen Nutzwerten (5,3) schadet sich somit j durch seine Irrationalität nicht nur selbst – er fügt auch seinem rationalen Mitspieler i einen erheblichen Schaden zu. Das Beispiel zeigt, dass ein rationaler Agent in bestimmten Spielen vollständig darauf vertrauen muss, dass seine Mitspieler ebenfalls rational sind, weil deren Irrationalität auch ihm schaden würde. Die in Konstantsummenspielen herrschende Gewissheit eines rationalen Agenten, dass er auch bei einem irrationalen Mitspieler *mindestens* seinen rationalen Nutzwert einfährt, fällt bei vielen anderen Spielen weg. Dies ist die Kehrseite dessen, dass rationale Interaktionen in Spielen wie Σ eben mehr als nur die MiniMax-Nutzwerte bergen: die rationalen Nutzwerte lassen sich dort nicht durch individuelle rationale Handlungen erzwingen. *Folglich kann ein rationaler Spieler mit einem irrationalen Gegenüber nicht rational umgehen.* In solchen Fällen hilft es nur noch, sich auf die eigene MaxiMin-Handlung zurückzuziehen – *nicht* weil dies rational wäre, sondern weil man sich dadurch wenigstens einen Mindestnutzwert sichert, auf den man sich verlassen kann. MaxiMin ist außerhalb von Konstantsummenspielen also eine Art *last resort*, die vor vollkommener Ungewissheit und damit vor ungewissen Schäden schützt.

Common Belief. Eine gemeinsame und einheitliche Rationalität – das Durchführen derselben rationalen Schlussfolgerungen aufgrund derselben Informationen und Annahmen – führt zu einer allgemeinen Abstimmung (Koordination) der Agenten. Man würde zu weit gehen, wenn man dies schon als eine Art Kooperation durch gemeinsame Rationalität interpretieren würde. Doch immerhin richten sich Agenten, die vor ihrer Interaktion dieselben Schlussfolgerungen einstellen, nicht nur auf dieselben Interaktionsergebnisse und damit verbundene Nutzwerte ein: dieselben Schlussfolgerungen *führen* auch zu eben diesen vorher *von Allen* erwarteten Ergebnissen.

Hierin spiegelt sich wider, dass solche rationalen Schlussfolgerung auf Grundlage gemeinsamer Informationen und Annahmen angestellt werden – und werden müssen. Diese recht starke Voraussetzung ist in der Spieltheorie als *common knowledge* bekannt (Aumann 1976, Monderer / Samet 1989). Der terminus geht auf eine philosophische Arbeit von Lewis (1969) zurück, und ist

etwas irreführend. Gemeint ist nicht ‚gemeinsames Wissen‘ i.S.v. sicherer Information (weder in der Spieltheorie noch bei Lewis), sondern vielmehr eine ‚gemeinsame Annahme‘: *common belief*. Beide Konzepte haben seither in der Modallogik, bzw. in deren Unterabteilung epistemische Logik, Karriere gemacht (Vgl. Chellas 1980). Ein *common belief* herrscht in einer Population von Agenten genau dann vor, wenn alle Agenten gleichermaßen glauben, dass ϕ wahr ist; wenn zudem alle Agenten glauben, dass alle Agenten glauben, dass ϕ wahr ist; wenn weiterhin Alle glauben, dass Alle glauben, dass Alle glauben, dass ϕ wahr ist ... etc. ad infinitum. Analoges gilt für *common knowledge*, nur dass dort auch ϕ wirklich gewusst und folglich auch wirklich wahr sein muss. Diese Konnotation ist in der Spieltheorie durchaus gewollt: man nimmt theoretisch an, dass das, was unter den Spielern gemeinsam angenommen wird (was also *common belief* ist) auch tatsächlich wahr ist.

Im Ausgangspunkt ist die Struktur des gemeinsamen Spiels Σ *common belief*, was im Besonderen bedeutet, dass jeder Spieler die den outcomes zugeordneten Nutzwerte *aller* beteiligten Agenten richtig einschätzt. Außerdem ist es *common belief*, dass alle Agenten rational sind. Folglich ist es *common belief*, dass alle Spieler angesichts von Σ dieselben rationalen Schlüsse ziehen. Daher sind schließlich auch die rationalen Interaktionen in Σ *common belief*.

Ohne diese gemeinsame Grundlage könnte es überhaupt keine rationalen Schlussfolgerungen geben, da diese, wie wir gesehen haben, darauf beruhen, dass jeder rationale Agent annimmt, dass auch alle rationalen Anderen seine Schlussfolgerungen durchschauen. Angenommen, dass in einem Spiel zwischen den Agenten i und j eine Interaktion (h^{i*}, h^{j*}) Nash-gleichgewichtig ist. Dann stellt sich die zugehörige selbstkonsistente rationale Überlegungen des Agenten i folgendermaßen dar: i nimmt vorläufig an, dass j seine Option h^{j*} realisiert; er folgert daraus $h^{i*} = RF^{i*}(h^{j*})$; er folgert, dass ein rationaler j diese rationale Folgerung von h^{i*} durchschaut und sich darauf einrichtet; er folgert schließlich $h^{j*} = RF^{j*}(h^{i*})$, was mit seiner ersten Annahme übereinstimmt. Analoges gilt für j . Der mittlere Schritt in dieser Kette von rationalen Schlussfolgerungen von besonderer Brisanz: „ i folgert, dass ein rationaler j diese rationale Folgerung von h^{i*} durchschaut und sich darauf einrichtet“. – *Es geht nicht um Telepathie, sondern darum, dass diese gesamte Kette von Schlussfolgerungen des i unter rationalen Agenten mit nutzenmaximierenden Reaktionsfunktionen common belief sein kann.*

Dasselbe lässt sich auch so ausdrücken: nur solche Interaktionen eines Spiels Σ zwischen i und j mit Reaktionsfunktionen RF^{i*} und RF^{j*} können von beiden Agenten sowohl angenommen als auch aufgrund dieser Annahme rational inferiert werden, die zur Menge $FIX(\Sigma, RF^{i*}, RF^{j*})$ gehören. *Es ist kein Zufall, dass sowohl unser Operator FIX als auch selbstkonsistente Schlussfolgerungen als auch common belief zirkuläre i.S.v. reflexive Strukturen aufweisen.* Nash-gleichgewichte sind spezielle Fälle aller Interaktionen, die mit dieser Reflexivität (oder Zirkularität) kompatibel sind. Was sie besonders auszeichnet,

ist die Annahme jener speziellen Reaktionsfunktionen, die aus der Entscheidungstheorie in die Theorie rationaler Interaktion übernommen wurden. Irrationale Schlussfolgerungen eines Agenten i enthalten entweder Widersprüche im engen Sinne enthalten oder aber die Annahme, dass andere Agenten diese Schlussfolgerungen *nicht* nachvollziehen können. Dies durchbricht die für common belief nötige Zirkularität aus ‚ i nimmt an, dass j annimmt, dass i annimmt ...‘. Ein rationaler Agent darf also nur Annahmen treffen, von denen er gleichzeitig annimmt, dass alle anderen Agenten eben diese Annahmen annehmen werden. Täuschungsversuche können folglich nicht in diesem Sinne rational sein. *Weil diese Konzeption von Rationalität auf common belief gegründet, geht sie immer mit Transparenz einher.*

Multiple Rationalitäten. Das auf common knowledge, eigentlich auf common belief, gegründete Rationalitätskonzept der Spieltheorie hat allerdings zwei Schönheitsfehler. Zwei sehr bekannte Spiele, das Chicken-Spiel und das Gefangenendilemma, sind für diese beiden Probleme sinnbildlich. Zunächst Chicken.

Λ (Chicken)	h_1^i	h_2^i
h_1^i	3, 3	2, 4
h_2^i	4, 2	1, 1

Das Spiel Λ wird gerne ‚Chicken‘ genannt, weil man es mit einer story um Mut und Feigheit verbinden kann; etwa so: Die Agenten i und j rasen mit ihren Autos aufeinander zu; wenn einer von beiden die Nerven verliert und alleine ausweicht (h_1), hat er verloren, was durch die Nutzwerte (1,4) zu seinen Ungunsten ausgedrückt wird; weichen beide gleichzeitig aus, ergibt sich ein Unentschieden (3,3); bleiben beide auf Kurs (h_2), gibt es schließlich einen fatalen Crash (h_2^i, h_2^j), den beide sicher bereuen werden (1,1).

Es gehört schon eine besondere ‚Rationalität‘ dazu, überhaupt an solch einem Spiel teilnehmen zu wollen, wenn man die Wahl hat. Denn dieses Spiel hat zwei Nash-Gleichgewichte $\text{FIX}(\Lambda, \text{RF}^{i*}, \text{RF}^{j*}) = \{(h_1^i, h_2^j), (h_2^i, h_1^j)\}$, die die Spieler verschieden gut finden, da in dem einen Gleichgewicht i und in dem anderen j gewinnt. Tragischerweise kommt es zum Crash, wenn beide Spieler darauf beharren, das für sie jeweils günstigere rationale outcome zu realisieren. Hier zeigt sich eine Verschärfung des oben festgestellten Problems, wonach die rationale Handlung eines Spielers bei irrationalen Mitspielern noch lange keinen rationalen Mindestnutzwert garantiert: *Selbst die Rationalität aller Spieler garantiert nicht, dass eine rationale Interaktion realisiert und damit rationale Nutzwerte eingefahren werden, da es in einem Spiele mehrere rationale Interaktionen geben kann.*

Das Problem ist, dass es sowohl common belief sein kann, dass j ausweichen und i durchfahren wird (h_1^i, h_2^j), als auch dass umgekehrt i ausweicht und j

gewinnt (h_2^i, h_1^j) . Die Schlussfolgerungen wären gleichermaßen rational – sie hätten aber nur dann den jeweils vorgesehenen Erfolg (Nutzwert), wenn beide Spieler gleichzeitig denselben Schluss zögen. Hierfür müssten die Spieler dieselbe rationale Interaktion – (h_1^i, h_2^j) oder (h_2^i, h_1^j) – als Prämisse ihre Inferenzen nehmen. *Welche der beiden gleichermaßen rationalen Annahmen ein Agent treffen soll ist aber nicht rational ableitbar.* Keine rationale Überlegung führt zu dem Schluss, dass unbedingt die Annahme und damit die Schlussfolgerung (h_1^i, h_2^j) common belief sein *müsse*; das gleiche gilt für das andere Nash-Gleichgewicht.

Dieser Sachverhalt wird gemeinhin unter dem Aspekt behandelt, dass es für ein Rationalitätskonzept wünschenswert wäre, wenn es pro Spiel genau eine Interaktion als rational auswies. Ansonsten drohen rational nicht aufzulösende Unbestimmtheiten wie hier, die am Ende dazu führen, dass nicht einmal garantiert ist, dass überhaupt eine rationale Interaktion stattfindet. – In diesem Sinne lässt das spieltheoretische Rationalitätskonzept *multiple Rationalitäten* zu. Für unser Thema ist noch eine weitere Eigenschaft des Spiels Λ aufschlussreich: dass die Interaktion (h_1^i, h_1^j) , die jeder der Spieler sowohl dem drohenden Crash als auch einer ‚Niederlage‘ vorzieht, nicht Nash-gleichgewichtig und daher nach diesem Konzept nicht rational ist. Diese Interaktion muss nicht als ‚beide sind feige‘ interpretiert werden; es geht allgemein darum, dass beide Nachgeben, anstatt auf ihren eigenen Triumph zu beharren und damit eine Katastrophe zu riskieren. Diese Interaktion ist ein *Kompromiss*, für den nur jeder der Spieler ein wenig nachgeben muss, um immerhin das zweitbeste global mögliche Ergebnis zu erzielen.

Den Kompromiss zu realisieren, erfordert aber ein gewisses gegenseitiges Vertrauen, da z.B. Spieler j dem Anreiz ausgesetzt ist, eine kompromissbereite Vorlage des i (h_1^i) mit einer aggressiven Handlung (h_2^j) zu beantworten – und genau diese unfreundliche Antwort schreibt die nutzenmaximierende Reaktionsfunktion vor: $RF^j(h_1^i) = h_2^j$. Dasselbe gilt analog für i . In der Folge wird eine solcher Kompromiss rationalerweise gar nicht erst versucht, wenn und weil die Funktionen RF^{i*} , RF^{j*} common belief sind. *Wir wollen diese einseitige Ausbeutbarkeit als Merkmal jeder Kooperation auffassen.* Diese Bestimmung von Kooperation macht schon deutlich, dass *nutzenmaximierende Reaktionsfunktionen systematisch Kooperationen verhindern.*

Dilemma. Das Dilemma Δ zeigt, dass spieltheoretisch rationale Schlussfolgerungen dazu führen können, dass Agenten eine offensichtliche Chance zu einer Kooperation verpassen.

Δ (Dilemma)	h_1^j	h_2^j
h_1^i	3, 3	1, 4
h_2^i	4, 1	2, 2

Das einzige Nash-Gleichgewicht liegt bei FIX $(\Delta, RF^{i*}, RF^{j*}) = \{(h^i_2, h^j_2)\}$, dem die Nutzwerte (2,2) zugeordnet sind – obwohl es common belief ist, dass die Agenten mit (h^i_1, h^j_1) die Nutzwerten (3,3) realisieren würden. Diese Interaktion (h^i_1, h^j_1) wäre eine *Kooperation* im eben bestimmten Sinne, da sie ein gewisses gegenseitiges Vertrauen erfordert. Für beide Agenten bestehen Anreize zur Ausbeutung des jeweils Anderen $((h^i_1, h^j_2)$ und $(h^i_2, h^j_1))$, wenn dieser kooperativ handeln sollte.

Der rationale Grund, der die Kooperation verhindert, ist offensichtlich: *wenn* RF^{i*} und RF^{j*} common belief sind, führt eine Annahme von (h^i_1, h^j_1) eben nicht konsistent zur Schlussfolgerung, dass (h^i_1, h^j_1) tatsächlich realisiert wird. Denn im Laufe einer Überlegung ergibt sich $RF^{i*}(h^j_1) = h^i_2 \neq h^i_1$ und außerdem $RF^{j*}(h^i_1) = h^j_2 \neq h^j_1$. Es kann m.a.W. nicht common belief sein, dass beide Agenten von einer Annahme von (h^i_1, h^j_1) und von (RF^{i*}, RF^{j*}) auf die Realisation von (h^i_1, h^j_1) schließen. Folglich ist die kooperative Interaktion (h^i_1, h^j_1) im Gefangenendilemma nicht rational – zumindest solange nicht, wie die Nutzenmaximierungen gemäß RF^{i*} , RF^{j*} common belief sind, was in der Spieltheorie nach Nash (1951) standardmäßig vorausgesetzt wird.

Nach allem, was in dieser Arbeit erörtert wurde, kann es kaum bezweifelt werden, dass es in Δ rational ist, das unkooperative Nash-gleichgewicht zu realisieren. Alle Argumente, die gegen die Rationalität dieser Nichtkooperation ins Feld geführt werden könnten, lassen sich leicht durch Verweise auf die Selbstkonsistenz der zugehörigen Überlegungen kontern. Dass die der Kooperation zugeordneten Nutzwerte höher sind als diejenigen im Nash-Gleichgewicht reicht nicht als Einwand. Die gesamte Struktur des Spiels ist zu berücksichtigen, und in dieser Struktur liefert die schiere Möglichkeit der einseitigen Ausbeutung von Kooperationsangeboten einen guten Grund gegen Kooperationsversuche und für das Nash-Gleichgewicht. Die Kehrseite der Tatsache, dass Nash-gleichgewichtige Rationalität Kooperationsgelegenheiten systematisch verpasst, ist nämlich, dass diese Rationalität *sicher verhindert, dass ein Agent durch Andere ausgebeutet wird*. In diesem Punkt liefert diese Rationalität eine Sicherheit, die sich mit der Sicherheit der MaxiMin-Vorgehensweise vergleichen lässt.

Solange man sich die Struktur des Spiels Δ nur abstrakt anschaut, kann es der Intuition zuwider laufen, und es kann sogar bedauert werden, dass die durch gegenseitiges Misstrauen bestimmte Nichtkooperation im Dilemma rational ist. Doch wenn man sich ein gängiges Anwendungsbeispiel aus der Wirtschaft anschaut, kann sich das Bild ändern: wenn sich an einem speziellen Markt nur einige wenige Anbieter befinden, dann bestünde eine Kooperation darin, sich zu einem Kartell zusammenzuschließen und ihre Kunden mit überhöhten Preisen zu konfrontieren – und dabei das Risiko einzugehen, dass einer der Beteiligten ausschert und die anderen unterbietet; die spieltheoretisch rationale Nicht-Kooperation besteht in der *volkswirtschaftlich* wünschenswerten Konkurrenz dieser Anbieter. Analog lässt sich der Fall einordnen, dass in einer Organisation leistungsschwächere Mitarbeiter verbünden und einen starken Konkurrenten

durch unlautere Mittel aus dem Weg zu räumen, ohne dass Vorgesetzte das bemerken – wobei die Kooperateure das Risiko eingehen, dass einer von ihnen die Chance erhält, diesen Plan zu seinem Vorteil zu verraten. Überhaupt überall dort, wo einige Agenten *zum Nachteil Dritter* kooperieren könnten, ist eine Kooperation aus globaler Sicht nicht unbedingt erwünscht.

7. Fazit: ist rationale Kooperation überhaupt möglich?

Wenn man also keine starke Alternativtheorie zu derjenigen anzubieten hat, die in der hier grob nachgezeichneten Tradition der Entscheidungstheorie steht, dann muss man durchaus im Einklang mit der philosophischen Bestimmung von Rationalität zugestehen: es ist rational, im Gefangenendilemma nicht zu kooperieren und ebenso rational, keine Kooperativität vom Anderen zu erwarten. – Dennoch bleibt eine Frage offen: *ob diese Rationalität die einzige mögliche Rationalität darstellt, oder ob darüber hinaus nicht noch weitere, insbesondere kooperativere einzubeziehen sind, die ebenfalls die generelle Selbstkonsistenz aufweisen.*

Die vorangehenden Ausführungen sollten nicht nur zeigen, dass die spieltheoretische Rationalität, die sich Nash-Gleichgewichtskonzept niederschlägt, zu Recht als Rationalität bezeichnet wird. Sie sollten, quasi unter diesem Vorwand, auch zeigen, auf welchen theoretischen Konzepten diese Einschätzung basiert, um mögliche Angriffspunkte für Kritik und Modifikationen sichtbar zu machen. Folgende Punkte könnten kritisch beurteilt werden:

- Die Annahme, dass die den outcomes zugeordneten Nutzwerte in strategischen Spielen, insbesondere in Nicht-Konstansummenspielen, im Wesentlichen dieselbe Rolle spielen wie die Nutzwerte, die ein einsamer Agent den outcomes seiner Handlungen in einer natürlichen Umwelt zuordnet. M.a.W. die Annahme, dass eine Analogie zum Agieren in verschiedenen, von Rationalität unabhängigen Weltzuständen auch in Nicht-Konstansummenspielen gegeben sei.
- Die Annahme, dass ein rationaler Agent alle möglichen Handlungen anderer Spieler unbedingt gemäß einer nutzenmaximierenden Reaktionsfunktion beantworten wolle und müsse. Die damit verbundene Annahme, dass es unter rationalen Spielern unbedingt common belief *und daher* korrekt – quasi common knowledge – sein müsse, dass sie alle nutzenmaximierende Reaktionsfunktionen aufwiesen. – Überhaupt die Annahme, dass eine entscheidungstheoretisch begründete Reaktionsfunktion ohne weiteres als die einzige mögliche rationale Option in die Theorie sozialer Interaktionen übernommen werden könne.
- Die Annahme, dass eine *lokal* nutzenmaximierende Bereitschaft, eine kooperative Vorlage des Anderen durch Nichtkooperation zu

beantworten, ohngeachtet ihrer globalen Folgen für das ganze Spiel schon als ‚rational‘ zu qualifizieren sei – obwohl die zugehörige Interaktion selbst nicht Nash-gleichgewichtig, also global irrational wäre.

- Die Annahme, dass das allgemeine Konzept ‚Rationalität‘ möglichst immer nur eine Lösung pro gegebener Problemstellung liefern müsse; und dass man daher erforschen müsse, wie die Menge der als rational qualifizierbaren Interaktionen in einem Spiel weiter eingeschränkt werden könnten – was dazu führt, sich auf die Nash-Gleichgewichte zu konzentrieren und einige davon aussondern zu wollen, anstatt die Menge rationaler Interaktionen z.B. auch auf Kooperationen auszuweiten.

Meines Erachtens (Kabalak 2009) lässt sich Kooperation durchaus mit Rationalität vereinbaren, ohne deswegen bezweifeln zu müssen, dass Nash-Gleichgewichte rational sind. Ohne das hier im Detail ausführen zu können: So wie das spieltheoretische Konzept längst multiple Rationalitäten zulässt, unter denen naturgemäß nicht rational diskriminiert werden kann, können Argumente dafür ins Feld geführt werden, dass in bestimmten Spielen neben den Nash-Gleichgewichten auch kooperative Interaktionen das Prädikat ‚rational‘ verdienen. Es folgt, dass Kooperationen *nicht zwingend* rational sind, sondern lediglich mögliche rationale *Optionen* neben ebenso rationalen unkooperativen Alternativen darstellen. Dass Nash-Gleichgewichte rational sind, heißt nicht, dass Kooperationen irrational seien, und umgekehrt. Es geht vielmehr darum, irrationale Interaktionen zu identifizieren und auszusondern. Um Kooperationen auf der rationalen Seite einzuordnen, müssen insbesondere Ausbeutungsversuche als irrational disqualifiziert werden; d.h. ein rationaler Agent darf nicht schon dadurch an einem Kooperationsversuch gehindert werden, weil er annehmen müsste, dass ein rationales Gegenüber diesen Versuch unbedingt ausbeuten würde, sobald er ihn erwartet. Hierbei hilft eine Einsicht, die oben an verschiedenen Stellen auftauchte: Alle Überlegungen, die auf Intransparenz und Undurchschaubarkeit, auf die Überraschung und Übervorteilung des Anderen, auf ein vom Anderen nicht akzeptables outcome hinauslaufen, sind mit der Reflexivität von Rationalität nicht vereinbar – selbst wenn ein Agent dadurch seinen individuellen Nutzen maximieren würde.

Es kann nach dieser Auffassung zwar nie rational inferiert werden, dass ein Gegenüber unbedingt rational sein müsse – es kann aber rational angenommen und konsistent geschlussfolgert werden, dass ein Gegenüber kooperativ *sein kann*. Voraussetzung dafür ist die vor-rational zu treffende gemeinsame Annahme aller Spieler, dass sich alle zur rationalen Kooperation statt zur ebenso rationalen Nicht-Kooperation entscheiden, ohne dazu gezwungen zu sein. Dies läuft darauf hinaus, dass neben nutzenmaximierenden Reaktionsfunktionen auch andere *Konstellationen von Reaktionsfunktionen aller Spieler* zulässig sind, solange diese Konstellationen im Fixpunkt zu allseits akzeptablen Ergebnissen führen. Nicht die aus der Entscheidungstheorie übernommene Nutzenmaximierung wäre also an den Anfang zu setzen, um daraus

Gleichgewichte abzuleiten; sondern *zuerst* wären akzeptable Interaktionen und outcomes zu bestimmen, die durch alle rationalen Spieler akzeptiert werden und damit common belief sein können, um hieraus die zu diesen outcomes passenden Reaktionsfunktionen samt dazugehöriger Annahmen und Inferenzen abzuleiten. Das in Kabalak (2009) vertretene Kriterium für akzeptable outcomes ist: jedes outcome, dem jeder der Spieler mindestens ebenso viel Nutzwert zuordnet wie dem für ihn selbst schlechtesten Nash-Gleichgewicht, kann durch alle Spieler akzeptiert werden. Im Dilemma sind das z.B. die Kooperation und das Nash-Gleichgewicht, nicht aber die einseitigen Ausbeutungen; und im Chicken-Spiel sind das die beiden Nash-Gleichgewichte und der Kompromiss, nicht aber der Crash.

Literatur

- Aumann, R. (1987): Correlated Equilibrium as an Expression of Bayesian Rationality, *Econometrica* 55/1: 1-18
- Aumann, R. (1995): Repeated Games with Incomplete Information, MIT Press.
- Bernheim, B. (1986): Axiomatic Characterisation of Rational Choice in Strategic Environments, *Scandinavian Journal of Economics* 88, 3: 473-488
- Binmore, K. (2000): [Game theory and the social contract](#), Cambridge: MIT Press.
- Brandenburger, A. / Dekel, E. (1987): Rationalizability and Correlated Equilibria, *Econometrica* 55,6: 1391-1402.
- Carnap, R (1952): The Continuum of Inductive Methods, Chicago: Univ. of Chicago Press.
- Chellas, B. (1980): Modal Logic, Cambridge: Cambridge Univ. Press.
- Gauthier, D. (1986): Morals by Agreement, Oxford: Oxford Univ. Press.
- Harsanyi, J. (1968): Games with incomplete information played by 'bayesian' players, Teile i-iii, in: *Management Science*, 14:159–182, 320–334, 486–502.
- Lewis, D. (1969): Convention: A Philosophical Study. Cambridge, Massachusetts: Harvard University Press.
- Luce, D. / Raiffa, H. (1967): Games and Decisions: Introduction and critical survey, New York
- Kabalak (2009): [Institutionelle Spiele: Ein neuerer akteurstheoretischer Zugang zu Rationalität und Institutionen](#), Marburg: Metropolis.
- Pumplün, D. (1999): Elemente der Kategorientheorie, Berlin: Spektrum.
- Ramsey, F. (1930): The Foundations of Mathematics and other Logical Essays (Repr. der Ausgabe London 1930), Saarbrücken: Dr. Müller, 2006.
- Savage, L.J. (1954): The Foundation of Statistics, New York: Wiley.
- Selten, R. (1965): Spieltheoretische Behandlung eines Oligopolmodells mit Nachfrageträgheit, in: *Zeitschrift für die Gesamte Staatswissenschaft* 121: 301-24 und 667-89.
- Monderer, D. / Samet, D. (1989): Approximating Common Knowledge with Common Beliefs, *Games and Economic Behavior*, 1: 170-190

- Nash, J. (1951): Non-Cooperative Games, in: Annals of Mathematics 54, No. 2, September 1951: 286-295.
- Neuman, J. (1928): Zur Theorie der Gesellschaftsspiele, in: Mathematische Annalen, Bd. 100: 295 – 320.
- Neumann, J. / Morgenstern, O. (1961): Spieltheorie und wirtschaftliches Verhalten, Würzburg: Physica.
- Pearl, J. (2010): Causality – Models, Reasoning, and Inference, Cambridge: Cambridge Univ. Press.
- Walliser, B. (1992): Epistemic logic and game theory, 197-225, in: Bicchieri et al (Hg): [Knowledge](#), [belief](#), and [strategic interaction](#), Cambridge: Cambridge Univ. Press.
- Weber, M. (1990): [Wirtschaft](#) und [Gesellschaft](#) : Grundriss der verstehenden Soziologie, 5., rev. Aufl., Nachdr., Studienausg., besorgt von Johannes Winckelmann, Tübingen: Mohr.
- Wittgenstein, L. (2010): Tractatus logico-philosophicus, München: Oldenbourg-Akademie-Verlag.