

E. Eigener experimenteller Ansatz zur algorithmischen Kollusion auf heterogenen Märkten

Die bisher verfügbaren Untersuchungen realer Märkte sowie Simulationen mit selbstlernenden Algorithmen deuten auf ein kollusives Potenzial algorithmischer Preissetzung hin.⁷⁸⁷ Allerdings wurde der wettbewerbsschädliche Effekt der Algorithmen überwiegend in Märkten mit homogenen Wettbewerbern aufgezeigt. Ergebnisse heterogener Marktzusammensetzungen deuten darauf hin, dass insbesondere komplexe selbstlernende Algorithmen Schwierigkeiten haben, kollusive Gleichgewichte zu erreichen.⁷⁸⁸ Empirische Befunde legen jedoch nahe, dass sich Händler in einem großen Teil der Märkte einem heterogenen Wettbewerb gegenübersehen. Die Untersuchungen verschiedener Plattformen deuten darauf hin, dass noch immer ein großer Teil der Händler keinen Gebrauch von algorithmischer Preissetzung macht.⁷⁸⁹ Im Einklang damit stehen die Ergebnisse der Sektoruntersuchung zum elektronischen Handel, bei der circa ein Drittel der befragten Händler angab, automatisierte Preissetzung einzusetzen.⁷⁹⁰ Deswegen ist davon auszugehen, dass es derzeit häufig zur Interaktion zwischen menschlichen und algorithmischen Anbietern kommt. Untersuchungen zur Auswirkung algorithmischer Preissetzung in heterogenen Marktzusammensetzungen können somit eine nützliche Ergänzung zu den bereits gewonnenen Erkenntnissen darstellen.⁷⁹¹ Hierzu habe ich zusammen mit *Hans-Theo Normann* ein Laborexperiment durchgeführt.⁷⁹²

In unserem Laborexperiment untersuchen wir die Entwicklung der Marktpreise auf Märkten, auf denen Algorithmen und menschliche Entscheider im Wettbewerb zueinanderstehen (im Folgenden auch hybride Märkte) und vergleichen diese mit Märkten, auf denen ausschließlich menschliche Anbie-

787 Vgl. Kapitel D. II. und III.

788 Vgl. Kapitel D., siehe hierzu *Hettich* (2021); *Kastius/Schlosser* (2022), J Rev Pricing Man 21, 50.

789 Vgl. Kapitel D. III.

790 Vgl. *Europäische Kommission*, Abschlussbericht über die Sektoruntersuchung zum elektronischen Handel.

791 Den generellen Bedarf für eine Betrachtung heterogener Märkte sehen auch *Chen et al.*, WWW '16, 1339 (1348); *Calvano et al.* (2020), AER 110 (10), 3267 (3296).

792 *Normann/Sternberg* (2021).

ter interagieren. Im Fokus steht dabei die Frage, inwiefern algorithmische Wettbewerber die Kollusionsbereitschaft menschlicher Anbieter beeinflussen können. Eine Kollusion ohne vorherigen Kommunikation ist mit großen Risiken für den Entscheidungsträger verbunden. Entscheidet er sich dazu, einen kollusiven Preis festzulegen, geht er das Risiko ein, von seinen Wettbewerbern ausgebeutet zu werden. Der menschliche Anbieter muss deshalb das Verhalten seiner Wettbewerber beobachten und zu einer Einschätzung bezüglich ihrer Marktstrategien gelangen. Nur wenn er es für hinreichend wahrscheinlich hält, dass die Wettbewerber an einer Kollusion interessiert sind, kann sich für ihn das Risiko einer höheren Preissetzung lohnen. Fraglich ist, inwiefern der Einsatz von Preissetzungsalgorithmen auf Seiten der Wettbewerber dazu beitragen kann, dieses Risiko zu reduzieren und *tacit collusion* auch außerhalb eines Duopols zu ermöglichen. Insbesondere ein statischer Algorithmus agiert systematischer und vorhersehbarer, als es menschliche Entscheider tun. Hierdurch könnte Kollusion befördert werden. Zugleich dürfte es menschlichen Entscheidern schwerer fallen, Vertrauen zu algorithmischen Wettbewerbern aufzubauen, wodurch eine Kollusion wiederum erschwert werden kann.

Für die Untersuchung der Marktinteraktion menschlicher und algorithmischer Anbieter betrachten wir einen Markt mit drei Unternehmen im Wettbewerb. Bisherige Untersuchungen haben gezeigt, dass menschliche Marktteilnehmer im Duopol kollusive Gleichgewichte erzielen können. Bei größeren Märkten mit mehr Wettbewerbern haben sie jedoch erhebliche Schwierigkeiten, sich auf höhere Marktpreise zu koordinieren.⁷⁹³ Ein Markt mit drei Wettbewerbern sollte daher hinreichend konzentriert sein, um *tacit collusion* überhaupt möglich zu machen und zugleich gering genug konzentriert sein, um mögliche Effekte durch die Interaktion mit algorithmischen Wettbewerbern zu untersuchen.

793 Vgl. Kapitel C. III. 2. a); für ein drei-Personen *prisoner's dilemma* zeigen Gerald Marwell und David R. Schmitt, dass diese bereits deutlich weniger kooperativ sind, als entsprechende Experimente mit zwei Spielern, Marwell/Schmitt (1972), JPSP 21 (3), 376.

I. Der Aufbau des Experiments

1. Das grundlegende Marktdesign

Dem vorgestellten Experiment liegt eine Erweiterung des *prisoners' dilemmas* mit $n = 3$ Spielern zugrunde. Zur Identifizierung eines kollusiven Effekts bietet sich ein Tripol-Markt an, da vergangene Experimente zu *tacit collusion* aufgezeigt haben, dass menschliche Teilnehmer in einem Duopol-Markt kollusive Gleichgewichte erzielen können, in Märkten mit mehr Unternehmen aber nur äußerst selten *tacit collusion* erreicht wird.⁷⁹⁴ Die Teilnehmer des Experimentes werden zufällig einem Markt zugeordnet und repräsentieren ein Unternehmen, welches in jeder Periode eine Preisentscheidung über ein von dem Unternehmen produziertes Gut zu treffen hat.⁷⁹⁵ Es handelt sich somit um einen *Bertrand*-Markt. Dieser Markt ist wie folgt modelliert: Jedes Unternehmen kann in jeder Periode zwischen einem hohen Preis von 100 ECU (*Experimental Currency Unit*)⁷⁹⁶ sowie einem niedrigen Preis von 60 ECU für das von ihm angebotene Gut wählen. Die Güter der Wettbewerber sind perfekte Substitute und es wird angenommen, dass keine Produktionskosten zu ihrer Herstellung anfallen. In jeder Periode werden 24 Einheiten des homogenen Gutes nachgefragt, wobei die Konsumenten die Einheiten stets zum niedrigsten Preis erwerben und jedes Unternehmen allein die Nachfrage bedienen könnte.

Tabelle 1: Auszahlungsmatrix für das Laborexperiment von Normann und Sternberg (2021).

		Preisentscheidung der Wettbewerber		
		2x hoher Preis	1x hoher, 1x niedriger Preis	2x niedriger Preis
Zu treffende Entscheidung	Hoher Preis (100 ECU)	800 ECU	0 ECU	0 ECU
	Niedriger Preis (60 ECU)	1.440 ECU	720 ECU	480 ECU

Die mögliche Auszahlung eines Teilnehmers in einer Periode ergibt sich entsprechend der *Tabelle 1*. Alle Marktteilnehmer verdienen 800 ECU pro

794 Vgl. Kapitel C. III. 2. a).
795 Normann/Sternberg (2021), S. 5.
796 ECU bezeichnet die Währung innerhalb des Experiments. 1.000 ECU entsprechen umgerechnet einem Euro.

Periode bei einer erfolgreichen Kollusion. Dieser Wert ergibt sich aus den getroffenen Preisentscheidungen: Wählen alle drei Wettbewerber den hohen Preis, teilen sie mit je 100 ECU als „niedrigstem“ Gebot die Nachfrage unter sich auf. Jedes Unternehmen wird sein Gut somit an jeweils einen Drittel der 24 Nachfrager verkaufen und für diese 8 Einheiten je 100 ECU erhalten. Entscheiden sich alle Unternehmen für einen niedrigen Preis, verdienen sie jeweils nur 480 ECU. Wählt ein Unternehmen den niedrigen, die anderen beiden den hohen Preis, verkauft das günstigere Unternehmen 24 Einheiten und erhält 1.440 ECU, die anderen beiden erhalten jeweils 0 ECU.

Das Experiment besteht aus 3 Spielrunden (*Supergames*), innerhalb derer die Teilnehmer über mindestens 20 Perioden wiederholt ihre Preisentscheidungen treffen. Würden die Teilnehmer das Ende des Spiels kennen, könnten sie mit Hilfe der Rückwärtsinduktion für jede Entscheidung das eindeutige teilspielperfekte Gleichgewicht bestimmen, das jeweils dem Bertrand-Gleichgewicht eines *one-shot-games* entspräche.⁷⁹⁷ Aus diesem Grund wird nach Erreichen der 20. Periode zufällig bestimmt, ob weitere Perioden gespielt werden oder das *Supergame* beendet wird. Die Wahrscheinlichkeit für das Ausführen weiterer Perioden beträgt jeweils 70%. Am Ende eines *Supergames* werden alle Teilnehmer zufällig neuen Mitspielern zugeteilt und das nächste *Supergame* beginnt.

2. Die vier verschiedenen Treatments

Das Experiment teilt sich in vier verschiedene *Treatments* auf, wobei zwei Dimension variiert werden (siehe *Tabelle 2*). Die erste Dimension bezieht sich auf die Zusammensetzung der Märkte und teilt die *Treatments* in *Human-* und *Algorithm-Treatments* auf.⁷⁹⁸ In den *Human-Treatments* treffen Teilnehmer aufeinander, die jeweils eigenständig bestimmen, welche Preise sie für ihr Unternehmen festlegen wollen. In den *Algorithm-Treatments* wird einer der Teilnehmer durch einen Algorithmus unterstützt, welcher die Preisentscheidung für das Unternehmen übernimmt. Der Gewinn, der sich aus den Entscheidungen des Algorithmus ergibt, verbleibt bei dem Teilnehmer. Die anderen Teilnehmer treffen ihre Entscheidungen weiterhin eigenständig.

⁷⁹⁷ Siehe hierzu Kapitel C. II. 2. c).

⁷⁹⁸ *Normann/Sternberg* (2021), S. 5.

Tabelle 2: *Treatments aus dem Laborexperiment von Normann und Sternberg (2021).*

	3 Teilnehmer	2 Teilnehmer, 1 Algorithmus
Unsicherheit	<i>Human Uncertain</i>	<i>Algorithm Uncertain</i>
Volle Information	<i>Human Certain</i>	<i>Algorithm Certain</i>

Die zweite Dimension betrifft die Marktinformationen und teilt die *Treatments* in *Certain*- und *Uncertain-Treatments* auf.⁷⁹⁹ In den *Certain-Treatments* sind die Teilnehmer über die Zusammensetzung der Märkte informiert, also wissen sie, ob eines der Unternehmen mit einem Algorithmus ausgestattet ist oder nicht. In den *Uncertain-Treatments* wissen die Teilnehmer nur, dass die Wahrscheinlichkeit, mit der eines der Unternehmen durch einen Algorithmus unterstützt wird, bei 50% liegt.⁸⁰⁰ In keinem der *Treatments* werden die Teilnehmer über die Funktionsweise des Algorithmus informiert.

3. Die Wahl des Algorithmus

In diesem Experiment wird ein statischer Algorithmus verwendet. Er verfolgt die Strategie des sogenannten *proportional tit-for-tats* (pTFT).⁸⁰¹

a) Vorteile statischer Algorithmen gegenüber selbstlernenden Algorithmen

Im Wettbewerb mit menschlichen Marktteilnehmern bieten sich statische Algorithmen für die Erzielung einer *tacit collusion* an.⁸⁰² Die Wettbewerbs-

799 Normann/Sternberg (2021), S. 6.
800 Diese Angabe entspricht auch der tatsächlichen Wahrscheinlichkeit an einem *Algorithm-Uncertain*- beziehungsweise an einem *Human-Uncertain-Treatment* teilzunehmen. Dieser Aufbau entspricht Farjam/Kirchkamp (2018), JEBO 146, 248.
801 Normann/Sternberg (2021), S. 6.
802 Die präsentierten Simulationen mit selbstlernenden Algorithmen deuten darauf hin, dass viele der bisher verwendeten Algorithmen in heterogenen Märkten Schwierigkeiten haben, kollusive Strategie zu erlernen und durchzusetzen. Hettich zeigt, dass zwei DQN-Algorithmen einen *Q-learning* Algorithmus im Wettbewerb ausbeuten und nur „unter sich“ kollusive Ergebnisse erzielen. Kastius und Schlosser zeigen, dass ein *deep reinforcement learning* Algorithmus mittels einer einfachen Preisregel in eine Kollusion „gezwungen“ wird, indem so lange der wettbewerbliche Preis festgelegt wird, bis der Algorithmus in die Kollusion „einwilligt“, zwei unterschiedliche

entscheidungen der Marktteilnehmer sind vor allem bei ausbleibenden Absprachen mit finanziellen Risiken verbunden. Entscheiden sie sich für einen kollusiven Preis, drohen sie ausgebeutet zu werden. Aus diesem Grund spielt bei der Preisentscheidung die Einschätzung des Verhaltens der Wettbewerber eine gewichtige Rolle. Nur, wenn ein menschlicher Entscheider zu der Überzeugung gelangt, dass seine Wettbewerber ein Interesse an einer Kollusion haben, scheint die Erzielung einer *tacit collusion* möglich. Das systematische und beständige Verfolgen einer Strategie unterscheidet statische Algorithmen von *on-the-job* lernenden Algorithmen sowie menschlichen Marktteilnehmern. Die Strategien sind in der Regel nicht überkomplex und können über die Zeit konstant beibehalten werden. Dies erleichtert es, menschlichen Wettbewerbern die Bereitschaft zur Kollusion zu signalisieren. Indem die Algorithmen systematischer handeln, kann die Risikoeinschätzung der Wettbewerber erleichtert und der Anreiz zur kollusiven Preissetzung erhöht werden. Hierdurch könnten statische Algorithmen die Kollusion mit menschlichen Wettbewerbern auf der kognitiven Ebene erleichtern.

Dementsprechend kommen viele Untersuchungen zu dem Ergebnis, dass insbesondere simple und konsistente Preissetzungsstrategien ein erhebliches Gefährdungspotenzial für den Wettbewerb bieten. Die Britische Wettbewerbsbehörde geht davon aus, dass die Gefahr einer *tacit collusion* vor allem dann hoch zu sein scheint, „wenn die Preissetzungsalgorithmen die Unternehmen dazu veranlassen, sehr einfaches, transparentes und vorhersehbares Preisverhalten anzunehmen.“⁸⁰³ *Wieting* und *Sapi* leiten ebenfalls aus ihren Beobachtungen ab, dass „ein Geheimnis erfolgreicher Absprachen [...] in der Fähigkeit der Manager liegen [könnte], sich auf einfache Strategien festzulegen.“⁸⁰⁴ Auch *David P. Byrne* und *Nicolas de Roos* kommen bei einer

deep learning Algorithmen entwickeln im Wettbewerb gegeneinander jedoch stets neue Strategien und erreichen kein kollusives Gleichgewicht. Die Ergebnisse von *Eschenbaum et al.* zeigen darüber hinaus Schwierigkeiten von Q-learning Algorithmen auf, die kollusiven Ergebnisse vom Trainingsumfeld in neue und unbekannte Marktumgebungen zu übertragen, siehe *Hettich* (2021); *Kastius/Schlosser* (2022), *J Rev Pricing Man* 21, 50; *Eschenbaum et al.* (2022).

803 Aus dem Englischen übersetzt, siehe *British Competition and Markets Authority*, Pricing Algorithms: Economic Working Paper on the Use of Algorithms to Facilitate Collusion and Personalised Pricing, 2018, Rn. 28; auch die Experimente von *Jacob W. Crandall et al.* deuten darauf hin, dass insbesondere „nicht-triviale, aber letztlich einfache algorithmische Mechanismen“ Kollusion in hybriden Zusammensetzungen befördern könnten (Zitat aus dem Englischen übersetzt); *Crandall et al.* (2018), *Nature Communications* 9 (233).

804 Aus dem Englischen übersetzt, siehe *Wieting/Sapi* (2021), S. 41.

Analyse des westaustralischen Tankstellenmarktes zu dem Ergebnis, dass „Unternehmen selbst bei perfekter Preisüberwachung einfache Preisstrukturen annehmen, weil sie leicht auszuprobieren und den Konkurrenten mitzuteilen sind.“⁸⁰⁵ Demnach erhöhe die Einfachheit einer Strategie die Transparenz und verringere die Fehlkommunikation.⁸⁰⁶ Ebenso kommt Musolff zu dem Ergebnis, dass „die Delegation der Preisbildung an einfache Algorithmen *tacit collusion* erleichtern kann, indem die Menge der verfügbaren Strategien reduziert wird.“⁸⁰⁷ Auch Eschenbaum et al. schlagen zur Erzielung überwettbewerblicher Gleichgewichte eine Beschränkung der Strategieoptionen der Algorithmen vor.⁸⁰⁸

Hinzu kommt, dass statische Algorithmen auf realen Online-Märkten derzeit weit verbreitet zu sein scheinen. Wieting und Sapi gehen davon aus, dass auf der Plattform *bol.com* überwiegend relativ einfache statische Algorithmen eingesetzt werden.⁸⁰⁹ Ebenso zeigt Musolff ihre Verbreitung bei seiner Auswertung des Preiswettbewerbs auf der Plattform *Amazon Marketplace* auf.⁸¹⁰ Auch die Monopolkommission kommt in ihrem Gutachten aus 2018 zu dem Schluss, dass „zur Preisgestaltung [derzeit] statische und seltener [selbstlernende][...] Algorithmen eingesetzt [werden].“⁸¹¹

aa) Die Strategie des *proportional tit-for-tats*

Das pTFT entspricht der Erweiterung der sogenannten *tit-for-tat* Strategie für mehrere Spieler.⁸¹² Die in den *Axelrod*-Turnieren⁸¹³ erfolgreiche *tit-for-tat* Strategie für zwei Spieler nutzt als Entscheidungsgrundlage die Entscheidungen der Wettbewerber aus der vorangegangenen Periode (*memory-one*)⁸¹⁴

805 Byrne/De Roos (2019), AER 109 (2), 591 (617).

806 Byrne/De Roos (2019), AER 109 (2), 591 (617).

807 Aus dem Englischen übersetzt, siehe Musolff (2021), S. 2.

808 Eschenbaum et al. (2022), S. 1.

809 Wieting/Sapi (2021), S. 40 f.

810 Musolff (2021), S. 10 f.

811 Monopolkommission, Wettbewerb 2018, Rn. 170.

812 Hilbe et al. (2015), Journal of Theoretical Biology 374, 115.

813 Axelrod, The Evolution of Cooperation, S. 27 ff.

814 Wie zuvor gezeigt hat werden *memory-one* Strategien auch bei einem Großteil der bisher getesteten *Q-learning* Algorithmen angewendet, siehe Kapitel D. II. 2. Calvano et al. zeigen, dass die Gewinne ihrer Algorithmen bei einer Berücksichtigung der letzten zwei Perioden (*memory-two*) deutlich geringeren ausfallen würden, Calvano et al. (2020), AER 110 (10), 3267 (Appendix A5.8).

und entspricht einem „wie du mir, so ich dir“.⁸¹⁵ Die Möglichkeit der Marktteilnehmer, ihre Preisentscheidung von den Entscheidungen der jeweiligen Wettbewerber abhängig zu machen, beruht auf der Annahme der gesteigerten Transparenz digitaler Märkte. Bereits 2016 gaben zwei Drittel der von der Kommission befragten Online-Händler an, die Preise der Wettbewerber mit Hilfe von Preisalgorithmen zu überwachen.⁸¹⁶ Die Beobachtung der Konkurrenten ist im digitalen Handel deutlich vereinfacht, wodurch Informationen über das vorangegangene Verhalten der Wettbewerber bei der Preisentscheidung berücksichtigt werden können.

Vorteile der *tit-for-tat* Strategie sind, dass sie zum einen eine Kollusion befördert, indem sie nicht einseitig von einem kooperativen Gleichgewicht abweicht. Zugleich lässt sie sich jedoch auch nicht ausbeuten, da sie auf eine Abweichung mit einer Bestrafung reagiert. Ist der Wettbewerber (versehentlich) von einem kollusiven Gleichgewicht abgewichen, kehrt sie in der Folge wieder zu einem kollusiven Gleichgewicht zurück, sofern der Wettbewerber ebenfalls zu einem hohen Preis zurückgekehrt ist.

Im vorgestellten Experiment verwendet der Algorithmus die Strategie des pTFTs. Diese Strategie funktioniert wie folgt: In der ersten Periode ($t = 1$) jedes *Supergames* beginnt der Algorithmus stets mit der Wahl des hohen Preises (100 ECU). Diese Entscheidung behält er bei, solange die beiden Wettbewerber ebenfalls den hohen Preis gewählt haben. Weichen beide Konkurrenten hingegen ab, wählt der Algorithmus in der folgenden Periode den niedrigen Preis (60 ECU). Wählt einer der Wettbewerber den hohen und der andere Wettbewerber den niedrigen Preis, wird der Algorithmus in der kommenden Periode mit gleicher Wahrscheinlichkeit (50%) den hohen beziehungsweise den niedrigen Preis auswählen. Somit hängt die Wahrscheinlichkeit dafür, dass der Algorithmus den hohen Preis (p_h) auswählt, davon ab, wie viele seiner Wettbewerber zuvor ebenfalls den hohen Preis gewählt haben [$j \in (0, 1, 2)$]:

$$\text{Wahrscheinlichkeit } (p = p_h) = \begin{cases} 1 & \text{für } t = 1 \\ \frac{j}{2} & \text{für } t > 1 \end{cases}.$$

815 Vgl. Pindyck/Rubinfeld, Mikroökonomie, S. 637, die die Strategie mit dem Ansatz des "Auge um Auge, Zahn um Zahn" umschreiben.

816 Europäische Kommission, Abschlussbericht über die Sektoruntersuchung zum elektronischen Handel, Rn. 13; siehe auch Kapitel A. III. 3. a).

II. Das theoretische Modell

1. Grundlagen des Modells

Der vorgestellte experimentelle Aufbau lässt sich als ein unendlich wiederholter ($t = 0, \dots, \infty$) drei Personen *Bertrand*-Wettbewerb modellieren. Der betrachtete Spieler (i) und seine Wettbewerber (j, k) können hierbei zwischen einem hohen sowie einem tiefen Preis $\{p_h, p_l\}$ wählen. Zukünftige Gewinne werden mit dem Faktor δ diskontiert. In jeder Periode dieses Spiels ergeben sich für Spieler i folgende Auszahlungsoptionen $[\pi(p_i, p_j, p_k)]$, welche in Abhängigkeit zu seiner Entscheidung sowie den Entscheidungen der Konkurrenten j und k stehen:⁸¹⁷

$$\pi^c = \pi(p_h, p_h, p_h) = 800$$

$$\pi^s = \pi(p_h, p_l, p_l) = \pi(p_h, p_h, p_l) = \pi(p_h, p_l, p_h) = 0$$

$$\pi^d = \pi(p_l, p_h, p_h) = 1440$$

$$\pi^f = \pi(p_l, p_l, p_h) = \pi(p_l, p_h, p_l) = 720$$

$$\pi^n = \pi(p_l, p_l, p_l) = 480.$$

In Anlehnung an das vorgestellte Kollusionsmodell⁸¹⁸ gilt es zu prüfen, unter welchen Bedingungen der Spieler i sich dazu entscheidet, von einer Koordination auf den hohen Preis (π^c) abzuweichen. Es ist zunächst davon auszugehen, dass die beiden Wettbewerber so lange die Kollusion aufrechterhalten, wie Spieler i ebenfalls den hohen Preis wählt. Weicht i von dieser Strategie ab (π^a), so kehren die Wettbewerber für alle folgenden Perioden zum wettbewerblichen Gleichgewicht zurück (π^w). Diese kooperative aber Abweichungen nicht verzeihende Strategie der Wettbewerber entspricht der *grim-trigger* Strategie (GT). Für Spieler i ist die Wahl des hohen Preises ein teilspielperfektes Gleichgewicht, wenn der langfristige Gewinn durch die Kollusion größer ist, als der Gewinn, den der Spieler durch eine Abweichung erhalten würde:

$$\pi^c + \pi^c * \frac{\delta}{1 - \delta} \geq \pi^d + \pi^n * \frac{\delta}{1 - \delta}$$

$$\delta \geq \frac{2}{3} \approx 0,667 \equiv \underline{\delta}_{GT}.$$

817 Die Werte entsprechen den möglichen Auszahlungen der *Tabelle 1*.

818 Siehe Kapitel C. II. 2. d) aa).

Wird einer der beiden Wettbewerber durch den Algorithmus ersetzt, verändert sich die Situation dahingehend, dass der Algorithmus des Konkurrenten anstatt GT die Strategie pTFT spielt. Entscheidet sich Spieler i dazu, von der Kollusion abzuweichen, wird der Algorithmus in der folgenden Periode mit jeweils 50-prozentiger Wahrscheinlichkeit den hohen beziehungsweise den niedrigen Preis wählen. Der andere Wettbewerber wird aufgrund seiner GT-Strategie sofort auf den niedrigen Preis wechseln, sodass Spieler i spätestens ab der darauffolgenden Periode π^n erhalten wird. Das Aufrechterhalten der Kollusion ist im Wettbewerb, bei dem für einen Gegenspieler ein pTFT-Algorithmus entscheidet, nur dann ein teilspielperfektes Gleichgewicht, sofern:

$$\pi^c + \pi^c * \frac{\delta}{1-\delta} \geq \pi^d + \frac{\pi^f}{2} + \frac{\pi^n}{2} + \pi^n * \frac{\delta^2}{1-\delta}$$

$$\delta \gtrsim 0,69 \equiv \delta_{\text{pTFT}}.$$

Vergleicht man δ_{GT} und δ_{pTFT} , so wird ersichtlich, dass der kritische Diskontierungsfaktor für zwei GT-Spieler und einen pTFT-Spieler (δ_{pTFT}) größer ist als bei drei GT-Spielern. Das Ergebnis ist nachvollziehbar, da pTFT im Gegensatz zu GT nachsichtiger ist und trotz Abweichung zurück zur Kooperation finden kann, was zugleich eine Abweichung attraktiver erscheinen lässt.

2. Strategische Unsicherheit

Im Kontrast zu den zuvor dargestellten Szenarien hat Spieler i regelmäßig keine Informationen über die gewählte Strategie seiner Gegenspieler. Im wiederholten Gefangenendilemma lässt sich hierfür das Konzept der Risikodominanz von *John C. Harsanyi* und *Reinhard Selten* anwenden,⁸¹⁹ welches sich mit der Entscheidungsfindung bei strategischer Unsicherheit befasst.⁸²⁰ Mit dem Konzept des geringsten Risikos wählen rationale Spieler zwischen mehreren Gleichgewichten eines Spiels.⁸²¹ Eine Strategie ist dabei

819 Harsanyi/Selten, A General Theory of Equilibrium Selection in Games, S. 82 ff.

820 Vgl. auch Blonski et al. (2011), AEJ: Micro 3 (3), 164; Blonski/Spagnolo (2015), IJGT 44, 61; Dal Bó/Fréchette (2011), AER 101 (1), 411; Dal Bó/Fréchette (2018), JEL 56 (1), 60; Green et al., in: Blair/Sokol (Hrsg.), Handbook of International Antitrust Economics, S. 464 (486 ff.).

821 Green et al., in: Blair/Sokol (Hrsg.), Handbook of International Antitrust Economics, S. 464 (487 f.).

risikodominant, sofern sie die beste Antwort auf einen anderen Spieler ist, der mit gleicher Wahrscheinlichkeit zwischen zwei verschiedenen Strategien wählt.⁸²² Hierbei zeigt sich ein Vorteil der automatisierten Preissetzung gegenüber menschlichen Entscheidern. Indem das Unternehmen mittels des Algorithmus eine einfache und konsistente Strategie verfolgt, ist es für die Wettbewerber leichter, diese zu entschlüsseln. Darüber hinaus bindet der Algorithmus das Unternehmen als *commitment device* in gewissem Umfang an seine Strategie, wodurch die Unsicherheit der Wettbewerber reduziert wird.⁸²³

Zur Vereinfachung wird im Folgenden davon ausgegangen, dass die menschlichen Spieler entweder zunächst kooperativ oder durchgehend kompetitive Spieler sind. Aus diesem Grund wählen sie zwischen zwei Strategien, dem kooperativen GT und einem kompetitiven dauerhaften Abweichen, dem sogenannten *always defect* (AD).⁸²⁴ Bei AD verhält sich der Spieler unabhängig von seinen Gegenspielern und wählt stets den niedrigen Preis. Die Konstellationen, in denen alle drei Spieler AD, bzw. GT, spielen stellen ein Gleichgewicht dar. Da im Gegensatz zu AD im Fall von GT zukünftige Gewinne aufgrund einer Kollusion in Betracht kommen, nimmt das Risiko der Strategie ab, je höher der Diskontierungsfaktor ist. GT ist somit die beste Antwort auf die unsichere Strategie der Gegenspieler und AD vorzuziehen, sofern:

$$\frac{\frac{1}{4} * \left(\frac{\pi^c}{1-\delta} \right) + \frac{3}{4} * \left(\pi^s + \frac{\delta \pi^n}{1-\delta} \right)}{\text{Erwarteter Gewinn bei GT}} \geq \frac{\frac{\pi^d}{4} + \frac{\pi^f}{2} + \frac{\pi^n}{4} + \frac{\delta \pi^n}{1-\delta}}{\text{Erwarteter Gewinn bei AD}}$$

$$\delta \geq \frac{8}{9} \approx 0,89 \equiv \delta_{GT}^*$$

822 Normann/Sternberg (2021), Appendix A.1; Harsanyi/Selten, A General Theory of Equilibrium Selection in Games, S. 83.

823 Vgl. Salcedo (2015); Z. Brown/MacKay (2022), AEJ: Micro (im Erscheinen); Dem Unternehmen steht es allerdings frei, den Algorithmus zu ändern oder abzuschalten, sodass sich die Wettbewerber nicht darauf verlassen können, dass ihr Konkurrent für immer an dieser Strategie festhält. Dennoch signalisiert der kollusive Algorithmus die generelle Bereitschaft des Unternehmens zur Kollusion. Die Umstellung des Algorithmus auf eine ebenfalls kollusive GT-Strategie würde das Ergebnis des Modells nicht beeinträchtigen, Normann/Sternberg (2021), Appendix A. 1.

824 So auch Blonski et al. (2011), AEJ: Micro 3 (3), 164; Blonski/Spagnolo (2015), IJGT 44, 61; Dal Bó/Fréchette (2011), AER 101 (1), 411; Dal Bó/Fréchette (2018), JEL 56 (1), 60; Green et al., in: Blair/Sokol (Hrsg.), Handbook of International Antitrust Economics, S. 464 (486 ff.).

Die Variable δ_{GT}^* gibt somit den kritischen Diskontierungsfaktor bei strategischer Unsicherheit dreier Spieler an, die entweder AD oder GT als Strategie wählen. Ist einer der Wettbewerber ein Algorithmus, welcher für einen der Spieler verbindlich pTFT spielt, so ist GT gegenüber AD risikodominant, sofern:⁸²⁵

$$\underbrace{\frac{1}{2} * \left(\frac{\pi^c}{1-\delta} \right) + \frac{1}{2} * \left(\pi^s + \frac{\delta}{2} * (\pi^f + \pi^n) + \frac{\delta^2 \pi^n}{1-\delta} \right)}_{\text{Erwarteter Gewinn bei GT}} \geq \underbrace{\frac{1}{2} * \left(\pi^d + \frac{\delta}{2} * (\pi^f + \pi^n) + \frac{\delta^2 \pi^n}{1-\delta} \right) + \frac{1}{2} * \left(\pi^f + \frac{\delta \pi^n}{1-\delta} \right)}_{\text{Erwarteter Gewinn bei AD}}$$

$$\delta \geq \frac{17}{21} \approx 0,81 \equiv \delta_{pTFT}^*.$$

Unter Hinzunahme strategischer Unsicherheit zeigt sich, dass der kritische Diskontierungsfaktor bei drei potenziellen GT Spielern größer ist, als bei zwei potenziellen GT Spielern und einem pTFT Algorithmus ($\delta_{GT}^* > \delta_{pTFT}^*$). Somit kann das strategische Risiko durch den Einsatz eines kollusiven Algorithmus reduziert werden.⁸²⁶

Das vorangegangene Ergebnis bedient sich der Annahme, dass die Strategie des Algorithmus ersichtlich oder zumindest vermutet wird. Angesicht der wiederholten Interaktionen in einem Supergame, der mehrfachen Wiederholung der *Supergames* sowie der relativ klaren Strategie des Algorithmus erscheint es naheliegend, dass die Spieler ein kooperatives Verhalten des Algorithmus über die Zeit feststellen. Allerdings können zu Beginn des Spiels unterschiedliche Vermutungen über die Strategie des Algorithmus eine Rolle spielen und die Spieler beeinflussen. Hierfür ließe sich die Variable γ ($\gamma \in (0, 1]$) für die Einschätzungen der Spieler ergänzen, welche die Wahrscheinlichkeit angibt, dass sie davon ausgehen, dass der Algorithmus eine kooperative Strategie spielt (hier pTFT)⁸²⁷ spielt. Im Umkehrschluss gibt $1 - \gamma$ die Wahrscheinlichkeit an, mit der der Spieler den Algorithmus kompetitiv (AD) einschätzt. Der kritische Diskontierungsfaktor stiege dann

825 Normann/Sternberg (2021), A. 1.

826 Wenngleich der kritische Diskontierungsfaktor ohne strategische Unsicherheit zu einem gegensätzlichen Ergebnis gelangt ($\delta_{pTFT} > \delta_{GT}$), zeigen experimentellen Ergebnisse auf, dass δ^* eine größere Aussagekraft als δ beizumessen ist, Dal Bó/Fréchette (2018), JEL 56 (1), 60.

827 Eine fälschliche Annahme der Teilnehmer, der Algorithmus würde GT spielen, führt zum gleichen Ergebnis.

mit der Abnahme von γ an. Die Wahl der Strategie GT wird umso attraktiver, je eher der Spieler davon ausgeht, dass der Algorithmus kooperativ spielt:⁸²⁸

$$\frac{\partial \delta_{pTFT}^*(\gamma)}{\partial \gamma} < 0.$$

III. Hypothesen

Im Folgenden werden Forschungshypothese aufgestellt, welche sich im Rahmen des Experiments testen lassen. Indem zwischen den *Treatments* eine unabhängige Variable variiert wird, lässt sich auf Grundlage des Modells prognostizieren, welchen Effekt die Veränderung auf die abhängige Variable haben sollte. Die abhängige Variable ist das kollusive Verhalten der Marktteilnehmer und somit der durchschnittlich erzielte Marktpreis. Die vorangegangene Analyse des wiederholten Spiels deutet darauf hin, dass ein Algorithmus das Marktgeschehen auf zwei verschiedenen Ebenen beeinflussen kann. Demnach scheint sowohl das Verhalten des Algorithmus, als auch die Vorstellungen der Teilnehmer über das Verhalten des Algorithmus einen Einfluss auf die Wahl der Strategien der Teilnehmer zu haben.

1. Das Verhalten des Algorithmus

Das Modell hat gezeigt, dass kollusives Verhalten bei Anwesenheit eines pTFT-Algorithmus ein teilspielperfektes Gleichgewicht darstellen kann. Indem der Algorithmus das Unternehmen auf eine konsequente Strategie verpflichtet, reduziert er das strategische Risiko des Wettbewerbs. Darüber hinaus belohnt der pTFT Algorithmus kooperatives Verhalten und bestraft Abweichungen. Insbesondere über die Zeit hinweg sollte der Algorithmus so die Preissetzung der Teilnehmer beeinflussen und höhere Preise ermöglichen. Betrachtet man ausschließlich das Verhalten der Algorithmen im Vergleich zu menschlichen Entscheidern, so ist – entsprechend dem Modell ($\delta_{GT}^* > \delta_{pTFT}^*$) – davon auszugehen, dass die *Algorithm-Treatments* zu einer höheren Kollusion gelangen müssten als die ihnen entsprechenden *Human-Treatments*.

⁸²⁸ Eine schrittweise Lösung für δ_{pTFT}^* findet sich im Appendix zu Normann/Sternberg (2021).

2. Vorstellungen über das Verhalten des Algorithmus

Darüber hinaus könnten aber auch Vorstellungen über das Verhalten des Algorithmus eine Rolle bei der Wahl der Strategie eines Teilnehmers spielen. Das Modell hat prognostiziert, dass *tacit collusion* umso wahrscheinlicher ist, je weniger kompetitiv das Verhalten des Algorithmus erwartet wird. Sofern ein Spieler davon ausgeht, dass ein Gegenspieler ihn ausbeuten möchte, wird er sich eher nicht für einen kollusiven Preis entscheiden. So scheint das Vertrauen der Spieler in ihr Gegenüber und seine Motivation eine wichtige Komponente für eine erfolgreiche Kollusion darzustellen. In diesem Zusammenhang zeigen diverse Experimente, dass Spieler in (*n*-Personen) Gefangenendilemmata weniger kooperieren, sofern sie auch von ihrem Gegenüber weniger Kooperation erwarten.⁸²⁹ Zwei Motive scheinen maßgeblich dafür verantwortlich zu sein, dass Spieler von einer kollusiven Strategie abweichen: Zum einen das Motiv, selbst auszubeuten in der Erwartung eines kooperativen Gegenübers. Zum anderen die Angst, von einem kompetitiven Wettbewerber ausgebeutet zu werden.⁸³⁰

Im vorgestellten Experiment erhalten die Teilnehmer keine Informationen über das Verhalten des Algorithmus. Aus diesem Grund könnten die Vorstellungen der Teilnehmer besonders zu Beginn des Experiments einen erheblichen Einfluss auf ihre Entscheidung haben, da es ihnen (noch) an Erfahrungen mangelt. Obwohl die Strategie des Algorithmus ebenso unbekannt ist, wie die Strategie der menschlichen Wettbewerber, deutet einiges darauf hin, dass Menschen das Verhalten von Algorithmen grundsätzlich anders einschätzen als das Verhalten menschlicher Gegenspieler.

Die Arbeiten von *Berkeley J. Dietvorst et al.* zeigen, dass Menschen für eine Vorhersage selbst dann menschliche Expertise einem Algorithmus vorziehen, wenn sie wissen, dass der Algorithmus dem Menschen überlegen ist.⁸³¹ Dieses Phänomen des Misstrauens wird von den Autoren als Aversion gegen Algorithmen beschrieben (*algorithm aversion*). Hinzu kommt, dass Algorithmen in der Vergangenheit vor allem dadurch aufgefallen sind, dass sie professionellen menschlichen Gegenspielern in Spielen wie Poker⁸³²,

829 Siehe die Meta-Studie von *Balliet/van Lange*, *Psychological bulletin* 139 (5) (2013), 1090.

830 Vgl. *Ahn et al.* (2001), *Public Choice* 106, 137; *Blanco et al.* (2014), *GEB* 87, 122; *Charness et al.* (2016), *GEB* 100, 113.

831 *Dietvorst et al.* (2015), *Journal of Experimental Psychology: General* 144 (1), 114; *Dietvorst/Bharti* (2020), *Psychological Science* 31 (10), 1302.

832 *N. Brown/Sandholm*, *Science* 359 (6374) (2018), 418.

Schach⁸³³, Dame⁸³⁴ oder Go⁸³⁵ deutlich überlegen waren. Diese Resultate lassen es nicht unwahrscheinlich erscheinen, dass ein Algorithmus auch dazu in der Lage ist, einen menschlichen Teilnehmer in einem Marktspiel zu „besiegen“, indem er höhere Gewinne auf Kosten seiner Wettbewerber einfährt. Mithin ist davon auszugehen, dass die Teilnehmer der Kooperationsbereitschaft der Algorithmen eher skeptisch entgegenstehen und dadurch – entsprechend dem Modell – Kollusion weniger wahrscheinlich wird. Die Treatments dürften somit umso weniger Kollusion erzielen, je höher die Wahrscheinlichkeit ist, dass einer der Wettbewerber einen Algorithmus verwendet.

3. Treatment Vergleich

Betrachtet man die *Algorithm-* und *Human-Treatments* jeweils untereinander, so spielt das tatsächliche Verhalten des Algorithmus keine Rolle. Demnach unterscheiden sich die *Treatments* ausschließlich in Hinblick auf das vermutete Verhalten des Algorithmus. Während im *Human-Uncertain-Treatment* Unsicherheit bezüglich des Einsatzes algorithmischer Preissetzung besteht, wissen die Teilnehmer im *Human-Certain-Treatment*, dass ausschließlich menschliche Wettbewerber Entscheidungen treffen – hier sollte die Gewissheit über menschliche Konkurrenten einen positiven Effekt auf die Kollusionsrate haben. In den *Algorithm-Treatments* ist der Effekt entgegengesetzt. Hier sollte die Gewissheit über algorithmische Preissetzung dazu führen, dass das *Certain-Treatment* eine geringere Kooperationsrate erreicht:

Hypothese 1.1: Human-Certain > Human-Uncertain

Hypothese 1.2: Algorithm-Uncertain > Algorithm-Certain.

Betrachtet man die *Uncertain-Treatments*, so ist die Wahrscheinlichkeit bezüglich der Verwendung algorithmischer Preissetzung sowohl im *Human-*, als auch im *Algorithm-Treatment* gleich. Die Vorstellungen über die Strategie des Algorithmus sollte somit keinen Einfluss auf die Entscheidung der Versuchsteilnehmer haben. Allerdings unterscheiden sich die *Treatments* bezüglich der tatsächlichen Anwesenheit eines Algorithmus, sodass das

833 M. Campbell et al. (2002), Artificial Intelligence 134 (1-2), 57; Silver et al., Nature 529 (7587) (2016), 484.

834 Schaeffer (1997), ICG 20 (2), 93.

835 Silver et al., Nature 529 (7587) (2016), 484.

tatsächliche Verhalten einen Einfluss auf das Marktergebnis haben sollte. Entsprechend der zuvor getroffenen Einschätzung ist davon auszugehen, dass der Algorithmus dazu führt, dass die Kooperationsraten in diesem *Treatment* höher ausfallen:

Hypothese 2: Algorithm-Uncertain > Human-Uncertain.

Vergleicht man das *Algorithm-Certain-Treatment*, mit dem *Human-Certain-Treatment*, so stehen sich hierbei zwei gegenläufige Effekte entgegen. Während das Verhalten des Algorithmus zu einer höheren Kooperationsrate führen sollte, könnten die Vorstellungen der Teilnehmerinnen bezüglich des vermuteten Verhaltens des Algorithmus einen negativen Effekt auf die Kooperation haben. Aufgrund der vielen Perioden und mehrfachen *Supergames* ist allerdings davon auszugehen, dass die Skepsis gegenüber dem Algorithmus aufgrund seines tatsächlich kooperativ geneigten Verhaltens abnimmt, sodass γ über die *Supergames* hinweg zunehmen sollte. Mit der Zeit dürfte der positive Effekt des tatsächlichen Verhaltens überwiegen, sodass wir insgesamt von einer höheren Kooperationsrate im *Algorithm-Treatment* ausgehen:

Hypothese 2.2: Algorithm-Certain > Human-Certain

Hypothese 2.3: Algorithm-Treatments > Human-Treatments.

Auf Grundlage der Hypothesen ergibt sich eine eindeutige überprüfbare Rangfolge der *Treatments*:

Hypothese 3:

Algorithm-Uncertain > Algorithm-Certain >

Human-Certain > Human-Uncertain.

IV. Durchführung des Experiments

Das Experiment wurde mit der Software *z-Tree*⁸³⁶ programmiert und in den Laboratorien für experimentelle Wirtschaftsforschung in Düsseldorf (*DICE-Lab*) und Bonn (*MPI Decision Lab*) durchgeführt. An den vier *Treatments* nahmen zwischen August 2019 und Oktober 2020 insgesamt 309 Teilnehmer

836 Z-Tree steht für Zurich Toolbox for Readymade Economic Experiments und ist eine von Urs Fischbacher entwickelte Software zur Programmierung und Durchführung von experimenteller Forschung in der Ökonomie, Fischbacher (2007), *Experimental Economics* 10, 171.

teil, wobei es sich hierbei überwiegend um Studierende aus Bonn und Düsseldorf handelte. Die Teilnehmer erhielten eine Aufwandsentschädigung für ihre Teilnahme, sowie eine Auszahlung entsprechend ihrem Gewinn aus einem zufällig ausgewählten Supergame. In jeder *Session* wurden die Teilnehmer zufällig einem Laborplatz mit einem Computer zugeordnet, auf welchem sie die Marktentscheidungen angezeigt bekamen und ihre Entscheidungen zu treffen hatten. Nach jeder Entscheidung wurde den Teilnehmern ihr Gewinn der vergangenen Periode sowie die Entscheidung der übrigen Marktteilnehmer angezeigt. War einer der Teilnehmer mit einem Algorithmus ausgestattet, so musste er selbst keine aktive Entscheidung treffen, nahm allerdings an dem Experiment teil und bekam den Gewinn entsprechend den Entscheidungen des Algorithmus nach jeder Periode angezeigt und am Ende des Experiments ausgezahlt. Zum Verständnis wurden zu Beginn jeder *Session* Instruktionen verteilt und anhand von Kontrollfragen sichergestellt, dass die Teilnehmer das Spiel und seine Regeln verstanden hatten. Teil der Instruktionen, war eine Auszahlungstabelle entsprechend der zuvor dargestellten *Tabelle 1*. Im Anschluss an das Experiment wurden Teilnehmer der *Uncertain-Treatments* darüber hinaus um eine incentivierte Einschätzung gebeten, wie sicher sie sich auf einer Skala von 1-100 sind, dass in der vorangegangenen *Session* je ein Teilnehmer in ihren Märkten durch einen Algorithmus unterstützt wurde. Lagen die Teilnehmer hierbei richtig, wurden ihnen bis zu 2 Euro zusätzlich ausgezahlt. Zum Abschluss folgte ein Fragebogen, welcher unter anderem das Alter, Geschlecht und den aktuellen Studiengang der Teilnehmer abfragte. Die Sessions dauerten im Schnitt 60 Minuten und die Teilnehmer erhielten durchschnittlich 17,73 Euro für ihre Teilnahme.

V. Ergebnisse des Experiments

1. Allgemeine Übersicht

Im Folgenden sollen die Ergebnisse des Experiments analysiert werden. *Abbildung 6* zeigt die Kooperationsraten der *Treatments* im Zeitverlauf.⁸³⁷ Hierbei fällt eine steigende Tendenz über die *Supergames* hinweg auf. Es zeigt sich, dass die Teilnehmer durch die wiederholten Interaktionen kollusives

⁸³⁷ Um die zu beobachteten Effekte zu Beginn und zum Ende eines jeden Supergames (*Restart-* und *Endgame-Effekt*) ausschließen, werden hierbei die Perioden sechs bis neunzehn je Supergame betrachtet.

Verhalten erlernen.⁸³⁸ Im ersten *Supergame* liegen die Kooperationsraten aller *Treatments* nah beieinander und die Märkte weisen eine hohe Wettbewerbsintensität auf. Im zweiten *Supergame* ist ein Anstieg der Kooperationsrate bei allen *Treatments* festzustellen. Im *Algorithm-Uncertain-Treatment* zeigt sich bereits eine relativ stabile und höhere Kooperationsrate über die Perioden hinweg. Die beiden *Certain-Treatments* weisen zu Beginn des *Supergames* eine hohe Kooperationsrate, welche jedoch im Zeitverlauf stark abfällt. Das *Human-Certain-Treatment* befindet sich von Beginn an auf niedrigem Niveau. Vergleicht man das zweite und dritte *Supergame*, so steigen die Kooperationsraten nur noch in den *Treatments* deutlich an, in denen eines der Unternehmen mit einem Algorithmus ausgestattet ist.⁸³⁹ Dadurch zeichnet sich im dritten *Supergame* ein differenziertes Bild ab, bei dem sich die *Algorithm-Treatments* klar von den *Human-Treatments* absetzen.⁸⁴⁰

Bildet man eine Rangfolge der *Treatments* über alle *Supergames* bezüglich der durchschnittlichen Kooperationsrate entspricht diese der Reihenfolge der Hypothese 3.⁸⁴¹ Demnach zeigt sich die höchste Kooperationsrate im *Algorithm-Uncertain-Treatment* (33,7%). Es folgen die *Algorithm-* (28,2%) und *Human-Certain-Treatments* (26,2%), wohingegen das *Human-Uncertain-Treatment* die geringste Kooperationsrate aufweist (18,7%).

838 Vgl. Bigoni et al. (2015), *Econometrica* 83 (2), 587; Dal Bó/Fréchette (2011), *AER* 101 (1), 411; Dal Bó/Fréchette (2018), *JEL* 56 (1), 60; Fudenberg et al. (2012), *AER* 102 (2), 720.

839 Das Treatment *Human-Certain* verzeichnet lediglich einen geringfügigen Anstieg von 0,2%P, während die Kooperationsrate im *Human-Uncertain-Treatment* leicht zurückgeht (-0,3%P).

840 Eine detaillierte Auflistung der Kooperationsraten findet sich im Appendix zu dieser Arbeit (Tabelle 1).

841 Die Rangfolge ist dabei unabhängig von den beobachteten *Restart-* und *Endgame-*Effekten auch für alle Perioden unverändert.

Abbildung 6: Ergebnisse aus dem Laborexperiment von Normann und Sternberg (2021). Dargestellt sind die Kooperationsraten (in %) über die Perioden 6-19 aller Supergames hinweg.

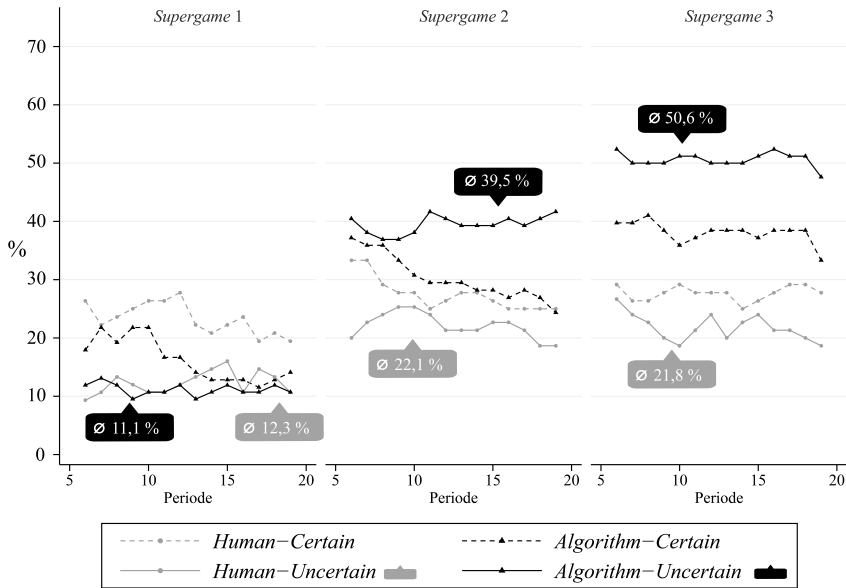
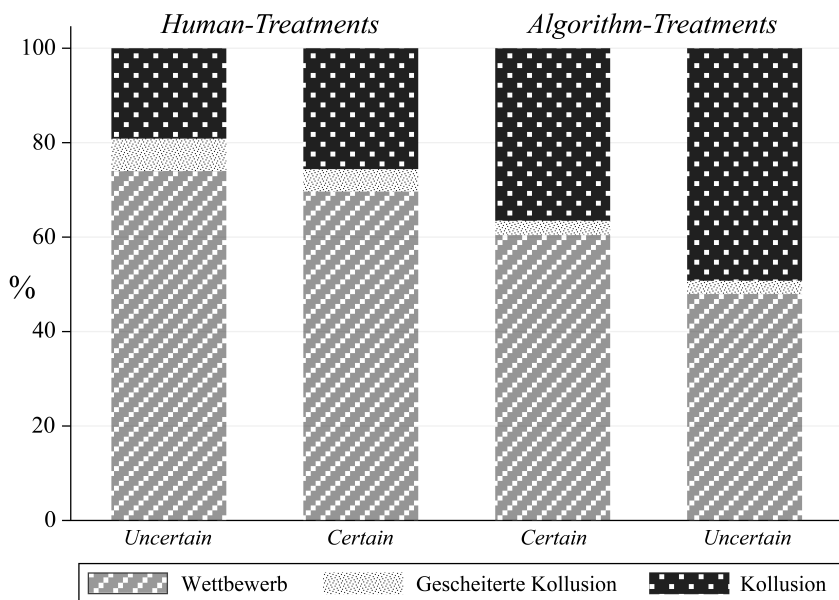


Abbildung 7 zeigt den Anteil erfolgreicher Kollusion in den Märkten des dritten Supergames. Haben alle Teilnehmer den niedrigen Preis gewählt, so entspricht das Marktergebnis einem wettbewerblichen Gleichgewicht (Wettbewerb). Haben alle Teilnehmer hingegen den hohen Preis gewählt, so liegt eine *tacit collusion* vor (Kollusion). Alle Marktergebnisse, bei denen mindestens einer der Teilnehmer einen anderen Preis als seine Wettbewerber gewählt hat, werden als gescheiterte Kollusion geführt.

Das Diagramm zeigt den deutlichen Unterschied zwischen dem Human- und Algorithm-Treatments im dritten Supergame auf. Während es in ungefähr der Hälfte der Entscheidungssituationen im Algorithm-Uncertain-Treatment zu einer erfolgreichen Kollusion kommt (49,23%), überwiegt im Human-Uncertain-Treatment eindeutig das wettbewerbliche Gleichgewicht (74%).

Abbildung 7: Ergebnisse aus dem Laborexperiment von Normann und Sternberg (2021). Dargestellt ist der Anteil der jeweiligen Marktergebnisse (in %) im dritten Supergame (Periode 6-19).



2. Statistische Auswertung

a) Zentrale Ergebnisse

Im Folgenden werden die im Rahmen des Experiments aufgestellten Forschungshypothesen mit Hilfe von statistischen Analysen überprüft. Hierfür wird mit Hilfe eines Regressionsmodells⁸⁴² überprüft, welchen Einfluss die modifizierten unabhängigen Variablen auf die abhängige Variable hatten. Im vorgestellten Experiment ist die abhängige Variable ob ein Unternehmen (Mensch oder Algorithmus) in einem bestimmten Zeitraum kooperiert oder nicht. Hierbei zeigt sich, dass der Algorithmus als erklärende Variable einen signifikanten positiven Einfluss ab dem zweiten *Supergame* ausübt (jeweils p

⁸⁴² Die Ergebnisse dieses linearen Wahrscheinlichkeitsmodells findet sich im Appendix zu dieser Arbeit (Tabelle 2). Im Folgenden werden die wichtigsten Befunde vorgestellt.

$< 0,05$)⁸⁴³. Daraus folgt, dass die beiden *Treatments*, in denen ein Algorithmus anwesend war zu einer höheren Kooperation geführt haben, als die beiden *Human-Treatments* (Hypothese 2.3). Betrachtet man die *Treatments* getrennt voneinander, so zeigt sich, dass *Algorithm-Uncertain* über alle *Supergames* betrachtet zu einem signifikanten Anstieg der Kooperationsrate führt (Hypothese 2.1.), der Unterschied zwischen den *Certain-Treatments* ist hingegen nicht signifikant (Hypothese 2.2).

Darüber hinaus lässt sich der Einfluss der Information über die Marktzusammensetzung überprüfen. Die Hypothese war, dass die Information in den *Human-Treatments* zu einer höheren (Hypothese 1.1), in den *Algorithm-Treatments* jedoch zu einer niedrigeren Kooperationsrate (Hypothese 1.2.). Allerdings zeigt eine statistische Auswertung weder bei dem einen, noch bei dem anderen Vergleich einen signifikanten Effekt. Dies deutet darauf hin, dass die Gewissheit über die Anwesenheit eines Algorithmus keinen ausschlaggebenden Effekt auf die Entscheidungen der Teilnehmer hatte. Lenkt man den Fokus auf die allererste Periode des Experiments, in der Erwartungen den größten Ausschlag geben sollten, zeigt sich ebenfalls kein signifikanter Effekt zwischen den *Treatments*.

Im Einklang mit unseren Annahmen steht hingegen die Rangfolge der einzelnen *Treatments*. Überprüft man ihren Einfluss mit Hilfe einer kardinalen Rangvariable für die *Treatments*, bei der *Algorithm-Uncertain* als am kooperativsten eingeschätztes *Treatment* den niedrigsten Wert (1) und *Human-Uncertain* als kompetitivstes *Treatment* den höchsten Wert (4) erhält, so zeigen die Daten über alle *Supergames* gesehen einen signifikanten negativen Effekt ($p < 0,05$) der Rang-Variable auf die Kooperationsrate (Hypothese 3).⁸⁴⁴

843 Der p -Wert ist ein Instrument zur Überprüfung einer Hypothese und zur Angabe der statistischen Signifikanz eines Ergebnisses. Der p -Wert ist die berechnete Wahrscheinlichkeit, dass das gefundene Ergebnis ein reiner Zufallsfund ist und der Effekt in der hypothetischen „realen“ Umgebung nicht auftritt, obwohl er im Experiment sichtbar war. Als Grenzwert für die Bewertung eines Ergebnisses als *signifikant* gilt $p < 0,05$ demnach läge die Wahrscheinlichkeit eines Zufallsfundes bei unter 5%. Spricht man von einem *stark signifikanten* Ergebnis, liegt die Wahrscheinlichkeit, dass es sich um einen Zufallsfund handelt sogar bei unter einem Prozent ($p < 0,01$), M. T. Harrison, in: Jung/Jäger (Hrsg.), *Encyclopedia of Computational Neuroscience*, S. 2695.

844 Die Regressionstabelle findet sich im Appendix zu dieser Arbeit (Tabelle 3).

b) Weitere Ergebnisse

aa) Vermutungen der Teilnehmer

Ein weiteres interessantes Ergebnis, welches unseren Annahmen entspricht, betrifft die abgefragte Vermutung der Teilnehmer am Ende des Experiments. In den *Uncertain Treatments* hatten alle Teilnehmer, die nicht mit einem Algorithmus ausgestattet waren, die Möglichkeit, zusätzlich bis zu 2 Euro zu verdienen. Die Höhe der Auszahlung hing davon ab, wie zutreffend ihre Einschätzung bezüglich der Anwesenheit eines Algorithmus auf ihren Märkten gewesen ist. Die Teilnehmer wurden gefragt, wie sicher sie sich auf einer Skala von 1-100 sind, dass in Ihrem Experiment einer (keiner) der Wettbewerber einen Algorithmus nutzte. Gaben die Teilnehmer 100 (0) an, so waren sie sicher, dass ein (kein) Algorithmus verwendet wurde. Hierbei zeigte sich, dass die Teilnehmer eher davon ausgingen, dass ein Algorithmus anwesend war, wenn dies gerade nicht der Fall gewesen ist ($p < 0,05$).⁸⁴⁵ Dies könnte damit zusammenhängen, dass die Teilnehmer kompetitives Verhalten einem Algorithmus zuschreiben, wenngleich in diesem Experiment das Gegenteil der Fall ist.

bb) Gewinner der Kollusion

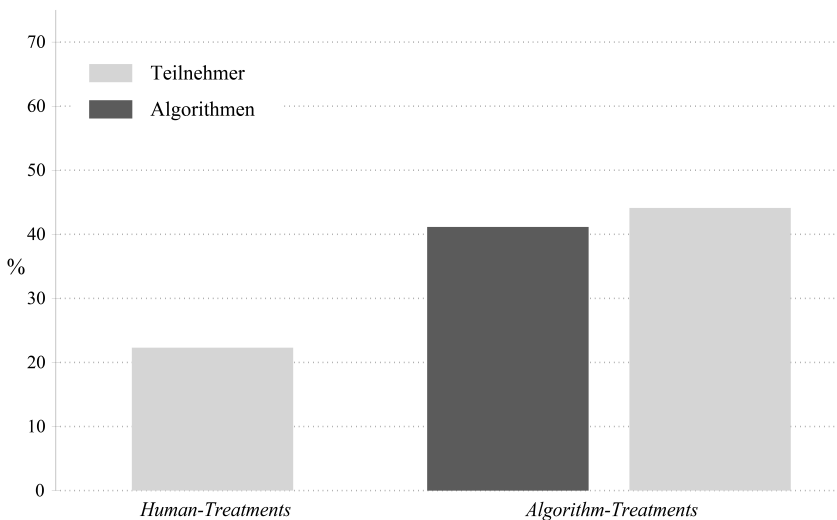
Neben der Kooperationsrate lässt sich auch der Gewinn der Unternehmen als abhängige Variable betrachten. Für das dritte Supergame zeigt *Abbildung 8* exemplarisch die Gewinne im Verhältnis zum statischen *Nash*-Gleichgewicht (0%) sowie einer erfolgreichen Kollusion (100%) und gruppiert in *Human*- und *Algorithm-Treatments*. Wenig überraschend zeigt sich auch hierbei der

845 Zur weiteren Untersuchung dieses Zusammenhangs wurde der Regression eine erklärende Variable hinzugefügt, die Ergebnisse finden sich auch im Appendix zu dieser Arbeit (Tabelle 5). Die Variable gibt an, wie häufig ein Teilnehmer eine gescheiterte Kollusion (mindestens einer der Teilnehmer wählte einen anderen Preis als seine Wettbewerber) erlebt hat. Die Variable hat einen stark signifikanten Effekt ($p < 0,01$) auf die Vermutung am Ende des Experiments. Demnach gehen die Teilnehmer umso mehr davon aus, dass ein Algorithmus anwesend war, je häufiger Kollusion durch Abweichung scheiterte. Die Variable schwächt auch den Einfluss der Algorithmus-Variable auf die Einschätzung der Teilnehmer ab, wenngleich auch sie signifikant bleibt ($p < 0,05$). Mithin stützt dieses Ergebnis die Einschätzung, dass Teilnehmer kooperatives Verhalten eher menschlichen Entscheidern zuordnen, *Normann/Sternberg* (2021), 19.

positive Effekt des Algorithmus, indem die Unternehmen, die in einem Markt mit einem Algorithmus im Wettbewerb stehen, höhere Gewinne erzielen.⁸⁴⁶ Unterteilt man die Gewinne innerhalb eines Marktes jedoch nach Unternehmen mit und ohne Algorithmus auf, so zeigt sich, dass der Algorithmus vor allem für seine Wettbewerber auf dem Markt vorteilhaft ist. Auch wenn alle Unternehmen in den *Algorithm-Treatments* höhere Gewinne erzielen als in den *Human-Treatments*, liegen die Gewinne bei Unternehmen ohne einen Algorithmus über alle *Supergames* gesehen signifikant höher ($p < 0,01$), als bei Unternehmen mit einem Algorithmus.

Die Ergebnisse deuten somit darauf hin, dass ein kollusiver Algorithmus in der Lage ist, die Gewinne zu erhöhen und eine *tacit collusion* zu erzielen, wenngleich sein Einsatz mit Kosten für das Unternehmen verbunden ist.

Abbildung 8: Ergebnisse aus dem Laborexperiment von Normann und Sternberg (2021). Dargestellt sind die Gewinne der Unternehmen (in %) im dritten Supergame (Periode 6-19).



846 Die Ergebnisse eines linearen Wahrscheinlichkeitsmodells findet sich im Appendix zu dieser Arbeit (Tabelle 4).

3. Diskussion der Ergebnisse

Während sich der bisherige wissenschaftliche Diskurs zur algorithmischen Preissetzung auf die Ausbeutung der Nachfrage durch Algorithmen konzentriert, setzt sich das vorgestellte Experiment mit der Interaktion menschlicher und algorithmischer Anbieter auseinander. Dafür werden Märkte miteinander verglichen, auf denen ausschließlich menschliche Teilnehmer interagieren und Märkte, auf denen ein Unternehmen mit einem Preissetzungsalgorithmus ausgestattet ist. Für die Untersuchung des menschlichen Verhaltens auf hybriden Märkten kommt ein Algorithmus mit einer relativ einfachen statischen Strategie zum Einsatz, welche Kollusion zunächst fördert und sich in der Folge dem Verhalten der Wettbewerber anpasst.

Bisherige Laborexperimente zu *tacit collusion* mit ausschließlich menschlichen Teilnehmern haben gezeigt, dass es den Teilnehmern regelmäßig nicht gelingt, kollusive Marktergebnisse bei mehr als zwei Unternehmen zu erzielen. In unserem Labormarkt mit drei Wettbewerbern führt der Einsatz eines *pTFT*-Algorithmus hingegen zu signifikant höheren Preisen und einhergehend höheren Gewinnen für alle Unternehmen auf dem Markt.

Die Ergebnisse zeigen, dass die Teilnehmer die Erfahrungen vorheriger *Supergames* nutzen. Während es innerhalb eines *Supergames* schwierig zu sein scheint, seine Wettbewerber von einem kollusiven Gleichgewicht zu überzeugen, ergibt sich in einem neuen *Supergame* – und so auch einem neuen Markt – eine neue Möglichkeit zur Kollusion. Insbesondere das systematische Verhalten des Algorithmus erhöht sukzessive das Level der Kollusion von *Supergame* zu *Supergame*. Im dritten *Supergame* zeigen sich deshalb große Differenzen zwischen *Treatments* mit und *Treatments* ohne algorithmische Preissetzung. Insgesamt deuten die Daten auf Vor- und Nachteile des algorithmischen Einsatzes für die Marktakteure hin. Indem die Algorithmen systematisch kooperatives Verhalten belohnen, erleichtern sie die Koordinierung auf ein kollusives Gleichgewicht auf kognitiver Ebene. Zugleich scheint jedoch die Kenntnis über einen algorithmischen Wettbewerber die Kollusionsbereitschaft der menschlichen Teilnehmer zu beeinträchtigen. So assoziieren sie kompetitives Verhalten und das Abweichen von einer kollusiven Strategie vor allem mit der Anwesenheit algorithmischer Wettbewerber. Dadurch erscheint es schwieriger, das für eine Kollusion förderliche Vertrauensverhältnis aufzubauen, sodass im Vergleich der *Algorithm-Treatments* insbesondere das *Uncertain-Treatment* eine höhere Kooperationsrate hervorbringt.

Desweiteren zeigt das Experiment die Kosten einer Kollusion auf. Die Unternehmen der Märkte, auf denen der Algorithmus zum Einsatz kommt, erzielen alle insgesamt höhere Gewinne als die Unternehmen, die auf Märkten ohne Algorithmus aktiv sind. Allerdings verdienen die Unternehmen, die einen Algorithmus einsetzen, in dem vorgestellten Experiment weniger als ihre direkten Wettbewerber. Das Angebot der Kollusion durch den Algorithmus und die Einrichtung eines kollusiven Marktes ist somit mit Kosten für das den kollusiven Algorithmus einsetzende Unternehmen verbunden (*setup* Kosten). Hieraus ergibt sich ein *second-order public good* Problem.⁸⁴⁷ Demnach profitieren alle Unternehmen vom eingesetzten Algorithmus, wenngleich jedes einzelne Unternehmen einen Anreiz hat, selbst nicht auf den kollusiven Algorithmus zurückzugreifen. Dadurch stehen die Wettbewerber vor einem Koordinierungsproblem, sobald es darum geht, welches Unternehmen algorithmische, beziehungsweise kollusionsfördernde, Preissetzung zur Erzielung suprakompetitiver Gewinne implementiert. Hierbei gilt allerdings zu berücksichtigen, dass andere, in diesem Experiment nicht betrachtete, Eigenschaften algorithmischer Preissetzung diesen Nachteil teilweise ausgleichen könnten. Vorteile in Bezug auf die Vorhersage zukünftigen Nachfrageverhaltens oder die höhere Frequenz in der Preissetzung könnten dazu beitragen, die Verluste aufgrund der kollusiven Preissetzung abzumildern.

Insgesamt zeigen die Ergebnisse dieses Experiments das kollusive Potenzial algorithmischer Preissetzung auf, welches durch die Heterogenität der Märkte nicht verhindert wird.

VI. Ein weiteres Experiment zu hybriden Märkten

Die bereits erwähnte und im Anschluss an dieses Experiment erschienene Arbeit von Tobias Werner setzt sich ebenfalls mit dem Vergleich menschlicher und algorithmischer Preissetzung im Labor auseinander.⁸⁴⁸ Neben der getrennten Untersuchung des Verhaltens von algorithmischer und menschlicher Preissetzung in derselben Duopol- und Triopol-Marktumgebung,⁸⁴⁹ betrachtet der Autor auch gemischte Märkte, in denen selbstlernende Algorithmen sowie menschliche Marktteilnehmer miteinander im Wettbewerb stehen. Zu beachten ist hierbei erneut, dass der *Q-learning* Algorithmus in

847 Zum *second-order public good* Problem siehe bereits Fn. 562.

848 Werner (2021).

849 Siehe Kapitel D. II. 2. c).

einer Simulationsumgebung mit einem anderen Algorithmus trainiert wurde und erst nach Abschluss des Lernens in der hybriden Marktumgebung zum Einsatz kommt. Mithin entspricht der Algorithmus in seinem Einsatz einem statischen Algorithmus, welcher sein Verhalten nicht an das beobachtete Verhalten der menschlichen Gegenspieler anpassen kann. Eine Analyse der Strategie der Algorithmen zeigt, dass diese einem *win-stay, lose-shift* entspricht.⁸⁵⁰

In heterogener Zusammensetzung erzielen dabei ein Algorithmus und ein Mensch im Duopol ähnliche Ergebnisse wie zwei menschliche Marktteilnehmer.⁸⁵¹ Im Markt mit drei Wettbewerbern schneiden ein Algorithmus kombiniert mit zwei Teilnehmern allerdings signifikant schlechter ab als drei menschliche Teilnehmer sowie zwei Algorithmen und ein Teilnehmer.⁸⁵² Die Ergebnisse zeigen, dass einige Teilnehmer die Strategie des Algorithmus ausnutzen, um auf seine Kosten höhere Gewinne zu erzielen.⁸⁵³

Werners Experiment zeigt sowohl die Stärken als auch die Schwächen der algorithmischen Preissetzung mit Hilfe von *Q-learning* auf. Insbesondere im homogenen Duopol – „unter sich“ – erreichen die Algorithmen suprakompetitive Marktergebnisse, welche den Versuchsteilnehmern überlegen sind. Zugleich benötigen die Algorithmen hierfür eine relativ lange Lernphase. Auf neues, zuvor unbeobachtetes Verhalten können sich die *off-the-job* trainierten Algorithmen nicht einstellen. Hieraus ergibt sich ein ausbeuterisches Potenzial, welches – wie bereits bei *Hettich* gezeigt⁸⁵⁴ – von überlegenen Algorithmen oder aufmerksamen Marktteilnehmern ausgenutzt werden kann.

VII. Die Übertragbarkeit (verhaltens-)ökonomischer Erkenntnisse

In den vorangegangenen Kapiteln wurden die ökonomischen Erkenntnisse zum Auftreten einer *tacit collusion* und algorithmischer Kollusion präsentiert und diskutiert. Im Folgenden soll die generelle Übertragbarkeit hierbei angewandter ökonomischer Modelle sowie verhaltensökonomischer Experimente diskutiert werden.

850 Werner (2021), S. 25; siehe hierzu Kapitel D. II. 2. c).

851 Werner (2021), S. 30.

852 Werner (2021), S. 30 f.

853 Werner (2021), S. 32.

854 Siehe Kapitel C. II. 2. b) dd) (2).

1. Die Realitätsferne ökonomischer Modelle

Ökonomische Modelle sehen sich häufig der Kritik ausgesetzt, dass sich aus ihnen „[keine] spezifischen Voraussagen von Einzelereignissen ableiten [lassen]“. ⁸⁵⁵ Insbesondere, die in ökonomischen Modelle getroffenen stark vereinfachten Annahmen über Marktgegebenheiten und das Verhalten der Marktteilnehmer können große Diskrepanzen zu den realen Märkten aufweisen, weshalb eine nur geringe Aussagekraft für die Praxis beklagt wird. ⁸⁵⁶ Hierbei ist jedoch festzustellen, dass Konzepte, wie das des perfekten Wettbewerbs, als Modell zu verstehen sind und nicht den tatsächlichen Wettbewerb auf realen Märkten darstellen sollen. ⁸⁵⁷ Außerdem ist der Zweck einer ökonomischen Analyse oft nicht die Untersuchung der tatsächlichen Marktgegebenheiten, sondern die Untersuchung des spezifischen Zusammenhangs zwischen dem Marktergebnis und potenziellen Einflussfaktoren. Dennoch hat sich die Verhaltensökonomie der Kritik angenommen und dient der Überprüfung und Weiterentwicklung entsprechender Modelle. Mit empirischen Erkenntnissen über das Verhalten nicht rein theoretischer und rational agierender Akteure möchte das Forschungsgebiet dazu beitragen, Grundannahmen fortzuentwickeln und Prognosen zu verbessern.

2. Die Validität verhaltensökonomischer Experimente

Auch die Verhaltensökonomie steht im Spannungsfeld zwischen einer notwendigen Vereinfachung der Modelle zur Analyse spezifischer Effekte und der Übertragbarkeit der Ergebnisse auf reale Gegebenheiten. Die Diskussion über die Nützlichkeit der Ergebnisse experimenteller Forschung beschäftigt sich mit der Frage der externen und internen Validität. Während sich die interne Validität damit auseinandersetzt, inwieweit ein Experiment „tatsächlich das Modell oder die Theorie abbildet“, fragt die externe Validität, ob das Ergebnis eines Experiments „auf die reale Welt außerhalb des Labors übertragen werden kann.“ ⁸⁵⁸

Zur einer „realistischen“ Betrachtung natürlicher Märkte bietet sich das Werkzeug des Feldexperiments an. Bei diesem werden die Experimente

855 Hayek, Freiburger Studien, S. 252.

856 Für eine Auseinandersetzung mit der Kritik am ökonomischen Ansatz, *Kasten*, Höchstpreisbindungen, S. 115 ff.

857 *Kasten*, Höchstpreisbindungen, S. 113.

858 *Weimann/Brosig-Koch*, Einführung in die experimentelle Wirtschaftsforschung, S. 30.

in einem natürlichen Umfeld durchgeführt.⁸⁵⁹ Aufgrund der realen Bedingungen können Feldexperimente im Besonderen bezüglich der externen Validität überzeugen. Allerdings lassen sich in natürlicher Umgebung die vielen äußeren Einflussfaktoren nur schwer kontrollieren, weshalb zwar Korrelationen aufgedeckt werden können, es jedoch häufig schwierig ist Kausalbeziehungen herzustellen.⁸⁶⁰ Mithin mangelt es bei Feldexperimenten an der internen Validität.

Die interne Validität ist jedoch die große Stärke des Standardwerkzeugs der Verhaltensökonomie, dem Laborexperiment.⁸⁶¹ Indem die Experimentatoren das Entscheidungsumfeld sowie die Handlungsoptionen vorgeben und kontrollieren können, lassen sich Einflussfaktoren isoliert betrachten und Kausalbeziehungen identifizieren.⁸⁶² Dies geht jedoch zugleich zu Lasten der externen Validität. Um sicherzustellen, dass es nicht mehr als den einen Unterschied zwischen *Baseline* und *Treatment* gibt, weicht die Laborumgebung in der Regel deutlich von den Gegebenheiten realer Märkte ab.⁸⁶³ Aus diesem Grund lassen sich Befunde aus dem Labor nicht ohne weiteres verallgemeinern und auf reale Märkte projizieren.

In den vergangenen Jahren hat es daher einige Kritik an der Aussagekraft von Laborexperimenten gegeben.⁸⁶⁴ Im wettbewerbsrechtlichen Kontext wird insbesondere die Rolle der Studierenden als Repräsentanten eines Unternehmens kritisch gesehen. Darüber hinaus wird die Diskrepanz zwischen der experimentellen und der realen Entscheidungssituation angemerkt. Im Folgenden sollen einige dieser Aspekte in Bezug auf die Aussagekraft verhaltenswissenschaftlicher Erkenntnisse im kartellrechtlichen Kontext diskutiert werden.

859 Haus, Das ökonomische Laboratop, S. 19; zum Instrument des Feldexperiments, siehe Engel, in: Zamir (Hrsg.), The Oxford handbook of behavioral economics and the law, S. 125 (131 ff.).

860 Engel (2015), JELS 12 (3), 537 (568).

861 Engel, in: Zamir (Hrsg.), The Oxford handbook of behavioral economics and the law, S. 125 (135).

862 Engel (2015), JELS 12 (3), 537 (568).

863 Engel (2015), JELS 12 (3), 537 (568).

864 So etwa Levitt/List (2007), JEP 21 (2), 153.

a) Die Studierenden als Unternehmer

Teilnehmer verhaltenswissenschaftlicher Experimente sind ganz überwiegend Studierende unterschiedlicher Fachrichtungen.⁸⁶⁵ Bezogen auf die zu untersuchenden Marktsituationen wird teilweise in Frage gestellt, inwieweit die Teilnehmer hinreichende Erfahrung in Bezug auf unternehmerische Marktentscheidungen mitbringen. Studierende könnten aus diesem Grund zum Beispiel risikoscheuer agieren als Managerinnen in einem Angestelltenverhältnis. Hinzu kommt, dass in manchen Unternehmen Entscheidungen in Arbeitsgruppen getroffen werden, im Labor Studierende jedoch allein für das Unternehmen entscheiden.⁸⁶⁶ Ein weiterer Punkt betrifft eine mögliche Selbstselektion unter den Studierenden. Teilweise wird vermutet, dass eine sehr bestimmte Gruppe von prosozialen Studierenden von Laborexperimenten angesprochen wird, während Managerinnen stärker auf das Interesse ihres Unternehmens bedacht sind.⁸⁶⁷ Diese Punkte stellen in Frage, inwiefern eine Verallgemeinerung der Ergebnisse – insbesondere im wettbewerblichen Kontext – möglich erscheint.

Zunächst ist den Bedenken entgegenzuhalten, dass Laborexperimente grundsätzlich nicht zum Ziel haben, reale Märkte abzubilden. Dementsprechend ist es nicht zwingend hilfreich, das Verhalten professioneller Akteure zu untersuchen. Diese könnten in einem Experiment mehr oder weniger blind entsprechend ihre Erfahrungen handeln, wenngleich der Versuchsaufbau Unterschiede zu ihrem Berufsalltag aufweist.⁸⁶⁸ Damit Experimente einen Mehrwert bieten, sollte sich das Verhalten der Studierenden allerdings auch nicht grundlegend von dem Verhalten der Marktakteure unterscheiden.

Ausgangspunkt ist die Frage, welchen Einfluss die Expertise auf die Entscheidungen von Menschen hat und inwiefern Entscheidungen der Expertinnen eines Gebiets von den Entscheidungen üblichen Versuchsteilnehmern abweichen.⁸⁶⁹ Hierbei steht der Verdacht im Raum, dass Expertinnen häufig rationaler agieren, als Studierende im Versuchslabor.

865 A. Morell, (Behavioral) Law and Economics im europäischen Wettbewerbsrecht, S. 201 f.

866 Tor, in: Zamir (Hrsg.), The Oxford handbook of behavioral economics and the law, S. 539 (549).

867 Levitt/List (2007), JEP 21 (2), 153 (165 f.).

868 Feltovich (2011), Journal of Economic Surveys 25 (2), 371 (373), mit Verweis auf Binmore, Does Game Theory Work?, S. 9 f.

869 Engel (2010), Journal of Institutional Economics 6 (4), 445.

In Experimente wurde in der Tat gezeigt, dass Expertinnen insbesondere in Umgebungen, in denen ihre Handlungen ein klare Rückmeldung erzeugen und sie ihre Erfahrungen nutzen können besonders rationale Entscheidungen treffen.⁸⁷⁰ In Situationen, in denen die Rückmeldungen auf eine Handlung jedoch nur schwer zu interpretieren ist, entfällt dieser Effekt.⁸⁷¹ Die Tatsache, dass Algorithmen die Qualität der Entscheidungen von Expertinnen häufig übertreffen,⁸⁷² zeigt darüber hinaus auf, dass auch Expertise nicht vor Fehleinschätzungen und irrationalen Erwägungen schützt. Auch Expertinnen weichen in ihren Entscheidungen vom Modell des *homo oeconomicus* ab. Unter der Annahme vollständiger Rationalität der Entscheider hätten einige Kartelle, welche in der Praxis aufgedeckt wurden, in der Theorie nicht bestehen können.⁸⁷³ Doch inwiefern weicht das Verhalten von Versuchsteilnehmern von dem Verhalten der Expertinnen ab?

Untersuchungen zu den unterschiedlichen Entscheidungen zwischen Expertinnen und Laien sowie üblichen Versuchsteilnehmern und andere Bevölkerungsgruppen liefern Anhaltspunkte für eine Beantwortung dieser Frage. Zunächst konnten verschiedene Experimente zeigen, dass auch Studierende im Vergleich zu Nicht-Studierenden häufig eigennütziger und rationaler agieren.⁸⁷⁴ Im Vergleich der Entscheidungen von Expertinnen und Studierenden in Bezug auf die im Labor auftretenden Effekte konnten in der überwiegenden Zahl der hierzu durchgeführten Experimente keine wesentlichen Unterschiede festgestellt werden.⁸⁷⁵

Auch der Verdacht möglicher Selektionseffekte im Rahmen von Laborexperimenten, wonach vorzugsweise prosoziale Probanden Teilnehmer eines Experiments würden, ließ sich in entsprechenden Untersuchungen nicht bestätigen.⁸⁷⁶ Unterschiede zwischen Entscheidungen von Einzelpersonen und Gruppenentscheidungen konnten in Experimenten festgestellt werden,

870 D. Kahneman/G. Klein, *The American Psychologist* 64 (6) (2009), 515 (523).

871 D. Kahneman/G. Klein, *The American Psychologist* 64 (6) (2009), 515 (523).

872 D. Kahneman/G. Klein, *The American Psychologist* 64 (6) (2009), 515 (523).

873 Tor, in: Zamir (Hrsg.), *The Oxford handbook of behavioral economics and the law*, S. 539 (555).

874 Belot et al. (2015), *JEBO* 113, 26; Falk et al. (2013), *JEEA* 11 (4), 839; Cappelen et al. (2015), *Scand. J. of Econ* 117 (4), 1306.

875 Lediglich bei einem der 13 beobachteten Experimenten, welche sich mit dem Vergleich von Studierenden und Expertinnen befassten, wiesen die Expertinnen ein Ergebnis auf, welches näher an dem von der Theorie vorhergesagten Verhalten lag. In neun der dreizehn Experimente konnten keine wesentlichen Unterschiede festgestellt werden, Fréchette, in: Fréchette (Hrsg.), *Fréchette* 2015, S. 360.

876 Falk et al. (2013), *JEEA* 11 (4), 839.

allerdings fielen die Entscheidungen der Gruppe in Abhängigkeit des konkreten Falls sowohl rationaler, als auch irrationaler aus.⁸⁷⁷

Die Entscheidungen von Unternehmen können stark von den handelnden Personen, den internen Entscheidungsstrukturen sowie den übrigen Gegebenheiten des Einzelfalls abhängen. Insgesamt konnten die bisherigen Untersuchungen keine generellen Unterscheidungsmerkmale zwischen Marktakteuren und Versuchsteilnehmern herausarbeiten. Die hierzu durchgeführten Experimente deuten darauf hin, dass Studierende grundsätzlich geeignet sind, das Verhalten von Expertinnen in Entscheidungssituationen abzubilden. Zur Berücksichtigung menschlicher Einflüsse in Marktentscheidungsprozessen scheinen Experimente mit Studierende deshalb grundsätzlich eine geeignete Methode darzustellen.

b) Die Besonderheiten der Laborumgebung

Weitere Kritikpunkte beziehen sich auf die Bedingungen der Laborexperimente. Bei einem Experiment sitzen die Teilnehmer in der Regel in abgetrennten Kabinen an einem Computer in einem Laborraum. Für ihre Teilnahme erhalten sie einen Geldbetrag. In unserem Experiment betrug die durchschnittliche Auszahlung an den Versuchsteilnehmer beispielsweise 17,73 Euro.⁸⁷⁸ Diese sich von Marktsituationen deutlich unterscheidenden Bedingungen sowie das Bewusstsein, Teil eines Experiments zu sein, könnten das Verhalten der Teilnehmer beeinflussen.⁸⁷⁹ Eine Sorge ist, dass Teilnehmer versuchen, ihre Entscheidungen so zu treffen, wie sie vermeintlich von der Experimentatorin gewünscht sind (*demand*-Effekte).⁸⁸⁰ Dieses Verhalten würde unter anderem dadurch befördert, dass Entscheidungen in einem Laborexperiment nicht hinreichend anonym getroffen werden könnten.⁸⁸¹ Außerdem wird die Sorge geäußert, dass die in einem Labor ausgezahlten Geldbeträge zu gering seien, um wirksame Anreize für die Teilnehmer zu setzen.⁸⁸²

877 Tor, in: Zamir (Hrsg.), The Oxford handbook of behavioral economics and the law, S. 539 (550); Engel (2010), Journal of Institutional Economics 6 (4), 445.

878 Normann/Sternberg (2021), 12.

879 A. Morell, (Behavioral) Law and Economics im europäischen Wettbewerbsrecht, S. 197 ff.

880 Weimann/Brosig-Koch, Einführung in die experimentelle Wirtschaftsforschung, S. 38.

881 Levitt/List (2007), JEP 21 (2), 153 (161).

882 Levitt/List (2007), JEP 21 (2), 153 (164).

Zur Überprüfung eines *demand*-Effekts finden *Jonathan de Quidt et al.* in ihrem Experiment, dass die explizite Erwähnung der erwarteten Handlung in den Instruktionen einen schwach positiven Effekt auf die Entscheidungen der Teilnehmer hat.⁸⁸³ Aus diesem Grund sollten Experimentatorinnen darauf achten, die Instruktionen eines Experiments immer ergebnisoffen zu formulieren. Im Rahmen der vorgestellten Markt-Spiele dürfte ein möglicher *demand*-Effekt aber keinen besonderen Einfluss auf das Verhalten der Teilnehmer haben, da nicht ersichtlich ist, in welche Richtung sich dieses anpassen sollte. Bezüglich der Anonymität der Teilnehmer deutet eine Reihe neuerer Experimente darauf hin, dass diese keinen besonderen Effekt auf das Handeln im Labor hat.⁸⁸⁴

Jeffrey Carpenter et al. untersuchen den Einfluss der Höhe des zu verdienenden Geldbetrags auf die Entscheidung der Spielerinnen in einem Ultimatumspiel sowie einem Diktatorspiel. In ihrem Experiment variieren sie den Geldbetrag zwischen 10 \$ und 100 \$, können allerdings keinen Einfluss der Höhe des Geldbetrags auf das Gebot des anbietenden Spielers feststellen.⁸⁸⁵ Allerdings scheint auch die Frage der finanziellen Anreize in Laborexperimenten vom jeweiligen experimentellen Kontext abhängig zu sein.⁸⁸⁶ Für Experimente mit Wettbewerbsbezug scheint die Höhe der Auszahlung nur einen geringen Einfluss auf das Verhalten unter Risiko zu haben.⁸⁸⁷

c) Schlussfolgerung

In speziell darauf ausgerichteten Experimenten wurde der Frage nachgegangen, inwiefern die Besonderheiten der Laborbedingungen sowie die Auswahl der Teilnehmer Einfluss auf die Effekte entsprechender Experimente nehmen können. Dies zeigt, dass die Verhaltensökonomie die aufgekommene Kritik aufnimmt und sich mit möglichen Schwächen des Instruments des Labor-experiments auseinandersetzt. Obwohl sich ein großer Teil der geäußerten Sorgen empirisch bisher nicht bestätigen ließ, sollten die vorgebrachten

883 *Quidt et al.* (2018), AER 108 (11), 3266.

884 *Barmettler et al.* (2012), GEB 75 (1), 17; *Kogler et al.* (2020), JEBO 177, 390; *Fries et al.* (2021), JEBO 189, 132; *Vorlauffer* (2019), JBEE 81, 216.

885 *Carpenter et al.* (2005), ECOLET 86 (3), 393.

886 *Falk/Heckman*, Science 326 (5952) (2009), 535 (537).

887 *Weimann/Brosig-Koch*, Einführung in die experimentelle Wirtschaftsforschung, S. 66.

Bedenken nicht ignoriert werden und sowohl bei der Durchführung als auch der Interpretation der Ergebnisse Berücksichtigung finden.

So bieten Laborexperimente die Möglichkeit, Effekte dahingehend zu untersuchen, ob innerhalb einer Gruppe von Probanden besondere Faktoren, wie das Geschlecht, Alter oder Studiengang einen Einfluss auf die Ergebnisse haben.⁸⁸⁸ Auch die Anonymität der Versuchsteilnehmer kann durch besondere Verfahren sichergestellt werden.⁸⁸⁹ Besondere Aussagekraft haben gefundene Effekte insbesondere dann, wenn sie in mehreren Experimenten zu demselben Themenbereich auftreten oder wenn der Effekt eines Experiments in weiteren Experimenten repliziert werden kann.⁸⁹⁰ Aus diesem Grund sind insbesondere die vorgestellten Meta-Studien ein hilfreiches Mittel, um die Ergebnisse eines Forschungsgebiets zu beurteilen.⁸⁹¹

Es bleibt jedoch festzuhalten, dass die in Experimenten getroffenen Beobachtungen immer „aus artifiziellen Umgebungen stammen und deshalb nicht ohne Weiteres auf die Realität übertragbar sind.“⁸⁹² Weder die Befunde der ökonomischen Theorie, noch Feld- oder Laborexperimente können für sich in Anspruch nehmen, ein Phänomen zweifelsfrei und allgemeingültig aufzuklären. Das bedeutet nicht, dass entsprechende Erkenntnisse keinen Wert für das Kartellrecht bieten, sondern im Gegenteil sollte sich das Kartellrecht sämtliche Methoden bei der Untersuchung potenzieller Effekte zu Nutzen machen. „Daten aus dem Feld, Daten aus Umfragen, Experimente im Feld sowie im Labor, ebenso wie die Standardmethoden der Ökonometrie können gemeinsam den Wissensstand der Sozialwissenschaften verbessern.“⁸⁹³

VIII. Zwischenergebnis

Das vorgestellte Experiment fügt sich in den wissenschaftlichen Diskurs zur algorithmischen Preissetzung ein und befasst sich mit der Interaktion algorithmischer und menschlicher Anbieter. Während menschliche Anbieter in den bisherigen Laborexperimenten auf Märkten mit mehr als zwei

888 *Falk/Heckman*, Science 326 (5952) (2009), 535 (537).

889 Siehe beispielsweise zum Doppelblindverfahren, *Weimann/Brosig-Koch*, Einführung in die experimentelle Wirtschaftsforschung, S. 118.

890 Zur Replikation sowie weiterer Verfahren, *Camerer*, in: Schram/Ule (Hrsg.), *Handbook of research methods*, S. 83.

891 Siehe zu *tacit collusion*, *Engel* (2015), JELS 12 (3), 537.

892 *Weimann/Brosig-Koch*, Einführung in die experimentelle Wirtschaftsforschung, S. 30.

893 Aus dem Englischen übersetzt, siehe *Falk/Heckman*, Science 326 (5952) (2009), 535 (537).

Unternehmen Schwierigkeiten hatten, Kollusion zu erzielen, zeigen die Ergebnisse dieses Experiments, dass der Einsatz eines *pTFT*-Algorithmus zu signifikant höheren Preisen und einhergehend höheren Gewinnen für alle Unternehmen auf einem Triopol-Markt führt. Somit deuten die Ergebnisse daraufhin, dass der Einsatz entsprechender Preisalgorithmen Kollusion auch in weniger stark konzentrierten Märkten erleichtert. Das systematische und kollusionsfördernde Verhalten des Algorithmus erhöht sukzessive das Level der Kollusion von *Supergame* zu *Supergame*. Insbesondere im dritten *Supergame* weisen die *Algorithm-Treatments* deutlich höhere Kooperationsraten als die *Human-Treatments* auf. Zugleich zeigen die Ergebnisse, dass Unternehmen vor einem Koordinierungsproblem stehen, sobald es darum geht, welches Unternehmen algorithmische beziehungsweise kollusionsfördernde Preissetzung implementiert (*second-order public good*).

Laborexperimente allein sind jedoch nicht geeignet das Phänomen algorithmischer Preissetzung allgemeingültig aufzuklären. In diesem Experiment wurde ein ganz spezieller Algorithmus in einem stark vereinfachten Marktgeschehen betrachtet, sodass sich diese Befunde nicht ohne Weiteres auf reale Märkte übertragen lassen. In Kombination mit den Befunden aus den vorangegangenen Kapiteln legt die vielfältige Literatur aus ökonomischer Theorie, Simulationen, Laborexperimenten und der Analyse von Daten aus dem Feld allerdings ein differenziertes Bild algorithmischer Preissetzung dar. Während einige Probleme in der Anwendung vor allem in Bezug auf selbstlernende Systeme offensichtlich noch zu lösen sind, deutet die Literatur insgesamt auf ein grundsätzlich wettbewerbsschädliches Potenzial algorithmischen Preissetzung hin.

Das vielfach vorgebrachte Argument, *tacit collusion* sei nur schwer zu erreichen und deshalb äußerst unwahrscheinlich, lässt sich aufgrund der besonderen Eigenschaften algorithmischer Preissetzung für digitale Märkte in Zweifel ziehen. Wenngleich die bisherigen Erkenntnisse aufzeigen, dass *tacit collusion* keine zwingende Folge des Einsatzes algorithmischer Kollusion ist, deutet vieles darauf hin, dass Algorithmen bei günstigen Marktbedingungen leichter zu einem kollusiven Marktergebnis gelangen als menschliche Entscheider. Insbesondere die Schlussfolgerung von *Kastius* und *Schlosser*, wonach Unternehmen lediglich Preise über den Markt senden müssten, um *tacit collusion* zu erzielen, sollte als Warnung verstanden werden.⁸⁹⁴ Besonders gefährlich scheinen bisher weniger die von algorithmischen Systemen entwickelten hoch komplexen Modelle, als vielmehr relativ einfache

894 Vgl. *Kastius/Schlosser* (2022), J Rev Pricing Man 21, 50.

Strategien in Verbindung mit den Vorzügen digitaler Märkte. Insbesondere die aufeinander abgestimmte Preissetzung statischer Algorithmen erhöht die Interdependenzen auf digitalen Märkten und sorgt darüber hinaus für eine bessere Vorhersage zukünftigen Verhaltens der Wettbewerber.

