

ESSAYS ON THE DIGITAL ECONOMY

DISSERTATION

zur Erlangung des akademischen Grades
doctor rerum politicarum (Dr. rer. pol.)

eingereicht an der Wirtschaftswissenschaftlichen Fakultät
der Heinrich-Heine Universität Düsseldorf

von

TOBIAS FELIX WERNER, M.Sc.

geb. am 13. April 1992 in Schwelm

Erstgutachter: PROF. DR. HANS-THEO NORMANN

Zweitgutachter: PROF. DR. MATTHIAS HUNOLD

Acknowledgements

I am grateful to my supervisor, Hans-Theo Normann, for his support and guidance. His passion for experimental methods in industrial organization influenced the focus of my dissertation, and I am thankful for his valuable feedback. Without his guidance, this dissertation would not have been possible.

Furthermore, I am grateful to my second supervisor, Matthias Hunold. From hiring me for a project-related position to providing valuable insights throughout the writing process, Matthias has been a constant source of support and guidance.

Also, I want to thank my co-authors and all other colleagues, especially Simon Cordes, Markus Dertwinkel-Kalt, Mats Köster, and Ulrich Laitenberger. Special thanks to Adrian Hillenbrand and Fabian Winter for sparking my interest in experimental economics and for their guidance. They could be considered my third and fourth supervisors.

Most importantly, I want to thank my family and friends. It would not have been possible without their love and support in the last few years. I am grateful to my mother for her constant encouragement throughout my academic journey. Last but not least, I thank my partner, Vasilisa Petrishcheva. She supported me whenever I got stuck and was always there for me in every moment. I am deeply grateful for her feedback and support throughout my doctoral studies.

Contents

Introduction	1
1 Algorithmic and Human Collusion	5
1.1 Introduction	6
1.2 Pricing algorithms	9
1.2.1 Basic Q-learning	10
1.2.2 Simulation setup	12
1.2.3 Performance evaluation	13
1.3 Market environment	16
1.4 Experimental design	17
1.4.1 Treatments	17
1.4.2 Procedure	18
1.5 Hypotheses	20
1.6 Results	22
1.6.1 Performance of the algorithms	23
1.6.2 Strategies of the algorithms	25
1.6.3 Comparing algorithmic and human collusion	29
1.6.4 Collusion between humans and algorithm	31
1.7 Concluding remarks	37
Appendices	
1.A Implementation details	39
1.A.1 Instructions	39
1.A.2 Screenshots of the experiment	42
1.A.3 Survey & control questions	43
1.B Strategy analysis	44
1.B.1 Incentive compatibility of the algorithm's strategy	44

1.B.2	Individual prices in the last supergame	45
2	Algorithmic Price Recommendations and Collusion: Experimental Evidence	49
2.1	Introduction	50
2.2	Theoretical framework and predictions	54
2.2.1	Setting and Nash equilibria	55
2.2.2	Algorithm that recommends Nash equilibrium actions	57
2.2.3	Behaviourally motivated soft punishment algorithms	59
2.3	Experimental design	61
2.3.1	Treatments	62
2.3.2	Procedure	63
2.4	Results	64
2.4.1	The influence of price recommendations on individual prices	65
2.4.2	Collusive effects of price recommendations	66
2.4.3	Heterogeneity and mechanisms	68
2.5	Concluding remarks	76
Appendices		
2.A	Additional theoretical results	81
2.B	Instructions and survey questions	84
2.B.1	Instructions	84
2.B.2	Survey questions	88
2.C	Further results	90
2.C.1	Randomization checks	90
2.C.2	Economic preferences and recommendations	90
2.C.3	Additional control treatments	94
2.C.4	Further figures and tables	95
3	Willingness to Volunteer Among Remote Workers is Insensitive to the Team Size	99
3.1	Introduction	100
3.2	Experimental design	103
3.2.1	Workers and the online labor market	105

3.2.2	The advertised job	106
3.2.3	The volunteering decision	108
3.2.4	Treatments	109
3.2.5	Belief elicitation	110
3.2.6	Questionnaire	110
3.2.7	Replication	111
3.3	Hypotheses	111
3.4	Results	114
3.4.1	Workers are sensitive to incentives and strategic situations	114
3.4.2	Heterogeneity in costs to volunteer	115
3.4.3	Insensitivity to team size in the volunteers dilemma	117
3.4.4	Beliefs	117
3.4.5	Conditional volunteering	119
3.4.6	Robustness of our results	121
3.5	Discussion and concluding remarks	122

Appendices

3.A	A formal model of volunteering at the workplace	125
3.A.1	Proofs	128
3.B	Additional figures and tables	132
3.C	Confirming the robustness of the results	134
3.C.1	Belief order effects	135
3.C.2	Prominence of group size	136
3.C.3	Workplace-related reputational concerns	136
3.C.4	Economic preferences	138
3.C.5	Gender differences in volunteering	139
3.C.6	Beliefs relative to other explanatory variables	139
3.C.7	Robustness checks with regard to control variables	141
3.D	Experiment instructions	149
3.D.1	The introduction	149
3.D.2	The task	152
3.D.3	The volunteering decision	153
3.D.4	Belief elicitation	155

3.D.5	Questionnaire items	158
4	What Drives Demand for Loot Boxes?	161
4.1	Introduction	162
4.2	Experimental design	167
4.2.1	Experimental setup	167
4.2.2	Conceptual framework	169
4.2.3	Implementation and logistics	171
4.3	Main results	172
4.4	Additional results	173
4.4.1	Treatment effects under realistic beliefs	173
4.4.2	Robustness experiment	175
4.4.3	Correlates of overspending on loot boxes	176
4.5	Conclusion	177
Appendices		
4.A	Additional tables	179
4.B	Experimental setup: Robustness experiment	182
4.C	Experimental details	183

List of Figures

1.1	Profitability of the Q-learning agents	23
1.2	Average market price of the Q-learning agents	24
1.3	Share of simulations that converge to a Nash equilibrium	25
1.4	Punishment behavior of the algorithms	26
1.5	Market prices for the algorithmic and human markets	30
1.6	Average market prices for all treatments	32
1.7	Average market prices by supergame and round	34
1.8	Decision screen	42
1.9	Profit information screen	43
1.10	Prices for each human participant in the last supergame in 1H1A . . .	46
1.11	Prices for each human participant in the last supergame in 1H2A . . .	47
1.12	Prices for each market in the last supergame in 2H1A	48
2.1	Market price for the main treatments	66
2.2	Market price for each treatment by supergame and round	74
2.3	Market price in RECTHEORY by High/Low split	75
2.4	Market price for all additional algorithms	95
2.5	Market price in RECSoft by High/Low split	96
3.1	Structure of the experiment	105
3.2	Volunteering rates in the baseline treatments	115
3.3	Volunteering rate across the different group sizes	117
3.4	Beliefs by group size	119
3.5	Relationship between beliefs and the volunteering rate	122
3.6	Volunteering rates in the initial experiment	132
3.7	Distributions of the different cost measures	132
3.8	Relationship between cognitive uncertainty and beliefs	133

3.9	Average values of several observables across treatments	135
3.10	Screenshot of the task	152
3.11	Screenshot of the volunteering decision	154
3.12	Additional screenshot of the volunteering decision	155
3.13	Screenshot of the belief elicitation	156
3.14	Additional screenshot of the belief elicitation	157
4.1	Typical design features of loot boxes	162
4.2	Decision screen in the Control condition after stating the belief	168
4.3	Information provided in the treatment Censored	168
4.4	Information provided in the treatment Sample	169
4.5	Average willingness-to-pay and beliefs by treatment	172
4.6	Average willingness-to-pay and beliefs by treatment with realistic beliefs	174
4.7	Average willingness-to-pay and beliefs by treatment in the robustness experiment	175
4.8	Attention check at the beginning of the experiment.	183
4.9	Welcome page.	183
4.10	Instructions on the experiment.	184
4.11	Additional info box with details on the payment mechanism.	184
4.12	Comprehension questions.	185
4.13	Start of the experiment.	185
4.14	Belief & WTP elicitation in the “Control” treatment.	186
4.15	Belief & WTP elicitation in the “Sample” treatment.	187
4.16	Belief & WTP elicitation in “Censored” treatment.	188
4.17	Belief & WTP elicitation in “Joint” treatment.	189
4.18	Additional info in “Info” treatment.	190
4.19	Belief & WTP elicitation in “Info” treatment.	191
4.20	General control questions.	192
4.21	Questions about loot box experience.	192
4.22	Questions for loot box users.	193
4.23	Self-control survey questions (Part 1).	193
4.24	Self-control survey questions (Part 2).	194
4.25	Gambling survey questions (Part 1).	195

4.26	Gambling survey questions (Part 2).	195
4.27	Last page in the experiment.	196

List of Tables

1.1	Treatment composition	18
1.2	Number of observations by treatment	20
1.3	Estimated proportion for each strategy	36
2.1	Number of observations by treatment	64
2.2	Individual prices explained by the recommendation in a linear regression	65
2.3	Linear regression for the treatment effects	68
2.4	Market price statistics by treatment	69
2.5	Market price explained by negative reciprocity and treatments	71
2.6	Market price explained by time preferences and treatments	72
2.7	Survey measures by treatment	90
2.8	Market price explained by altruism and treatments	91
2.9	Market price explained by positive reciprocity and treatments	92
2.10	Market price explained by risk preferences and treatments	93
2.11	Market price explained by trust preferences and treatments	94
2.12	Linear regression with the treatment effects by supergame	95
3.1	Number of observations for each treatment.	106
3.2	Logistic regression to estimate the volunteering choice as a function of different cost measures	116
3.3	Belief overview by treatment	118
3.4	Volunteering choice explained by Beliefs and all treatments	120
3.5	Logistic Regression estimating the probability to volunteer	133
3.6	Estimated average marginal effects for the main treatment variable . .	133
3.7	Estimated average marginal effects for model specification (3) in Table 3.2	134
3.8	Mean of different socioeconomic and demographic variables	134

3.9	Volunteering choice explained by the order of the belief elicitation . . .	137
3.10	Belief overview by treatment	138
3.11	Volunteering choice explained by different control variables	140
3.12	Volunteering choice explained by REPUTATION	141
3.13	Volunteering choice explained by gender	142
3.14	Volunteering choice explained by EFFICIENCY	143
3.15	Volunteering choice explained by ALTRUISM	144
3.16	Volunteering choice explained by N. RECIPROCITY	145
3.17	Volunteering choice explained by RISK	146
3.18	Volunteering choice explained by TRUST	147
3.19	Volunteering choice explained by TIME	148
4.1	Regression results — main specification	173
4.2	Regression results — predictors for real world overspending	176
4.3	Regression results — main specification — realistic beliefs	179
4.4	Regression results — robustness experiment	180
4.5	Summary statistics	180
4.6	Loot box statistics	181

Introduction

In this cumulative dissertation, my coauthors and I study different aspects of the digital economy. The digitalization of markets alters various economic interactions, which may change market outcomes compared to classical analog markets. I use methods from experimental economics, as well as theoretical modeling and simulations studies, to uncover different aspects of this digitalization and its implications. The first two chapters of this dissertation focus on the influence of pricing algorithms on markets. In Chapter 1, I investigate the ability of self-learning pricing algorithms to collude on non-competitive prices. I compare them to human behavior in a market experiment and highlight in which situations algorithms are more collusive than humans. Chapter 2 discusses the possibility that online platforms can use price recommendation algorithms to make the seller market more collusive. The third and fourth chapters shift their focus away from algorithms and examine other aspects of the digital economy. Chapter 3 focuses on the cooperation among remote online workers, whereas Chapter 4 discusses methods that video game developers use to monetize their games via so-called “loot boxes”. In the following, I provide a summary of each chapter.

Chapter 1 discusses the effect of self-learning pricing algorithms on market outcomes. When firms use those tools, the pricing decision for a specific good in a market is not made by a human pricing manager anymore, but it is outsourced to an algorithm that takes all pricing decisions. Often those algorithms do not follow any predefined strategy but learn by themselves how to price the product (Ezrachi and Stucke, 2017). With the rise in popularity, there have been increasing concerns from competition authorities and academics that those algorithms could also learn to tacitly collude on non-competitive prices (see, for instance, Bundeskartellamt and Autorité de la concurrence, 2019). Building on previous work from Calvano et al. (2020a), I show that popular self-learning algorithms are collusive in market simulations. Then, I conduct market experiments with humans in the same

environment to derive a counterfactual that resembles traditional tacit collusion. Across different treatments, I vary the market size and the number of firms that use a pricing algorithm. I demonstrate that oligopoly markets become more collusive if algorithms make pricing decisions instead of humans. In two-firm markets, prices are weakly increasing in the number of algorithms in the market. In three-firm markets, algorithms weaken competition if most firms use an algorithm and human sellers are inexperienced. The results highlight the anti-competitive effect of pricing algorithms, and increased scrutiny by competition authorities is warranted.

Also, Chapter 2 discusses the influence of algorithms on market outcomes. Many online platforms like *Airbnb.com* or *Ebay.com* provide sellers competing on their platform with a price recommendation for the product they are selling. We analyze if those non-binding pricing recommendations can be used to make the seller markets more collusive. As platforms often extract a share of the revenue from the sellers through sales commission rates, collusion among the sellers can be in the interest of the platform. We develop a stylized theoretical model and derive two rule-based algorithms that aim to make the seller market more collusive. Utilizing a laboratory experiment, we find that sellers condition their prices on the recommendation of the algorithms. The algorithm with a soft punishment strategy lowers market prices and has a pro-competitive effect. The algorithm that recommends a subgame perfect equilibrium strategy does not lead to higher average market prices compared to a benchmark without any price recommendations. However, it increases the range of market outcomes. As a result, specific markets are more collusive while others become more competitive. Variations in economic preferences lead to heterogeneous treatment effects and explain the results. The results suggest that algorithm-based non-binding price recommendations influence sellers' pricing behavior but that the effects are heterogeneous, and price recommendations would have to be well-targeted to make markets more collusive.

Chapter 3 presents a field experiment on volunteering in an online labor market. Volunteering is a widespread allocation mechanism in many workplaces. It emerges naturally in many work environments within the digital economy, like in software development or the generation of online knowledge platforms. In contrast to our theoretical predictions and previous research, we find no effect of team size on volunteering behavior. Using control treatments, we can show that workers generally react to free-riding incentives

provided by the volunteering setting but do not react strategically to the team size. We replicate our results and elicit workers' beliefs about their co-worker's volunteering. It reveals a form of conditional volunteering as the primary driver of volunteering, which can explain our results. Furthermore, we discuss the implications of those findings for organizations that rely on volunteering as a task allocation mechanism and highlight in which situations it may create inefficiencies.

Lastly, Chapter 4 focuses on the design features of "loot boxes". Loot boxes are digital goods that video game developers use to monetize their games. They are often designed as lotteries that gamers can buy and offer random rewards to be used in-game. Regulators are increasingly concerned that they might induce consumers to overspend on video games. We consider popular design features like opaque disclosure of the odds of winning and positively selected feedback in a stylized experimental framework. We find that those design features lead participants to overspend on the lotteries. In combination, these features double the average willingness-to-pay for lotteries. Furthermore, we highlight that the beliefs about the winning probability are an important channel that leads to this overspending. The findings emphasize the need to regulate the design of loot boxes to protect consumers from overspending.

Chapter 1

Algorithmic and Human Collusion

1.1 Introduction

The use of autonomous pricing algorithms is on the rise in various industries.¹ When firms use those tools, the pricing decision for a given product is outsourced from the human decision-maker to a computer algorithm. While in the past most pricing algorithms have been rule-based with rules defined by the seller, there is a recent evolution towards self-learning algorithms (Ezrachi and Stucke, 2017). These self-learning algorithms develop the strategies to achieve a specific goal, for instance, maximizing the firms' profits, without explicit instructions.

There are concerns among competition authorities (e.g., Bundeskartellamt and Autorité de la concurrence, 2019; Competition & Markets Authority, 2021) and academic scholars (e.g., Ezrachi and Stucke, 2016, 2017; Mehra, 2016) that pricing algorithms could not necessarily learn to price products more efficiently but also that there exists a possibility that they learn to collude tacitly.² In other words, algorithms could learn by themselves that tacit collusion benefits the firm.

While recent papers by Calvano et al. (2020a) and Klein (2021) show that algorithms can learn to be collusive, it is unclear whether pricing algorithms are more collusive than humans and, therefore, harm competition. Tacit collusion in traditional markets amongst human decision-makers is a well-documented phenomenon in both empirical and experimental economics.³ To assess the (anti-)competitive effects of algorithms, it is, therefore, necessary to establish a suitable baseline.

This chapter provides a counterfactual for algorithmic collusion for a wide range of possible market compositions and highlights the impact of algorithms on competition. To examine whether commonly used self-learning algorithms make markets more collusive relative to the status quo of human collusion, I apply a two-step approach. In the first step, I consider self-learning pricing algorithms in an extensive simulation study to test whether algorithms learn to set supracompetitive prices and suitable strategies to support

¹The European Commission (2017) finds that two-thirds of sellers in digital markets use pricing tools. Prominent examples are Amazon (Chen et al., 2016b; Musolf, 2022) or the gasoline market (Assad et al., 2022).

²For recent discussions about the possible policy and legal implications of algorithmic pricing see Kühn and Tadelis (2017), Schwalbe (2018), Harrington (2018), Calvano et al. (2019), Calvano et al. (2020b) and Assad et al. (2021).

³Empirical evidence is provided by, for instance, Borenstein and Shepard (1996), Davies et al. (2011), Miller and Weinberg (2017) and Byrne and De Roos (2019). See Engel (2015) and Horstmann et al. (2018) for experimental evidence.

those prices as a collusive outcome. Here, I closely follow the approach from Calvano et al. (2020a) but consider a different market environment that is more tractable. In the second step, I conduct market experiments in which humans compete either against each other or with self-learned pricing algorithms.⁴ In the experiments, I closely mimic the market environment from the simulations. Across different treatments, I vary the market composition between algorithms and humans and the number of firms in the market. The experimental approach allows me to consider tacit collusion and study the underlying mechanics in a controlled setup. My design enables me to observe humans and algorithms in the same environment and, thus, to analyze whether algorithms promote collusion.

I find evidence that algorithms foster tacit collusion in duopolies. Two-firm markets with algorithms are always more collusive than markets with humans. In “mixed” markets, in which humans and algorithms compete with each other, self-learned pricing algorithms are as good as humans when colluding with the other market participant. Hence, pricing algorithms never promote competition but foster collusion if all firms use one. In triopolies, there exists a non-linear relationship between the number of firms with a pricing algorithm and the level of tacit collusion. Markets in which a single firm uses a pricing algorithm are more competitive than markets with only humans. Yet, as more firms use pricing algorithms, market prices can increase and may even exceed prices in human markets, especially if humans are inexperienced. Similar to Calvano et al. (2020a) and Klein (2021), algorithms learn to punish price deviations. As I consider a stylized market environment, I can interpret the strategies of the algorithms. The most successful algorithms learn a win-stay lose-shift strategy that is common for the iterated prisoner’s dilemma. The outcomes in mixed markets have a large variance as humans choose heterogeneous strategies when playing against the algorithms.

While there exists reoccurring support for the hypothesis that algorithms can learn to set non-competitive prices and develop reward-punishment strategies (Klein, 2021; Calvano et al., 2020a, 2021; Abada and Lambin, 2022; Johnson et al., 2022; Asker et

⁴Many modern markets do not consist of just algorithms or only humans, but both can interact with each other in the same market environment. For example, according to Chen et al. (2016b), only one-third of the vendors selling the most popular products on amazon.com use some form of pricing algorithm, which gives rise to a mixture in market composition. Also, in the German gasoline industry Assad et al. (2022) identify local markets in which algorithms compete against firms in which arguably human managers make pricing decisions.

al., 2022), it is unclear how algorithmic collusion compares to human collusion.⁵ Market environments in previous studies on algorithmic collusion deviate substantially from the setting used in experimental market games. My design allows me to compare the outcomes of pricing algorithms to human pricing directly, as I can observe both in the identical market environment. A recent paper by Assad et al. (2022) identifies the adoption of algorithms in the German gasoline market. Within duopoly markets, price margins rise as both firms begin to utilize an algorithm. The effect is comparable to my findings for two-firm markets. For the market studied by Assad et al. (2022), the exact algorithms are unobservable as they are usually proprietary. The combination of simulations and laboratory experiments enables me to examine human and algorithmic strategies to study the underlying mechanics that may drive those effects.

This research also allows studying cooperation between humans and algorithms, which is a topic in computer science and experimental economics. In computer science, the design of cooperative algorithms in repeated games is an active research area (e.g., Crandall et al., 2018; Lerer and Peysakhovich, 2017). Here, cooperative algorithms are often the explicit objective. Similar to Calvano et al. (2020a), I consider a popular self-learning algorithm, which can be attractive as a pricing tool. While cooperation might be an outcome, it is not the initial design objective.

In experimental economics, on the other hand, the focus is often on deterministic algorithms, which do not learn themselves (see March (2021) for a recent literature review). Moreover, collusion in mixed markets with humans and algorithms is rarely studied. A notable exception is a recent paper by Normann and Sternberg (2023). They consider a tit-for-tat algorithm and find that three-firm markets with an algorithm are more collusive than classical human markets. The authors vary whether participants know if they play against a computer or a person and find no differences in this domain. My approach differs as I study self-learning algorithms and the limit strategies which they develop by themselves. Furthermore, my design allows analyzing the entire array of market composition as I can observe algorithmic and human markets as well as mixed markets. Hence, I

⁵Besides self-learned collusion, algorithms could also influence competition by offering better demand predictions (Miklós-Thal and Tucker, 2019; O'Connor and Wilson, 2021) or by serving as commitment devices (Brown and MacKay, 2022; Leisten, 2022). Furthermore, Harrington (2022) argues that outsourcing the development of algorithms to a third-party developer can affect market outcomes. For a general discussion of the possible effects algorithms can have on the economy, see Agrawal et al. (2019).

can directly compare algorithmic and human collusion and investigate the effect of pricing algorithms on a wide range of scenarios.

The remainder of the chapter is structured as follows. In Section 1.2, I discuss the main concepts of the algorithm, which I consider in this study. Then, in Section 1.3, I explain the market environment that I use in all simulations and experiments. After discussing the experimental design in Section 1.4 and the hypotheses in Section 1.5, I present my results in Section 1.6. Section 1.7 discusses the implications of my findings and concludes.

1.2 Pricing algorithms

Following the approach by Calvano et al. (2020a) and Klein (2021), I utilize Q-learning algorithms to study the collusive effects of self-learning pricing algorithms.⁶ Q-learning is a reinforcement learning algorithm designed to solve Markov decision processes with an ex-ante unknown environment (Watkins, 1989). In other words, Q-learning algorithms must learn everything about the environment alone and are not instructed to follow a particular strategy.

Many of the most successful state-of-the-art reinforcement learning algorithms build on the main ideas of Q-learning (e.g., Mnih et al., 2015; Silver et al., 2016; Arulkumaran et al., 2017). Therefore, it appears reasonable to assume that self-learning pricing tools also use some form of Q-learning. As Q-learning is still interpretable in contrast to more modern approaches, it makes it a natural choice when studying algorithmic collusion.⁷ In the following, I discuss some of the general concepts in Q-learning.⁸

⁶Earlier work by Waltman and Kaymak (2008) shows that Q-learning algorithms can converge to non-competitive quantities in a Cournot framework. However, they do not obtain collusion as algorithms also learn this behavior if they are memoryless. Hence, punishment strategies, which are essential for collusion to be sustainable in the long run, could never arise.

⁷While most of the literature on algorithmic collusion focuses on Q-learning algorithms, there are some exceptions. For instance, Hansen et al. (2021) study algorithmic pricing as a multiarmed bandit problem, in which each firm uses an Upper Confidence Bound Algorithm. Supracompetitive prices can arise in this setup as firms tend to run correlated experiments. Recent studies by Hettich (2021) and Jeschonke (2021) consider reinforcement learning algorithms that use function approximation methods.

⁸For a more detailed discussion of Q-learning, see Sutton and Barto (2018).

1.2.1 Basic Q-learning

Optimization problem In each period t , a Q-learning algorithm, often referred to as an agent, observes the current state $s_t \in \mathbf{S}$ of its environment and chooses some action $a_t \in \mathbf{A}$. Here, $\mathbf{A}$ is the set of feasible actions and $\mathbf{S}$ the set of possible states. Picking the action results in a reward signal $\pi_t \in \mathbb{R}$ and the next state $s_{t+1} \in \mathbf{S}$. The objective of the agent is to maximize the sum of discounted future expected rewards given the current state s_t over $\mathbf{A}$. The Bellman equation commonly expresses this maximization problem

$$V(s_t) = \max_{a_t} \{ \mathbb{E}[\pi_t | s_t, a_t] + \delta \mathbb{E}[V(s_{t+1}) | s_t, a_t] \} \quad (1.1)$$

with $\delta \in [0, 1)$ being the discount rate. The Bellman equation described by Equation 1.1 is recursive. The value of being in state s_t is given by the current reward signal π_t plus the discounted value of the continuation state s_{t+1} .

Conditional on having perfect knowledge over a stationary environment the agent is interacting in, Equation 1.1 could be solved using dynamic programming techniques. Yet, in Q-learning, the environment is typically unknown to the agent. Before learning, the agent does not know which actions result in which states or which state-action combinations lead to which rewards. Furthermore, the environment might be non-stationary because the same state-action combinations in period t may lead to a different reward and another continuation state in a different period t' . Thus, classical dynamic programming techniques, like recursively solving the Bellman equation, do not work given the unknown and possibly non-stationary environment.

In Q-learning, the Bellman equation is rewritten as the Q-function

$$Q(s_t, a_t) = \mathbb{E}[\pi_t | s_t, a_t] + \delta \mathbb{E}[\max_a Q(s_{t+1}, a) | s_t, a_t] \quad (1.2)$$

In this chapter, $\mathbf{A}$ and $\mathbf{S}$ are finite, and hence the Q-function is given by a $|\mathbf{S}| \times |\mathbf{A}|$ matrix, where $Q(s_t, a_t)$ represents the expected net present value of picking action a_t in state s_t .⁹ The goal of the Q-learning agent is to repeatedly interact with the environment to iteratively update the cells of its Q-matrix to obtain an approximation for the state-action value $Q(s_t, a_t)$ for each state-action combination.

⁹Note that the $V(s_t) \equiv \max_{a_t} Q(s_t, a_t)$.

In all simulations, the Q-matrix is initialized with random numbers drawn from a uniform distribution with support on the unit interval.¹⁰ For each subsequent iteration t , the agent picks some action a_t conditional on the current state s_t , which yields π_t and s_{t+1} . Then, the Q-matrix gets updated as the weighted average of the past estimate of $Q(s_t, a_t)$ and the newly learned value

$$Q_{t+1}(s_t, a_t) = (1 - \alpha)Q_t(s_t, a_t) + \alpha(\pi_t + \delta \max_a Q_t(s_{t+1}, a)) \quad (1.3)$$

where $\alpha \in (0, 1)$ is referred to as the *learning rate*. For the subsequent states, the same procedure is applied in each iteration until some convergence rule is met.

Exploration versus exploitation When selecting the action in s_t , the agent faces a trade-off. On the one hand, picking the action $a_t^* = \arg \max_a Q(s_t, a)$ yields the highest expected payoff given the current approximation of the Q-matrix. On the other hand, the agent can explore the environment only by picking new actions. By exploring the action space for each state of the Q-matrix the agent gradually learns the value of each state-action combination. As a result, the agent may find actions, which were underestimated beforehand. To balance exploration and exploitation, I use the ε -greedy algorithm. When using the ε -greedy algorithm, the current optimal action a_t^* is picked with probability $(1 - \varepsilon_t)$ with $\varepsilon_t \in [0, 1]$. With the probability ε_t the algorithm selects a random action from $\mathbf{A}$. I follow the approach by Calvano et al. (2020a) and Johnson et al. (2022) and define $\varepsilon_t = e^{-\beta t}$ for some small $\beta > 0$. Note that ε_t decays over time. At the beginning of the learning process, when the Q-matrix is rather uninformative about the value of picking an action in a specific state, the agent picks actions at random with a high probability and explores the action space. Overtime ε_t decreases and the agent chooses the action, which offers the highest expected long-run reward, more often. Eventually, the agent does not explore anymore but always picks the action with the highest expected value given that $\lim_{t \rightarrow \infty} \varepsilon_t = 0$.

¹⁰Klein (2021) and Abada and Lambin (2022) initialize the Q-matrix with zeros. Calvano et al. (2020a) and Johnson et al. (2022) use an initialization that corresponds to the discounted profit if all firms randomize their prices. Calvano et al. (2020a) show that outcomes are similar for different initialization of the Q-matrix.

1.2.2 Simulation setup

Learning and convergence In each simulation, the agents play against independent copies of itself and learn by interacting simultaneously with each other. For a stationary environment, Q-learning agents converge to the optimal solution under mild conditions.¹¹ For my design, this is, however, not the case, as multiple agents learn in the same environment at the same time. The environment of each agent is influenced by the decisions of other agents taking actions and learning at the same time by picking action simultaneously. Furthermore, the learning procedure when using the ε -greedy algorithm is stochastic for finite t . Given this constant change in strategies for the competitors, the environment is non-stationary, and thus convergence is not guaranteed. Therefore, I rely on simulations to derive results on algorithmic collusion in this non-stationary setup.

To determine the state of convergence in each training iteration, I follow the approach by Calvano et al. (2020a) and look at the cells of the Q-matrix. If the best action for each state does not change for 100,000 subsequent periods, I assume that the agents in the market have converged and found a stable strategy.

Uniqueness For a given environment, three hyperparameters define the learning process of the agent: α , β and δ . The learning rate α controls how much the agent values new information relative to the current approximation of the Q-matrix. For large values of α , the agent does not put much weight on past interactions with the environment. If α is small, the agent does not learn much from the newly arriving information in each period. The value of β captures the extent of exploration of the agent. The larger β , the faster ε_t converges to zero. Lastly, the discount rate δ captures the importance of future rewards. All three hyperparameters are not learned by the agents but are set exogenously. While there are some rules of thumb, their choice remains essentially unclear from a theoretical perspective. Given that an agent's outcome usually depends on those hyperparameters, the state of convergence is not unique for a given environment. This problem gets amplified further as the learning process of all agents is stochastic, which hinders uniqueness.

¹¹For a stationary Markov decision process, Q-learning converges to the optimal policy when using the ε -algorithm if the rewards signals are bounded. Furthermore, the learning rate must decrease over time, the sum of learning rates must diverge, while the squared sum of learning rates must converge. For a proof see Watkins and Dyan (1992).

In the entire study, I keep the discount rate fixed at $\delta = 0.95$. It corresponds to the continuation probability in the market experiments with human participants described in Section 1.4.2. For α and β , I consider a parameter grid with $\alpha \in [0.025, 0.25]$ and $\beta \in [1 \times 10^{-8}, 2 \times 10^{-5}]$ with 100 points in each dimension evenly spaced from one another. For each grid-point, I simulate 1,000 distinct markets which differ in the underlying stochastic process. Hence, I run in total 10,000,000 simulations for each market size. The grid and the simulation setup, again, follows Calvano et al. (2020a) and Johnson et al. (2022). I evaluate the results from this grid simulation using different performance measures, which I describe below. Furthermore, I use the grid to find a specific algorithm that competes with humans in the experiments.

1.2.3 Performance evaluation

Profitability and prices Given the definition of the Bellman equation, $V(s)$ provides an unbiased estimate of the expected sum of future discounted rewards for a given state s .¹² Thus, $V(s)$ is a natural measure for the profitability for agent i in a given state s . I utilize this direct interpretability and use $V(s)$ as a profitability measure. Additionally, I consider the (average) market price upon convergence as a way to measure tacit collusion. Market prices are of particular interest when comparing outcomes in algorithmic to experimental markets with humans as the value function is unobserved here.

Optimality High profitability alone does not necessarily imply that the agent has learned an optimal strategy against its competitors. To derive a measure of how well the agent does in comparison to how well it could do, I derive the optimality of the agent for state s as

$$\Gamma_i(s) = \mathbb{1}\{\arg \max_a Q_i(s, a) = \arg \max_a Q_i^*(s, a)\} \quad (1.4)$$

where $Q_i^*(s, a)$ is the optimal Q-matrix assuming that the opponents play their limit strategy. Here, $\mathbb{1}$ denotes the indicator function.

I can obtain this optimal Q-matrix in the following way: After convergence, the strategies of the other agents, which learned in the same environment, are held con-

¹²I evaluate the performance of the agent only upon convergence. Therefore, for ease of notation, I omit the period subscript t .

stant. Thereby, the environment becomes stationary as the state-transition probabilities do not change anymore. Next, I initialize a new agent, which competes against the limit strategy of the opponents. I utilize dynamic programming to repeatedly iterate over each state-action combination of its Q-matrix until convergence. Within this stationary environment, convergence is guaranteed. The Q-matrix of the new agent corresponds to the optimal Q-matrix $Q_i^*(s, a)$ that the actual agent could have learned against the limit strategy of the other agents.

Note that $\Gamma_i(s) = 1$ implies that the agent has learned to play a Nash equilibrium for state s . The agent has learned a subgame perfect Nash equilibrium if and only if $\bar{\Gamma} = \frac{1}{|\mathbf{S}|} \sum_s \Gamma_i(s) = 1$. For $\bar{\Gamma} < 1$, there are states for which the agent did not learn to play a Nash equilibrium.

Selecting a specific algorithm I study experimental treatments in which humans compete against algorithms in the same market. To conduct laboratory experiments in those “mixed” markets, I have to decide on a specific parameterization of the agent with a specific stochastic process. To select an algorithm for this purpose, I take the perspective of a firm that would want to deploy a pricing algorithm to a market. It is reasonable to assume that this firm would want the pricing algorithm to be (i) profitable and (ii) optimal in the sense that it is not easily exploitable by other market participants. Considering both objectives in isolation is not sufficient when selecting an agent for the deployment to the market. Importantly, high profitability does not necessarily imply high optimality or vice versa. It is vital to rule out that a lack of sophistication of the algorithms drives the high profitability.¹³ For instance, the agents in the environment could jointly converge to a seemingly collusive outcome in which they price at the monopoly price but fail to learn a strategy that accounts for certain deviations. Such a myopic strategy is unlikely to perform well against new competitors, which would result in a lower profitability than in the simulation. When both performance measures are maximized jointly, the agent has learned a profitable strategy that also accounts for the strategic element of the environment. Accordingly, it increases the likelihood that the algorithms perform well against new (possibly human) competitors. Therefore, I propose the following criterion for agent

¹³The situation can arise, for example, if the exploration of the agent is limited.

i that combines both performance measures

$$\Psi_i = \frac{\bar{\Gamma}_i}{|\mathbf{S}|} \sum_s V_i(s). \quad (1.5)$$

Since $\bar{\Gamma}_i \in [0, 1]$, the selection criterion is the average profitability over the entire state space shrunk towards zero by the degree of suboptimality. Hence, $\bar{\Gamma}_i$ can be interpreted as a shrinkage penalty in this context. The intuition is that high profitability is only valuable if other players cannot exploit the agent easily with a possibly more sophisticated strategy. In the simulation, algorithms usually converge to a specific state. Within the experiments with humans, different states are potentially relevant as human and algorithmic strategies can differ. For that reason, I consider the entire state space and not only the state of convergence when defining Ψ_i . For treatments in which algorithms interact with humans, I select the algorithm from the simulation in which $\bar{\Psi} = \frac{1}{N} \sum_i^N \Psi_i$ is maximized over the parameter grid for α and β . I will refer to it as the selected algorithm.

In mixed markets, the algorithms learn “offline” in advance. Put differently, they develop their strategies in a simulated market environment against other algorithms. Once they interact with humans in the actual experimental market, they do not learn any more from new market information but use the strategy obtained during the training in the simulation.¹⁴ While this approach might appear restrictive, it is arguably realistic. Q-learning is a slow learning algorithm as it updates only one cell of the Q-matrix in each training step. Furthermore, each cell has to be visited by the agent multiple times to obtain an accurate estimate of the value of this state-action combination. As a result, Q-learning algorithms usually learn too slow to be trained “online”, meaning in the actual market environment. Furthermore, this learning strategy has been used in other groundbreaking Q-learning applications and is a common standard for successful reinforcement learning applications.¹⁵

Also, from an industry perspective, offline training makes intuitive sense. A firm using a pricing algorithm would likely want to evaluate its performance before deploying it to the market to avoid any possible loss, given a potential suboptimal pricing strategy. The

¹⁴Thus, I do not consider how humans and algorithms may interact during the learning process. Yet, it can be an interesting avenue for future research.

¹⁵For instance, AlphaGo, which is a reinforcement learning algorithm that outperforms humans in the board game Go, uses offline learning (Silver et al., 2016). For a discussion of offline learning, also see Calvano et al. (2020a).

risk that the agent learns a suboptimal strategy can be partially mitigated by offline training as an ex-ante evaluation is feasible.¹⁶

1.3 Market environment

I consider a stylized Bertrand market environment, which is commonly used in the experimental economics literature on collusion (see, for instance, Fonseca and Normann, 2012; Horstmann et al., 2018).

There are $N \in \{2, 3\}$ firms in the market, which face a perfectly inelastic demand function and have zero marginal costs. Each firm produces the same homogeneous good. The market consists of $m = 60$ computerized consumers, which are all willing to purchase exactly one unit of this good in each round and have a maximum willingness to pay of $\bar{p} = 4$. The price of firm i in period t is denoted by $p_t^i \in \mathcal{P} := \{0, 1, 2, \dots, 5\}$. Consumers buy the good at the lowest offered price. If multiple firms offer the lowest price in a given round, the market is shared equally. Firms are always either represented by a human or by a Q-learning algorithm. This market environment is the same for the simulation study and all experimental treatments. It allows me to directly compare the simulation and outcomes and derive a counterfactual for algorithmic collusion. In the simulation treatments, firms compete in an infinitely repeated game with a discount rate of $\delta = 0.95$. To mimic the features of an infinitely repeated game in the experimental treatments, I use a repeated game with random stopping, where the continuation probability for playing another round is given by 95% (Roth and Murnighan, 1978). While this environment is less complex than many actual markets, it yields a suitable setting for my design as it distills the main components of price competition when studying collusion.

There exists a stage game Nash equilibrium at $p^{NE} = 1$. The monopoly price of $p^M = \bar{p} = 4$ maximizes joint profits.¹⁷ When all firms charge the same price, $\pi_t^i = p m / N$ gives the profit. The profit for a single deviating firm is $\pi_t^i = p m$. Collusion is sustainable at the monopoly price for the given discount factor, for instance, by grim-trigger strategies.

¹⁶When using offline training for pricing tools that build on reinforcement learning algorithms, the market environment has to be known to the developing firm. However, with modern tools in demand estimation and supervised machine learning, it is reasonable to assume that firms can derive this environment.

¹⁷There are two prices ($p = 5$ and $p = 0$) which are (weakly) dominated. I include both in the set of possible prices $\mathcal{P}$ to rule out that convergence to the boundaries of the price set is equivalent to collusion or competition at the stage game Nash equilibrium.

Crucially, the environment is stylized and tractable, which allows for an analysis of the strategies algorithms learn in the game. Furthermore, it is arguably easy to understand for experimental subjects due to its simple mechanics. Also, the environment offers a different extension to algorithmic price competition to markets with a perfectly inelastic demand function, which has not been studied before.

Q-learning and the market environment If a firm uses a Q-learning algorithm, the algorithm takes over the pricing decision in each period. Hence, the action a_t of the agent corresponds to a price and the set of possible actions corresponds to the price set. The economic profit obtained in each respective period is the reward signal for the Q-learning algorithms.

Similar to Calvano et al. (2020a) and Johnson et al. (2022), I define the state of the environment for each agent as the set of past prices from the previous period $s_t = \{p_1^{t-1}, \dots, p_N^{t-1}\}$. Notably, this state representation corresponds to memory-one strategies, which humans predominantly use in the prisoner’s dilemma and market games (Dal Bó and Fréchette, 2019; Romero and Rosokha, 2018; Wright, 2013). Calvano et al. (2020a) also consider state representations that allow for a two-period memory. This larger memory does not improve the collusive abilities of the algorithms. Importantly, it is straightforward to construct memory-one strategies that make collusion incentive-compatible in two and three-firm markets.¹⁸

1.4 Experimental design

1.4.1 Treatments

I consider five experimental treatments and two treatments based on simulations. Across treatments, I vary the market composition between algorithms (A) and humans (H) and the number of firms in the market (see Table 1.1). I label the treatments with the number of human firms followed by the number of firms that use an algorithm. For example, treatment 2H1A stands for two human players and one algorithmic player operating in

¹⁸For instance, consider a memory-one strategy that mimics a grim-trigger strategy in the sense that the agent always plays the monopoly price in the state $s = (p^M, p^M, p^M)$ but chooses the stage game Nash equilibrium in any other state. This strategy is a possible outcome for the simulations from an ex-ante perspective and makes collusion sustainable for the given discount factor of $\delta = 0.95$.

one market. Thus, I consider treatments without any algorithms (2H0A and 3H0A) and

Table 1.1: Treatment composition

Number of Human Players	Number of Algorithms			
	0	1	2	3
3	3H0A	-	-	-
2	2H0A	2H1A	-	-
1	-	1H1A	1H2A	-
0	-	-	0H2A	0H3A

without any humans (0H2A and 0H3A).¹⁹ Comparisons between those treatments reveal whether algorithms are more collusive than humans. Additionally, I consider treatments in which humans compete against algorithms (1H1A, 1H2A and 2H1A). I utilize those treatments to examine the way humans and pricing algorithms interact with each other. Furthermore, they show if an increase in the share of algorithms in the market fosters tacit collusion for different market sizes.

1.4.2 Procedure

Each experimental treatment is repeated for three supergames to observe learning effects. Within each supergame, there is a fixed group composition. Across supergames, I use a perfect stranger matching scheme. This matching scheme is common knowledge. Hence, participants know they interact with each person only within one supergame throughout the entire experiment. It rules out any reputation effects that could be present elsewhere.

In my experiment, each round has a continuation probability of 95% in each supergame. Hence, with a 5% chance, a given supergame ends after each respective round. The instructions mention the continuation probability to the subjects at the beginning of each supergame. It corresponds to the discount rate of $\delta = 0.95$ that is used for the algorithms in the simulation treatments.²⁰ To allow for different experimental sessions with the same supergame lengths, the round numbers are pre-drawn with a random number generator.²¹ At the end of each round, participants receive information about all prices in the market. Furthermore, they see their own profit in the given round.

¹⁹Note that the latter are simulation studies.

²⁰For a risk-neutral player the continuation probability is theoretically equivalent to the discount rate (see for instance Roth and Murnighan, 1978; Dal Bó, 2005).

²¹The exact round numbers are 25 (supergame 1), 17 (supergame 2) and 11 (supergame 3).

Each participant has complete information on which firms use a pricing algorithm. While this is arguably not the case in actual markets, Normann and Sternberg (2023) show that participants are insensitive to the knowledge of playing against an algorithm or a human player.

Following Normann and Sternberg (2023), all profits obtained by a firm that uses a pricing algorithm are given to a passive human player, who does not take any active decision. It rules out any differences in social or distributional preferences that might arise elsewhere across treatments.

The framing regarding the algorithmic decisions in the experiment is neutral. In all treatments, participants do not know the objective of the algorithm, they do not know that it is self-learned or how it learned its strategy.²² Subjects receive the instructions at the start of the session, but they are also available during the experiment at any point. After the participants read the instructions, I ask them a set of control questions.²³ If a participant gives three times a wrong answer, I show an additional explanation for the respective question. One person dropped out of the experiment due to technical problems. I exclude the entire matching group of this subject from the analysis.

As described in Section 1.3, the algorithms always condition their current pricing decision on a state which is the set of prices from the previous period. In the first round of each supergame, the algorithm has no state to condition upon as the state $s_{t=0}$ is undefined. To circumvent this initial condition problem, I define $s_{t=0}$ as the state of convergence from the learning process. Thus, the algorithm always begins each supergame with the same action it played last in the simulated environment.

The experiments were conducted online in May and June 2021. I recruited the participants using ORSEE (Greiner, 2015) from the subject pool of the DICE Lab, University of Düsseldorf. A web-conference call accompanied each session in which participants could ask clarifying questions and receive technical assistance if required.²⁴ In total, 313 participants were recruited, with between 60 to 64 subjects in each experimental treatment (see

²²It mimics the information structure in actual markets, in which firms do not know much about the algorithms of other market participants. Yet, varying the information participant get about the objective and the learning process of the algorithm can be an interesting path for future research.

²³See Appendix 1.A for the full set of instructions, the control questions, and screenshots of the relevant decision screens.

²⁴The procedure is similar to Zhao et al. (2020) and Danz et al. (2021). Li et al. (2021) find that this procedure offers comparable results to lab experiments in different economic games.

Table 1.2). For each treatment without any humans, I use 1,000 independent simulation runs for each parameterization of the algorithms as the respective comparison unit.

Table 1.2: Number of observations by treatment

Treatment	Number of participants	Number of independent observations
3H0A	63	7
2H0A	60	10
1H1A	64	32
1H2A	63	21
2H1A	63	7
0H2A	-	1,000
0H3A	-	1,000

* The number of independent observations for the experimental treatments refers to the last supergame. It is determined by the size of the matching group. A matching group consists of six (nine) firms for markets with two (three) firms. Since the algorithms do not learn anymore during the experiment and the respective participants do not take any active decisions, there are no reputation effects across super games. Hence, I exclude the algorithmic firms from the matching scheme.

Each session lasted for approximately 30 minutes, and subjects earned on average 11.3 Euro, including a show-up fee of 4 Euro. Within the experiment, I used an experimental currency unit (ECU) where one Euro corresponded to 130 ECU. The experiment employed a between-subject design and thus, each subject participated only in one treatment. I programmed the experiment in oTree (Chen et al., 2016a).

1.5 Hypotheses

The experimental design allows for the comparison of human and algorithmic collusion for different market compositions. Furthermore, I can investigate the interaction between humans and pricing algorithms if both populate the same market. I use the respective market price to measure the degree of tacit collusion. Within the experimental treatments, an independent observation is the average market price for a given matching group. For the simulations, I average the market prices for each independent simulation over 1,000

rounds after convergence. All hypotheses and the corresponding statistical tests have been pre-registered.²⁵

From a theoretical perspective, punishment strategies are vital for collusion to be sustainable in the long run (Friedman, 1971; Abreu, 1988). While humans often fail to employ punishment strategies that appear desirable from the theoretical perspective²⁶, Calvano et al. (2020a) and Klein (2021) find that self-learned pricing algorithms learn harsh punishment strategies that make collusion incentive compatible. I expect the pricing algorithms in my design to learn comparable punishment strategies, which would theoretically foster collusion compared to lenient strategies often used by humans.

Algorithms may also reduce the strategic uncertainty within the game. After convergence, algorithms follow the same strategy for all of the subsequent rounds. In the mixed market treatments, the algorithm plays according to their limit strategy against humans. Hence, they play the strategy that they learned in the simulations. Crucially, after the learning process of the algorithms, their strategy is deterministic and does not change anymore. Normann and Sternberg (2023) argue that playing against a deterministic algorithm reduces strategic uncertainty compared to playing against a human, who might change the strategy during the game and are less committed to a particular behavior. They demonstrate theoretically that the postulated reduction in strategic uncertainty fosters collusion. In line with this prediction, they show that in experimental market games where humans compete against a tit-for-tat algorithm, markets become more collusive compared to markets with only humans. As the algorithms in my design are also deterministic after convergence, I expect higher degrees of tacit collusion in mixed markets relative to human markets. Since humans first have to learn about the algorithm's strategy, it appears natural that this effect especially materializes in later supergames. Harsher punishment strategies and reduced strategic uncertainty by algorithms should foster collusion. Thus, I hypothesize that market prices increase as more firms use a pricing algorithm for a given market size.

Hypothesis 1. *The level of tacit collusion increases in the share of firms using self-learned pricing algorithms for a given market size.*

²⁵The pre-registration uses the template from *AsPredicted.org* and can be found here <https://osf.io/yd32b> and here <https://osf.io/uxdcp>. Ethical approval was granted by the German Association for Experimental Economic Research e.V. (No. vzRbKXHq).

²⁶For a discussion of the strategies that humans use in experimental market games see for instance Wright (2013).

It is a well-documented finding in the literature on experimental market games that tacit collusion becomes less likely as the number of firms in the market increases (Engel, 2007; Huck et al., 2004; Harrington et al., 2016). Within my design, a larger market size implies higher deviation profits. That, in turn, increases the incentive to deviate from a collusive price level. Furthermore, the strategic complexity of the game grows as the number of firms increases. With more firms in the market, market participants have to condition their behavior on additional factors such as the previously chosen prices from the extra competitor. This increase in strategic complexity may further hinder collusion.²⁷ Similar to the findings in experimental market games, Calvano et al. (2020a) and Johnson et al. (2022) find decreasing prices in algorithmic markets in their simulations as the number of firms increases. I expect comparable results in my experimental design, which leads to the following hypothesis.

Hypothesis 2. *The level of tacit collusion decreases in the number of firms in the market for human and algorithmic markets alike.*

It needs to be clarified how those number effects differ between algorithmic and human markets. While the decline in market prices in the previous studies on algorithmic collusion appears smaller than in human markets, the market setup deviates substantially from the environments usually used in experimental market games. Hence, it is an open question whether algorithms are better at colluding as the market size increases. I investigate this question in the following sections.

1.6 Results

This section discusses the results and examines whether algorithms foster tacit collusion for different market compositions. I begin by considering the performance measures for the algorithms in Section 1.6.1. Then, I discuss the exact strategies that the selected algorithms learn upon convergence in Section 1.6.2. In Section 1.6.3, I investigate how algorithmic collusion compares to human collusion. Lastly, I discuss the results on mixed market compositions in Section 1.6.4 to shed light on the interaction between pricing algorithms and human decision-makers if both populate the same market at the same time.

²⁷For a discussion on the influence of strategic complexity on cooperation see Jones (2014) and Gale and Sabourian (2005).

1.6.1 Performance of the algorithms

In the following, I discuss the performance of algorithm in the simulation treatments without any humans for duopolies (0H2A) and triopolies (0H3A).

Profitability Figure 1.1 shows the profitability of the algorithms in the state of convergence²⁸ for a parameter grid over the learning rate α and the exploration decay β for the simulations 0H3A and 0H2A.²⁹ As a benchmark, I provide the value function under collusion V^C at the monopoly price and under competitive pricing V^{NE} . Those can be derived by considering the individual fixed profit π in those states and then rewriting the Bellman equation as the arithmetic series $V = \frac{1}{1-\delta}\pi$. Lighter colors in Figure 1.1 show a profitability close to collusion at the monopoly price, while darker colors indicate that the algorithms are not profitable.

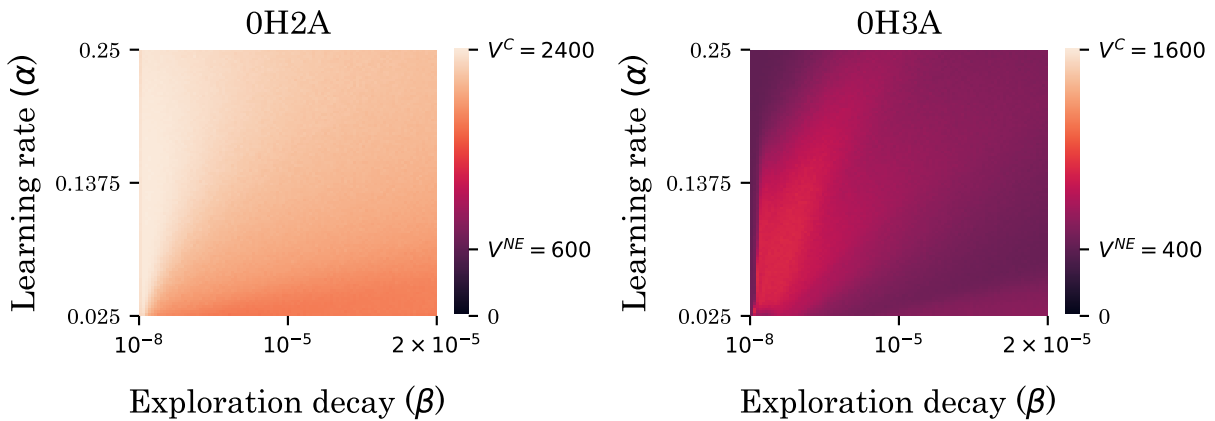


Figure 1.1: Profitability of the Q-learning agents in the state of convergence averaged over 1,000 simulation runs for different grid points.

In both treatments, the algorithms have a high profitability after convergence. For most grid points, the algorithms learn a strategy that is more profitable than with competitive pricing. This becomes evident by comparing the value function for the grid points in Figure 1.1 to V^{NE} . Notably, the profitability is usually greater for 0H2A compared to 0H3A. This pattern is also confirmed when considering the average prices the algorithms play upon convergence shown in Figure 1.2.

²⁸While it is not guaranteed, the algorithms always converge using the respective convergence criterion defined in Section 1.2.2.

²⁹Computational support and infrastructure was provided by the “Centre for Information and Media Technology” (ZIM) at the University of Düsseldorf (Germany).

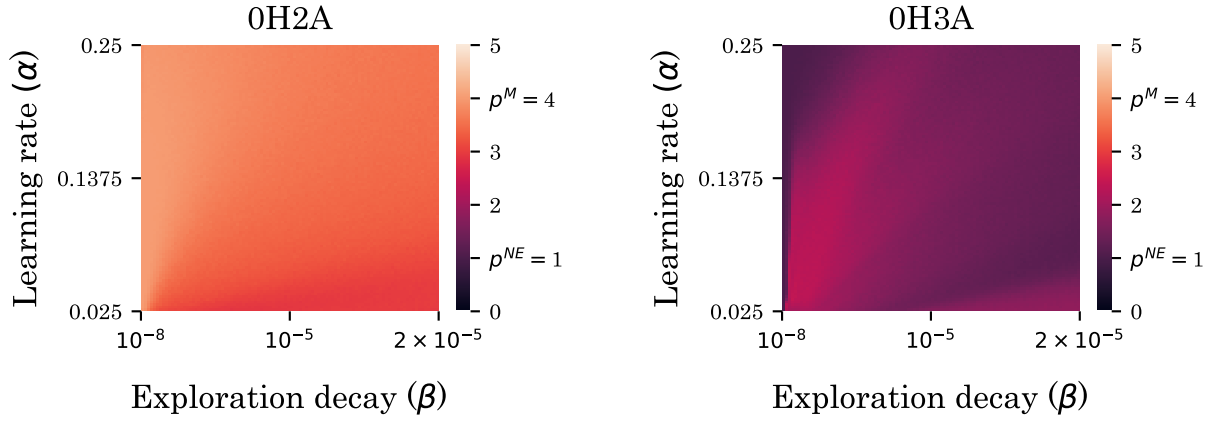


Figure 1.2: Average market price of the Q-learning agents after convergence averaged over 1,000 simulation runs for different grid points. This market price average is obtained by considering 1,000 periods after convergence.

On average, the algorithms learn to set non-competitive prices for a wide range of parameterizations. For each grid point, the average market price is above the stage game Nash equilibrium in 0H2A. It is also the case for 99.2% of all grid points in 0H3A. Hence, in both treatments, market prices are above the competitive benchmark. Notably, the market prices in 0H2A are, on average, higher than 0H3A. While in 0H2A, the average market price is above $p = 3$ for more than 93.4% of all grid points, it never exceeds $p = 2.4$ in 0H3A. Indeed, for each grid point, the average price in 0H2A is statistically significantly higher than in 0H3A (Two-sided Mann–Whitney U test, $p < 0.01$ for each grid point separately). Accordingly, the level of tacit collusion is higher in two-firm algorithmic markets. This result is in line with Hypothesis 2 and previous findings on algorithmic collusion.

Result 1. *Algorithms learn to set non-competitive prices in two and three-firm markets. In these markets, tacit collusion is significantly higher with two than with three firms.*

Optimality Figure 1.3 shows the share of all simulation runs in which both algorithms converge mutually to a Nash equilibrium.

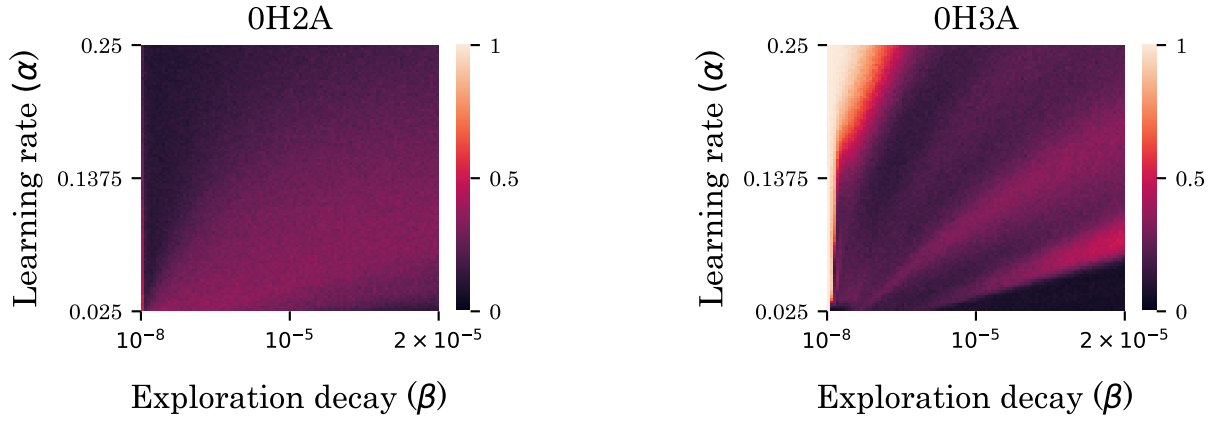


Figure 1.3: Share of simulations that converge to a Nash equilibrium.

On average, learning to play a Nash equilibrium appears difficult for the algorithms in 0H2A and 0H3A. While the algorithms manage to play a mutual best response for specific parameterizations, the optimality measure is below one for most grid points. Also, compared to previous findings by Klein (2021) and Calvano et al. (2020a), the share of outcomes converging to a Nash equilibrium seems smaller. The market environment with a perfectly inelastic demand function proves to be challenging for the algorithms. A possible reason for this is that small price changes lead to drastic shifts in profits, which may hinder a smooth convergence to an equilibrium strategy.

1.6.2 Strategies of the algorithms

Next, I examine the limit strategies that the algorithms learn once they converge. I focus on the parameterizations of the algorithm that perform best, given the selection criterion discussed in Section 1.2.3.³⁰ Thus, I concentrate on the algorithm that is likely to have high profitability while being harder to exploit by other market participants. It appears reasonable to assume that a firm would select such an algorithm when deploying a pricing tool to an actual market environment.

Punishment behaviour Next, I examine the limit strategies that the algorithms learn once they converge. I focus on the parameterization of the algorithm that performs best, given the selection criterion discussed in Section 1.2.3.³¹ Thus, I concentrate on the

³⁰For 0H2A those parameters are $\alpha \approx 0.027$ and $\beta \approx 1 \times 10^{-8}$. The values for 0H3A are $\alpha \approx 0.029$ and $\beta \approx 6.16 \times 10^{-7}$.

³¹For 0H2A those parameters are $\alpha \approx 0.027$ and $\beta \approx 1 \times 10^{-8}$. The values for 0H3A are $\alpha \approx 0.029$ and $\beta \approx 6.16 \times 10^{-7}$.

algorithm that is likely to have high profitability while being harder to exploit by other market participants. It appears reasonable to assume that a firm would select such an algorithm when deploying a pricing tool to an actual market environment.

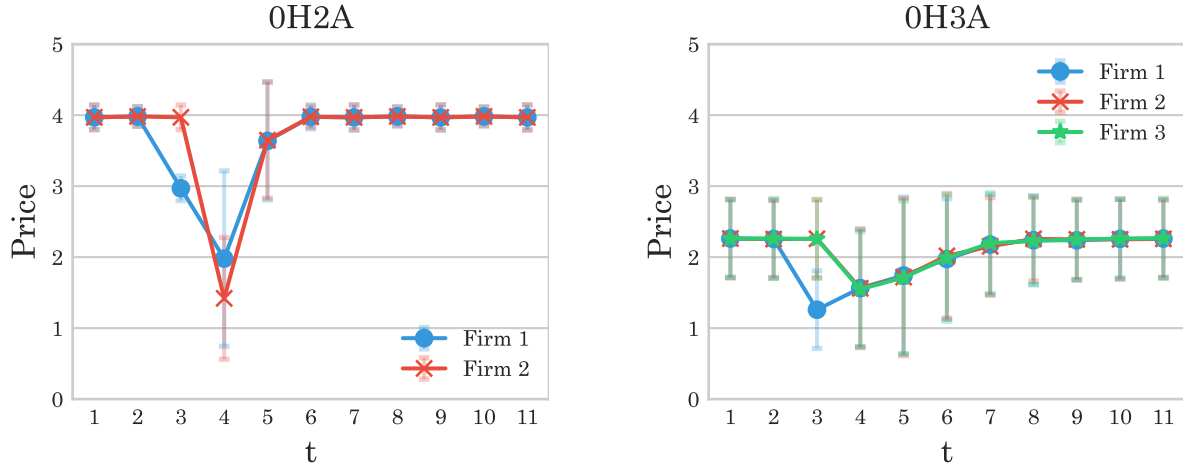


Figure 1.4: Punishment behavior of the algorithms after convergence. Starting from the state of convergence, the algorithms play according to their limit strategy. I induce an exogenous deviation from Firm 1 in $t = 3$ to observe the reaction of the other firms. I use 1,000 independent simulation runs. The error bars represent the standard deviation.

The left-hand side of Figure 1.4 shows the behavior of the algorithms in 0H2A after a deviation by one of them. The algorithms play according to their limit strategy in the first two periods. Then, in period $t = 3$, I force Firm 1 to deviate by undercutting the competitor's price. The deviation price is always just below the price it would have played according to its limit strategy. Thus, Firm 1 always chooses the most profitable one-period deviation.³² Afterwards, I allow both algorithms again to play according to their limit strategy to observe their response to the deviation.

In the initial two periods, algorithms choose prices close to the monopoly level. After the exogenously induced deviation of Firm 1 in the third period, Firm 2 lowers its price below the price of Firm 1. Then, after this phase of lower prices, both firms revert to the initial price level within the next couple of periods. The behavior of the algorithms is consistent with a punishment scheme. By undercutting the price of Firm 1 after the deviation, Firm 2 makes the deviation from the initial price level unprofitable. Indeed, for 81.4% of all simulation runs, algorithms learn a limit strategy that makes deviations, as shown in Figure 1.4 unprofitable. Thus, algorithms do not only learn to play high prices but also strategies that make collusion incentive compatible.

³²In more than 99% of simulation runs, the algorithms learn to play a price above the stage game Nash equilibrium.

The right-hand side of Figure 1.4 shows the results for the same exercise for 0H3A. As discussed in Section 1.6.1, prices are significantly lower in three-firm markets compared to two-firm markets. Similarly to 0H2A, the other firms in the market decrease their price after the deviation by Firm 1. After a punishment phase of multiple periods, the algorithms return to the initial price level they played before the deviation.³³ The punishment behavior of the algorithms is incentive compatible in 62.7% of all simulation runs.

Result 2. *Algorithms in 0H2A and 0H3A learn punishment strategies that can make collusion incentive compatible.*

Limit strategy of the selected algorithm In the previous sections, I considered the average behavior of the algorithms with a fixed parameterization over multiple simulation runs with different underlying stochastic processes. Next, I present the exact limit strategy of the algorithm which maximizes the selection criterion Ψ as discussed in Section 1.2.3 in 0H2A and 0H3A. It is also the strategy the algorithm will use within the experimental treatments with humans.

Equation 1.6 describes the core idea of the strategy. It is nearly identical for the algorithm in two and three-firm markets.³⁴

$$p_i^t(s^t) = \begin{cases} p^M & \text{if } s^t = \{p_i^{t-1} = p^M | \forall i\} \\ p^M & \text{if } s^t = \{p_i^{t-1} = p^{NE} | \forall i\} \\ p^{NE} & \text{otherwise} \end{cases} \quad (1.6)$$

Upon cooperation at the monopoly price p^M in the previous period, the algorithm always chooses the monopoly price again. Any deviation from the cooperative outcome is punished by playing the stage-game Nash equilibrium p^{NE} . If and only if all firms played p^{NE} in the previous period, the algorithm reverts to playing p^M . In every other relevant state,

³³While the average punishment strategy appears like a smooth transition between periods of punishment and cooperation, it is usually not the case for each simulation in isolation. Large and sudden price jumps after deviations are common. The transition only appears smooth when averaged over all simulation runs.

³⁴There are minor differences between the strategies in 0H2A and 0H3A. Namely, in 0H3A, a small number of states trigger a different response by the algorithm after deviations from the monopoly price. For instance the state $s_t = (4, 3, 0)$ leads to $a_t = 4$ or $s_t = (4, 4, 2)$ yields $a_t = 3$. However, those states are never reached after the algorithms converge. Furthermore, those states only account for approximately 1% of all rounds in mixed market experiments. Additional details are provided in Appendix 1.B.2.

the algorithm plays p^{NE} .³⁵ While algorithms in 0H2A learn this strategy frequently, it only arises occasionally in 0H3A, which is also indicated by the overall lower price level in this treatment.

Interestingly, this strategy is similar to the win-stay, lose-shift strategy (WSLS) discussed by Nowak and Sigmund (1993) in the context of the iterated prisoner's dilemma. Whenever an agent uses WSLS in the iterated prisoner's dilemma, she conditionally cooperates. Upon any deviation, the agent defects and reverts back to cooperation if and only if both players defected in the previous period. WSLS has several desirable properties from an (evolutionary) game-theoretical perspective.³⁶ If actions are noisy, WSLS can correct for unintended deviations when playing with another agent that uses WSLS. That is not the case for other popular strategies like tit-for-tat. Furthermore, WSLS can detect and exploit unconditional cooperators after unintended deviations, which may arise if the action implementation is noisy. However, depending on the exact payoff structure, agents that always defect can exploit WSLS. Nowak and Sigmund (1993) show that WSLS arises naturally as the most widespread strategy in an evolutionary simulation in a noisy iterated prisoner's dilemma.³⁷

The algorithm's strategy is an application of WSLS to the market environment. Similar to WSLS in the iterated prisoner's dilemma, the selected Q-learning algorithm restricts its attention to two actions: Cooperation at p^M and defection at p^{NE} . Just as the classical WSLS, the algorithm only cooperates if all firms in the market jointly played p^M or p^{NE} in the previous round and defects otherwise.³⁸

Importantly, the WSLS strategy does not arise by construction but as a result of the learning procedure of the algorithms. Even with a state representation that is restricted to the prices of the previous period, it would be straightforward to construct strategies that

³⁵This refers to all possible states that are reachable given the limit strategy of the algorithm. Thus, I do not consider states requiring the algorithm to play prices that it never plays itself when following its limit strategy.

³⁶For a discussion of win-stay, lose-shift also see Posch (1999) and Imhof et al. (2007).

³⁷For Q-learning algorithms, actions are also implemented with noise during the learning process as the exploration of the environment is stochastic. Hence, there might exist a possible relation between the evolutionary processes that lead to WSLS in simulation studies and the learning of Q-learning agents.

³⁸Calvano et al. (2019) find that Q-learning algorithms often learn similar one-period-punishment strategies in the iterated prisoners dilemma.

punish deviations for more than one period.³⁹ However, the algorithms do not coordinate on strategies that outperform WSLS.

Result 3. *The algorithm that maximizes the selection criterion learns a win-stay lose-shift strategy.*

While this strategy can correct unintended deviations and punish intended deviations by other firms, it is also possible to construct strategies that exploit the algorithm. As an example, consider a three-firm market where two firms use the strategy of the algorithm described by Equation 1.6 and firm k uses the following strategy

$$p_k^t(s^t) = \begin{cases} p^D = 3 & \text{if } s^t = \{p_i^{t-1} = p^{NE} | \forall i\} \\ p^{NE} & \text{otherwise} \end{cases} \quad (1.7)$$

Given that the algorithm always plays p^M after all firms played p^{NE} in the previous round, the strategy by firm k triggers the algorithms to cooperate every other round only to exploit their cooperative phase by choosing the most profitable deviation p^D . It is straightforward to show that in an infinitely repeated game with $\delta = 0.95$, the strategy of firm k strictly dominates cooperation at the monopoly price in three-firm markets. However, this strategy is dominated by always cooperate for two-firm markets.⁴⁰ While Q-learning algorithms never learn the strategy described by equation 1.7 during their simultaneous and dynamic learning process, I will analyze if humans manage to exploit the limit strategy of the algorithm in mixed markets in Section 1.6.4.

1.6.3 Comparing algorithmic and human collusion

While algorithms converge to non-competitive prices and learn punishment strategies, it is unclear whether algorithms are more collusive than humans. Therefore, in this section, I compare the market outcomes from the experiments with humans to the algorithmic markets. Figure 1.5 shows the average market prices by supergames (SG) for the treat-

³⁹A simple example is a strategy that mimics the behavior of a grim-trigger strategy. Yet, also other strategies are feasible. For instance, consider strategies that do not revert to the monopoly price immediately after playing the stage game Nash equilibrium but to intermediate values of the price set.

⁴⁰For details see Appendix 1.B.

ments with only humans (2H0A and 3H0A) and outcomes for the selected algorithmic markets (0H2A and 0H3A).

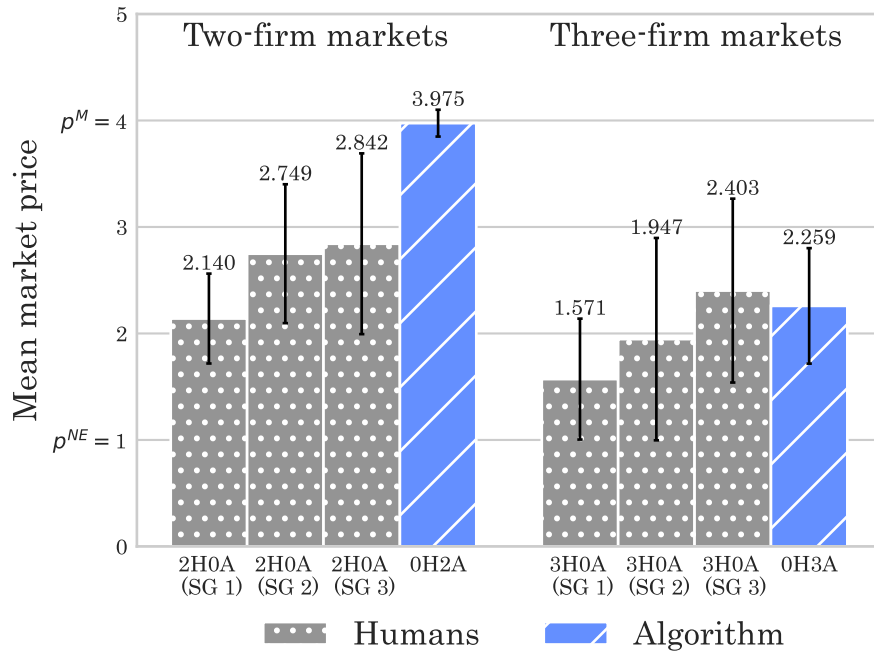


Figure 1.5: Market prices for the algorithmic and human markets for each supergame (SG). I derive the prices for the algorithmic markets upon convergence as an average over 1,000 subsequent periods for 1,000 independent simulations. The error bars represent the standard deviation.

Similar to previous findings in the literature (e.g., Huck et al., 2004), collusion becomes more difficult for humans as the market size increases. Average market prices are higher for each supergame in 2H0A compared to 3H0A. Those differences are (weakly) statistically significant for the first and second supergame but insignificant for the last supergame (SG1 $p = 0.045$; SG2 $p = 0.055$, SG3 $p = 0.283$; two-sided Mann–Whitney U tests). Thus, while the disparity in prices becomes smaller after learning, prices are always higher in two-firm markets compared to three-firm markets, which is in line with Hypothesis 2. While both algorithms and humans see a drop in price due to the expanded market size, the decline is greater for algorithmic markets. It suggests that the market size within the discussed environment might be more harmful to algorithmic than to human collusion.⁴¹

Result 4. *Similar to algorithmic markets, the level of tacit collusion declines for humans as the market size increases. Price drops due to the increase in market size are higher for algorithmic compared to human markets.*

⁴¹For future research, it can be interesting to see the development of those number effects for even larger markets.

In two-firm markets, algorithms outperform humans at colluding. Average market prices in 0H2A are statistically significantly higher than in 2H0A for each supergame when using the selected algorithm as a comparison unit ($p < 0.01$ for all supergames; two-sided Mann–Whitney U tests). Also, when considering the average market price of the entire parameter grid discussed in Section 1.6.1, prices in 2H0A are smaller ($p < 0.05$ for all supergames; one-sample two-sided t-tests against the average grid price of 3.51).

Furthermore, the selected algorithms in three-firm markets are more collusive than humans in the first two supergames. However, this advantage entirely fades after the first two supergames as there are no differences between algorithms and humans in the third and last supergame (SG1 $p < 0.01$; SG2 $p < 0.01$, SG3 $p = 0.980$; two-sided Mann–Whitney U tests). Hence, after humans had the chance to learn about the game, they are as good as self-learned algorithms at colluding in three-firm markets. If I compare prices in human markets to the average algorithmic outcome in the parameter grid, no statistically significant price differences exist for the first and second supergame. Moreover, in the last supergame, the average grid price within three-firm human markets even exceeds the average price of the parameter grid (SG1 $p = 0.991$; SG2 $p = 0.340$, SG3 $p = 0.044$; one-sample two-sided t-tests against the average grid price of 1.57). Hence, trained algorithms can outperform inexperienced humans at colluding in markets with three firms. Yet, humans are as good as algorithms at colluding after they gain experience. If a firm fails to pick an optimal algorithm, humans can even surpass algorithmic performance in this environment.

Result 5. *Algorithms are more collusive than humans in two-firm markets. In three-firm markets, algorithms outperform inexperienced humans at colluding but there are no price differences if humans are experienced.*

Result 5 is in line with Hypothesis 1 for two-firm markets. It is also the case for three-firm markets, if humans are inexperienced. There is no evidence that algorithms hurt competition in three-firm markets after humans adapt and learn themselves.

1.6.4 Collusion between humans and algorithm

In this section, I consider the outcomes for mixed markets in which humans compete against algorithms. The algorithms always play according to the limit strategy of the

selected algorithm. Hence, similarly to Normann and Sternberg (2023) who consider a tit-for-tat algorithm, the humans compete against a fixed strategy in the experiments. In contrast to Normann and Sternberg (2023), the strategy is a result of the learning procedure of the algorithms instead of being chosen by a researcher. Furthermore, the WSLs algorithm maximizes the selection criterion Ψ_i . Thus, it would arguably be used by firms.

Figure 1.6 shows the average market price pooled across all supergames for each treatment with human involvement.

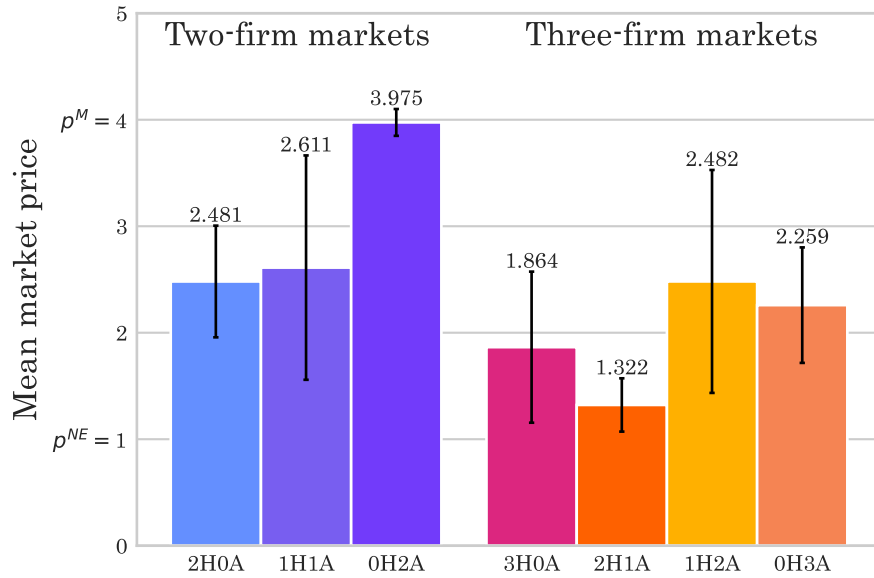


Figure 1.6: Average market prices for all treatments. For treatment with humans, I pool market prices across all super games. For algorithmic markets, I use the parameterization of the selected algorithm as a comparison unit. The error bars represent the standard deviation.

Within two-firm markets, there are no statistically significant differences in market prices between two humans (2H0A) and one human competing with one algorithm (1H1A) ($p = 0.84$, two-sided Mann–Whitney U test). Thus, contrary to Hypothesis 1, the pricing algorithm does not foster collusion. Nevertheless, on average, a single algorithm is as good at colluding with a human as another human player. Furthermore, prices in 1H1A are significantly lower than in the fully algorithmic market 0H2A ($p < 0.01$, two-sided Mann–Whitney U test). Hence, while algorithms never foster competition in a duopoly, they can make markets more collusive if all firms utilize them.

In three-firm mixed markets, I observe a non-linear relationship between the level of tacit collusion and the number of algorithms in the market. Market prices in 2H1A are

lower than in 3H0A ($p = 0.07$, two-sided Mann–Whitney U test).⁴² Adding another algorithm to the market (1H2A) increases prices again compared to 2H1A ($p < 0.01$, two-sided Mann–Whitney U test). There are no statically significant differences between 1H2A and 0H3A using algorithms with the parameterization of the selected algorithms as a comparison unit ($p = 0.76$, two-sided Mann–Whitney U test). However, average prices in 0H3A are higher than market prices in 3H0A if the outcomes are pooled across supergames ($p < 0.01$, two-sided Mann–Whitney U test).

Result 6. *Humans manage to cooperate with pricing algorithms. In duopolies, algorithms (weakly) foster tacit collusion. In triopolies, there exists a non-linear relationship between the level of tacit collusion and the number of algorithms in the market. If most firms use pricing algorithms, markets can become less competitive.*

Within my framework, firms have a clear incentive to use pricing algorithms in a duopoly. If only a single firm adopts it, prices do not change. Yet, if both firms outsource their pricing decisions to an algorithm, markets become more collusive, which in turn increases firms' profits.⁴³ This effect resembles recent findings on algorithm pricing in the German gasoline market. Assad et al. (2022) show that market-level margins do not increase if only one gas station in a local market adopts a pricing algorithm. However, if both gas stations in the duopoly adopt it, the price algorithm margins increase by 28%. For triopolies, I find vastly different outcomes depending on the exact market composition in mixed markets, but also, with three firms in the market, algorithms can hurt competition. It is especially the case if most firms decide to use pricing algorithms and humans lack experience. However, adoption incentives are less pronounced compared to a duopoly, as firms' profits can decrease if only a single firm utilizes them.

Heterogeneous strategies in mixed market In Figure 1.7, I plot the average market price by round and supergame for each experimental treatment. While 1H2A and 1H1A have a similar trend as 2H0A in the first supergame, 3H0A and 2H1A have noticeably

⁴²Note that the results differ from Normann and Sternberg (2023) who find that a single tit-for-tat algorithm fosters collusion with three firms in a simpler market environment. The strategies of the human sellers drive the results, and I analyze them in the subsequent paragraph.

⁴³Köbis et al. (2021) argue that the decision to delegate the pricing to an algorithm can be particularly relevant as it allows the firms' manager to morally distance herself from the unethical behavior of collusion. In my experiment, firms cannot decide whether to adopt a pricing algorithm as it is determined exogenously. However, it can be a path for future research to examine the adoption decision of participants.

lower market prices. In fact, after some initial rounds, average prices in 2H1A are at the stage game Nash equilibrium. Some interesting patterns emerge in the later supergames after the participants learn about the game. Prices in 2H1A are still close to the stage game Nash equilibrium. While average market prices in 1H1A and 1H2A are similar to 2H0A, there are sharp spikes every other round.

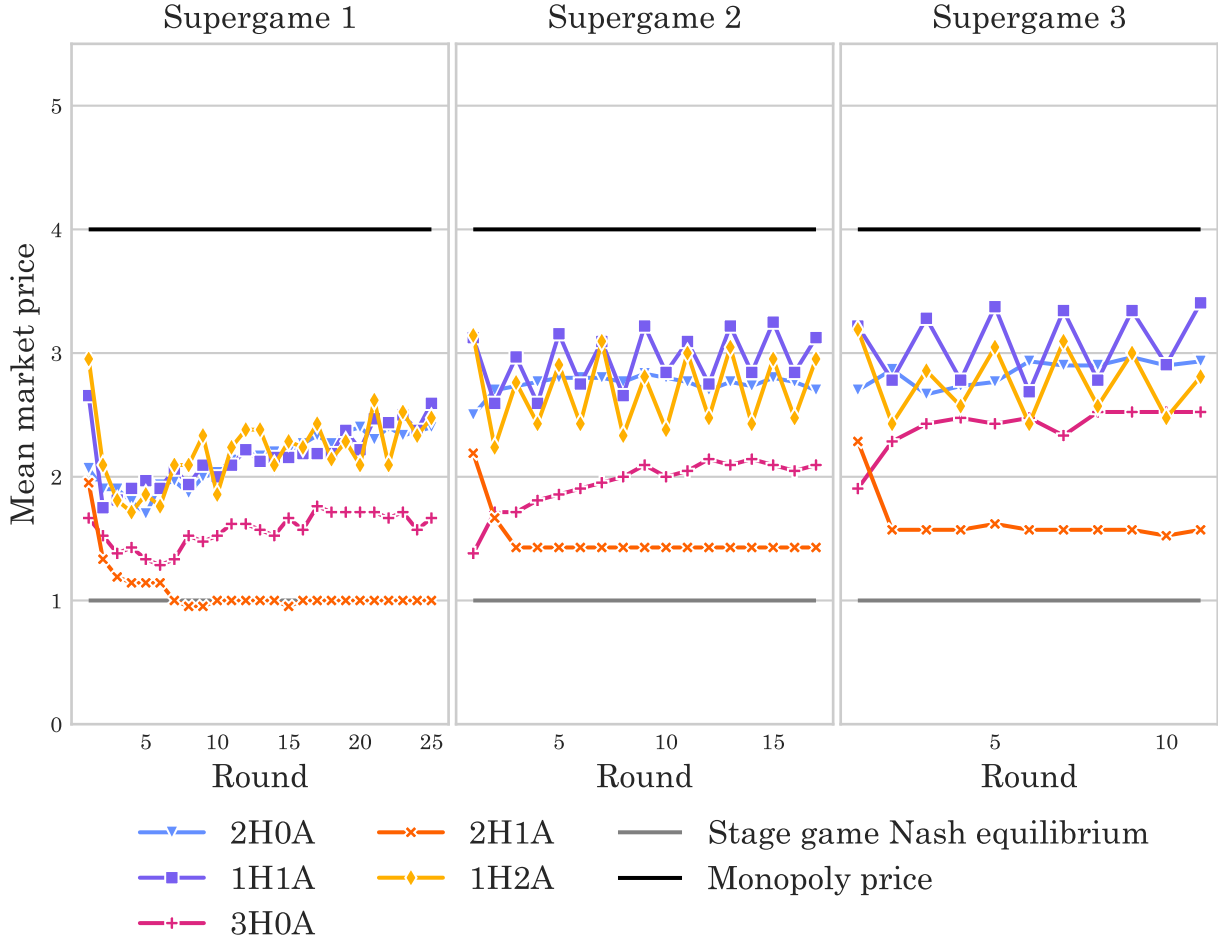


Figure 1.7: Average market prices by supergame and round for all experimental treatments.

To understand those price patterns, it is essential to remember that participants during the experiment play against the limit strategy of the selected algorithm, which is described by Equation 1.6. In other words, participants play against a variation of a win-stay lose-shift (WSLS) strategy. Normann and Sternberg (2023) demonstrate that the algorithm's strategy is a significant determinant of outcomes in human-machine interactions. The expectations that participants have about the algorithm's behavior are mostly irrelevant. Hence, it is essential to understand how participants respond to the algorithm's strategy in the presented setup.

While participants do not know the algorithm's strategy initially, they can learn about it during the first supergame. Once a participant understands how the algorithm works, there are different ways to adapt her strategies as a response. First, she can ALWAYS COOPERATE with the algorithm at the monopoly price. Second, she can try to EXPLOIT the algorithm by playing the price cycle strategy described in Equation 1.7. This strategy dominates ALWAYS COOPERATE in 1H2A and is dominated by ALWAYS COOPERATE in 1H1A.⁴⁴ Other strategies are possible. Namely, participants can ALWAYS DEFECT at the stage game Nash equilibrium. Furthermore, they can play an imperfect exploitation strategy by playing a price of $p = 2$ in the cooperative phase of the algorithm and $p = 1$ otherwise. I denote this strategy by EXPLOIT2. The strategies ALWAYS DEFECT and EXPLOIT2 are dominated by ALWAYS COOPERATE and the EXPLOIT, respectively.

To investigate which strategies the participants use in 1H1A and 1H2A against the algorithm, I estimate a mixture model using the Strategy Frequency Estimation Method (SFEM) proposed by Dal Bó and Fréchette (2011). The method is highly influential for estimating strategy choices in infinitely repeated games, especially the Prisoner's Dilemma (e.g., Fudenberg et al., 2012; Romero and Rosokha, 2018; Dal Bó and Fréchette, 2019). Starting from a predefined set of strategies, SFEM assumes that subject i chooses strategy s^k with probability ϕ^k and follows this strategy for all rounds of the game. In each period, participant i selects her price according to strategy s^k with probability $\sigma \in (1/2, 1)$ but makes an error with probability $1 - \sigma$. The individual likelihood that participant i plays according to strategy k is given by $P_i(s^k) = \prod_t \sigma^{I_{t,i}} (1 - \sigma)^{1-I_{t,i}}$. The identifier variable $I_{t,i}$ is equal to 1 if the price of participant i in period t corresponds to the price she would have played if she followed strategy s^k . Otherwise, $I_{t,i}$ equals zero. The log-likelihood function is given by $\mathcal{L} = \sum_i \ln(\sum_k \phi^k P_i(s^k))$. The estimate of ϕ^k represents the share of participants in the population that uses strategy k . The value of σ can be interpreted as a goodness of fit parameter. The model is noisy if σ is close to its lower bound of 0.5. The model describes the data well for values of σ that are close to 1.

For the estimation procedure, I focus on the strategies that are reasonable when competing against the algorithm (ALWAYS COOPERATE, ALWAYS DEFECT, EXPLOIT, and EXPLOIT2). Moreover, I restrict the analysis to the last supergame. Table 1.3 shows the results of the estimation procedure.

⁴⁴For details see Appendix 1.B.

Table 1.3: Estimated proportion for each strategy

Strategy	Treatment	
	1H1A	1H2A
ALWAYS COOPERATE	0.61 (0.09)	0.48 (0.11)
ALWAYS DEFECT	0.10 (0.05)	0.22 (0.10)
EXPLOIT	0.29 (0.08)	0.29 (0.10)
EXPLOIT2	0.00 (0.00)	0.02 (0.05)
σ	0.92	0.84

* The mixture model is estimated by maximum likelihood estimation. I restrict the data to the last supergame. The bootstrapped standard errors are in parentheses.

The most frequent strategy that participants play against the algorithm is ALWAYS COOPERATE in both treatments. The estimated proportion is, however, smaller in 1H2A compared to 1H1A. Also, EXPLOIT is prevalent in the population, but the estimates do not differ between 1H1A and 1H2A. Notably, the share of ALWAYS DEFECT is higher in 1H2A compared to 1H1A. The imperfect exploitative strategy EXPLOIT2 is never played in 1H1A, and it only accounts for a share of 0.02 of the data in 1H2A.

In line with the shift in incentives when increasing the market size, fewer participants play a cooperative strategy against the algorithm in 1H2A. Nevertheless, participants often fail to learn the best response as EXPLOIT, and ALWAYS COOPERATE dominate ALWAYS DEFECT in 1H2A. A possible reason is that learning about the environment is more difficult in 1H2A due to higher strategic complexity. While both algorithms in 1H2A use a WSLS strategy, participants still have to consider additional information compared to 1H1A. That can impede learning for a subset of participants. Individual prices reveal that some participants circle between prices of 1, 2, and 3 without a clear pattern. It appears that those participants did not learn to follow a fixed strategy (see Appendix 1.B.2 for the price patterns on an individual level). This argument is also supported by the smaller value of σ and higher standard errors in 1H2A, as it indicates a more noisy behavior of the participants. While average market prices in 1H1A (1H2A) and 2H0A (3H0A) are similar, it is usually not the case for individual markets. Depending on the

particular strategies that humans learn, mixed markets can be more or less collusive than their entirely human counterparts.⁴⁵

In 2H1A, it is also crucial to consider the possible strategies humans can use against the algorithm. While ALWAYS COOPERATE and EXPLOIT are both still viable options to play against the algorithm, they now require joint coordination by two humans. Indeed, low prices in 2H1A can be explained by a frequent failure to coordinate simultaneously against the algorithm. While some markets manage to collude at the monopoly price with the algorithm, most participants fail to coordinate on any other strategy than ALWAYS DEFECT against the algorithm.⁴⁶

Result 7. *Most humans always cooperate with the algorithm or try to exploit it in 1H1A and 1H2A. In 2H1A, most humans always defect as they fail to coordinate on a joint strategy against the algorithm. Market outcomes differ substantially conditionally on the exact strategies that humans learn.*

1.7 Concluding remarks

In this chapter, I study the collusive potential of self-learning pricing algorithms and show that pricing algorithms can weaken competition. Similar to previous results by Calvano et al. (2020a) and Klein (2021), I observe that algorithms learn to set prices above the competitive benchmark and develop reward-punishment strategies in simulations. As the market environment is stylized and, therefore, highly tractable, I can analyze the strategies of the algorithms. The most successful algorithms learn a win-stay lose-shift strategy. To derive a counterfactual for algorithmic collusion and observe the interaction of humans and pricing algorithms, I conduct laboratory experiments with the same market environment as in the simulations. Across different treatments, I vary the market size and the number of firms that use a self-learned pricing algorithm. This approach allows me to pin down the anti-competitive effects algorithms can have across a wide range of market compositions.

In duopolies, algorithmic markets are always more collusive than human markets. Markets with one human and one algorithm have similar average market prices compared

⁴⁵Also Wieting and Sapi (2021) find heterogeneous market outcomes depending on the exact number of algorithms in the market using data from the Dutch online retailer *bol.com*.

⁴⁶Figure 1.12 in Appendix 1.B.2 highlights those price patterns.

to entirely human markets. In three-firm markets, market prices decrease if a single firm uses a pricing algorithm. It is driven by the specific strategy the algorithms learn and the failure of humans to coordinate with the algorithm. As more firms utilize pricing algorithms, prices increase again in three-firm markets. If all firms in the market use an algorithm, market prices can be higher than in human markets. However, the effect fades after humans have the chance to learn about the market environment. Most participants cooperate with the algorithm, but the strategies are heterogeneous, and some participants try to exploit the algorithm.

My results highlight the potential anti-competitive effects of self-learning algorithms. While market outcomes vary depending on the exact parameterization and market composition, algorithms rarely foster but often weaken competition if they populate the market. The considered pricing algorithms are simple, and the experimental environment is stylized. Yet, it appears probable that more complex algorithms can achieve similar results and scale to more complex real-world markets.⁴⁷ Within the presented framework, the fear from competition authorities that algorithms can harm the competitive landscape is justified.

Current research in computer science focuses on explainable artificial intelligence (see Barredo et al., 2020). The development objective for those algorithms is that humans can understand their results and the decision process. Also, for pricing algorithms, explainable artificial intelligence is desirable. Understanding why algorithms learn to be collusive and how they must be designed to prevent collusive market outcomes is critical. Asker et al. (2022) underlines the significance of the algorithmic design on its collusive behavior, but more research is needed to determine a suitable procedure to regulate pricing algorithms. Furthermore, recent work by Calvano et al. (2020b) proposes to audit algorithms before firms can use them as a pricing tool. Competition authorities could examine the algorithm in a simulated market environment to evaluate its potential for tacit collusion and ban specific algorithms if necessary. Auditing is not feasible for tacit collusion among humans. My findings indicate that it is not only possible for algorithms but also necessary to prevent harm to competition.

⁴⁷Hettich (2021) shows that deep reinforcement learning algorithms can be collusive in a different market environment with up to ten firms.

Appendix

Appendix 1.A Implementation details

The instructions were translated from German. Section 1.A.1 provides a translation for the 2H1A treatment.⁴⁸

1.A.1 Instructions

Particularly important: If you have any questions, contact the administrator using the chat function in the Webex conference.

Once you took a decision on the respective page and read all the information, **please click on the "Next" button so that the experiment can continue.** If you **do not make an input for an extended time** or temporarily leave this website, we will remove you from the experiment.

In this case, you will **not receive any payment and will be banned from future online experiments.**

Instructions

In this experiment, you will repeatedly make price decisions. These allow you to earn real money. How much you earn depends on your decisions and those of the other participants. **Regardless, you will receive 4.00 euros for participating.** In the experiment, we use a fictional monetary unit called ECU. After the experiment, the ECUs are converted into euros and you will be paid accordingly. **Here, 130 ECUs correspond to one euro.**

In this experiment, you represent a firm in a virtual product market. In each market, two

⁴⁸The complete oTree application to conduct the experiment is available here: <https://github.com/ToFeWe/AlgoCollusionApp>

other firms sell the same product as you do. These firms are represented by two other experiment participants. All firms offer 60 units of the same product. There occur no costs of production to the firms. The game has multiple rounds, with the exact number being decided by a random mechanism. You play the game for three sessions. In each round of a session, you meet the same firms (i.e., experiment participants). However, the firms in your market change after each session.

The market has 60 identical customers. Each customer in each round of a session intends to buy **one unit** of the product as cheaply as possible. Each customer is willing to spend a maximum of **4 ECUs** for this unit of the product. All firms decide in each round again **at the same time** for how many ECUs they want to sell their product. You can sell your product for 0 ECU, 1 ECU, ... or 5 ECUs (only whole ECUs). Your profit is the price multiplied by the number of units sold. Formally expressed:

$$\text{Profit} = \text{Price} \times \text{Units sold}$$

The firm with **the lowest price in the given round** sells its products as long as the price is not greater than 4 ECUs. Firms with a higher price do not sell their product. **The lowest price is the market price in the given round.** Firms with a **price higher than the market price do not sell their products in that round.** If two or all three firms want to sell their products for the same market price, the demand is divided equally between the two or three firms.

Examples:

Example 1: Firm A sets a price of 3 ECUs, firm B sets a price of 3 ECUs, firm C sets a price of 4 ECUs. Thus, firms A and B together set the lowest price. Firm A and B both sell the same number of products. Both firms have 30 customers each and thus get the same profit of 90 ECUs. Firm C sells nothing and has a profit of 0 ECUs.

	Firm A	Firm B	Firm C
Price	3	3	4
Profit	90	90	0

Example 2: Firm A sets a price of 2 ECUs, firm B sets a price of 2 ECUs, firm C sets a price of 2 ECUs. Thus, firms A, B, and C together set the lowest price. They all sell the same number of products (20 each) and thus get the same profit of 40 ECUs.

	Firm A	Firm B	Firm C
Price	2	2	2
Profit	40	40	40

Example 3: Firm A sets a price of 1 ECUs, firm B sets a price of 2 ECUs, firm C sets a price of 3 ECUs. Thus, firm A set the lowest price. Firm A is the only one that sells the product for 1 ECU to all 60 customers and thus gets a profit of 60 ECUs. Firms B and C both sell nothing and have a profit of 0 ECUs.

	Firm A	Firm B	Firm C
Price	1	2	3
Profit	60	0	0

Example 4: Firm A sets a price of 5 ECUs, firm B sets a price of 5 ECUs, firm C sets a price of 5 ECUs. Thus, firms A, B, and C together set the lowest price. However, customers are only willing to buy the product for 4 ECUs. Therefore, no firm sells its products and all firms get a profit of 0 ECUs.

	Firm A	Firm B	Firm C
Price	5	5	5
Profit	0	0	0

Market decisions by algorithms:

In your markets, two participants decide at which price they want to sell their firm's product and are paid the profit their firms earn at the end of the experiment.

Firm C will be equipped with an algorithm in all rounds, which will make the necessary price decisions for the participant. In this case, the participant does not take any decisions, but **still receives the profit** that his or her firm earns.

The procedure of the experiment:

After each round, all firms are informed about the prices chosen by each firm and about

their profit. In the next round, each firm again has the chance to re-select its price. You interact with the same participants in each round within one session.

After each round, a random mechanism decides whether another round is played or the session ends. The probability that another round will be played is 95%. Thus, the session ends after each round with a probability of 5%. The session continues until the end is determined randomly.

Figuratively, the computer throws a virtual dice with 20 sides before each possible further round. The result decides whether another round is played or not. If the number is 20, the session is over; for all other numbers, another round is played.

Note:

You will play the game described for a total of three sessions. After each session, you will be paired with other participants to form a new market. This means that you interact with other participants in each of the three sessions. **After all, sessions are completed, it will be randomly decided which of the three sessions will be paid for.** You will receive this profit after the experiment. You will also receive additional 4.00 euros for participating in this experiment.

1.A.2 Screenshots of the experiment

Wählen Sie Ihren Preis

Bitte wählen sie Ihren Preis zwischen 0 und 5 Talern:

Weiter

Zeige Instruktionen

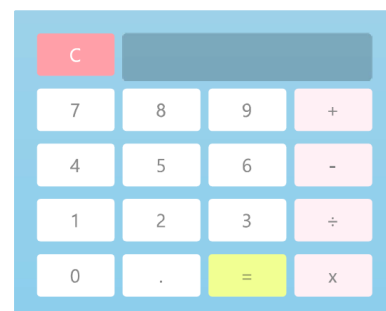


Figure 1.8: Decision screen

Ergebnisse

Ergebnisse der aktuellen Runde

	Sie (Firma B)	Firma A	Firma C
Preise	4 Taler	4 Taler	4 Taler

Ihr Preis in dieser Runde betrug 4 Taler. Der Marktpreis betrug 4 Taler. Sie erhalten in dieser Runde einen Gewinn von 80 Taler.

Weiter

Zeige Instruktionen

Figure 1.9: Profit information screen

1.A.3 Survey & control questions

Control Questions

Question 1: How many consumers are in the market who want to buy the product?

- 35
- 30
- 40
- 60

Question 2: What is the probability of playing another period after completing one?

- 95%
- 20%
- 50%

Question 3: What is the maximum price consumers are willing to pay for the product?

Question 4: You are firm A and choose a price of 2, firm B chooses a price of 3, firm C chooses a price of 5. What is your profit in ECU in this round?

Question 5: You are firm A and choose a price of 3, firm B chooses a price of 3, firm C chooses a price of 3. What is your profit in ECU in this round?

Appendix 1.B Strategy analysis

1.B.1 Incentive compatibility of the algorithm's strategy

I want to show that the algorithms' strategy can be exploited by the strategy described by Equation 1.7. I focus on the case in which a player k uses this strategy, and all other players use the limit strategy of the algorithm (Equation 1.6). When using the EXPLOIT strategy, the firms enter a price cycle. After a deviation to $p_k = 3$, all firms play the stage game Nash equilibrium. In the following period, firms $-k$ play the monopoly price while firm k plays again $p_k = 3$ to restart the cycle. Hence, in every other period firm k receives the deviation profit of $\pi_D = 3m = 180$. In the other periods, all firms share the market and firm k receives $\pi_{NE}/N = m/N = 60/N$. Cooperation at the monopoly price yields $\pi_M/N = 4m/N = 240/N$. Thus, the exploitative strategy dominates cooperation in an infinitely repeated game with discount factor δ if

$$\begin{aligned}
 & V^{Cooperate}(N) \geq V^{Exploit}(N) \\
 \Leftrightarrow & \sum_{t=0}^{\infty} \delta^t \frac{\pi_M}{N} \geq \pi_D + \delta \frac{\pi_{NE}}{N} + \delta^2 \pi_D + \delta^3 \frac{\pi_{NE}}{N} + \dots \\
 \Leftrightarrow & \frac{1}{1-\delta} \frac{\pi_M}{N} \geq \frac{1}{1-\delta^2} \pi_D + \frac{\delta}{1-\delta^2} \frac{\pi_{NE}}{N} \\
 \Leftrightarrow & \frac{1}{1-\delta} \frac{240}{N} \geq \frac{1}{1-\delta^2} 180 + \frac{\delta}{1-\delta^2} \frac{60}{N}
 \end{aligned} \tag{1.8}$$

Within the experiment and the simulations, the discount rate is $\delta = 0.95$. Hence, for a two-player game we have $V^{Cooperate}(N = 2) = 2400$ and $V^{Exploit}(N = 2) \approx 2138.46$. It implies that cooperation with the algorithm at the monopoly price dominates the exploitative strategy.

For a three-player game it is $V^{Cooperate}(N = 3) = 1600$ and $V^{Exploit}(N = 3) \approx 2041.03$. Thus, exploiting the algorithm dominates cooperation.

1.B.2 Individual prices in the last supergame

Figure 1.10 and 1.11 show the individual prices for the treatments 1H1A and 1H2A respectively. Furthermore, I plot the price of the algorithm for each round. Note that this prices is the same for both algorithms in 1H2A.

Both figures reveal the price patterns associated with the strategies described in Section 1.6.4. In very few rounds in 1H2A, the strategy of the algorithm diverges from the win-stay lose-shift strategy described by Equation 1.6. While this may delay efficient learning in the game, participants never play a strategy that exploits those negligible deviations.

Figure 1.12 shows all prices for each market in the last supergame in 2H1A. In few markets, both humans manage to collude with the algorithm. The other markets usually have a market price that is equal to the stage game Nash equilibrium.

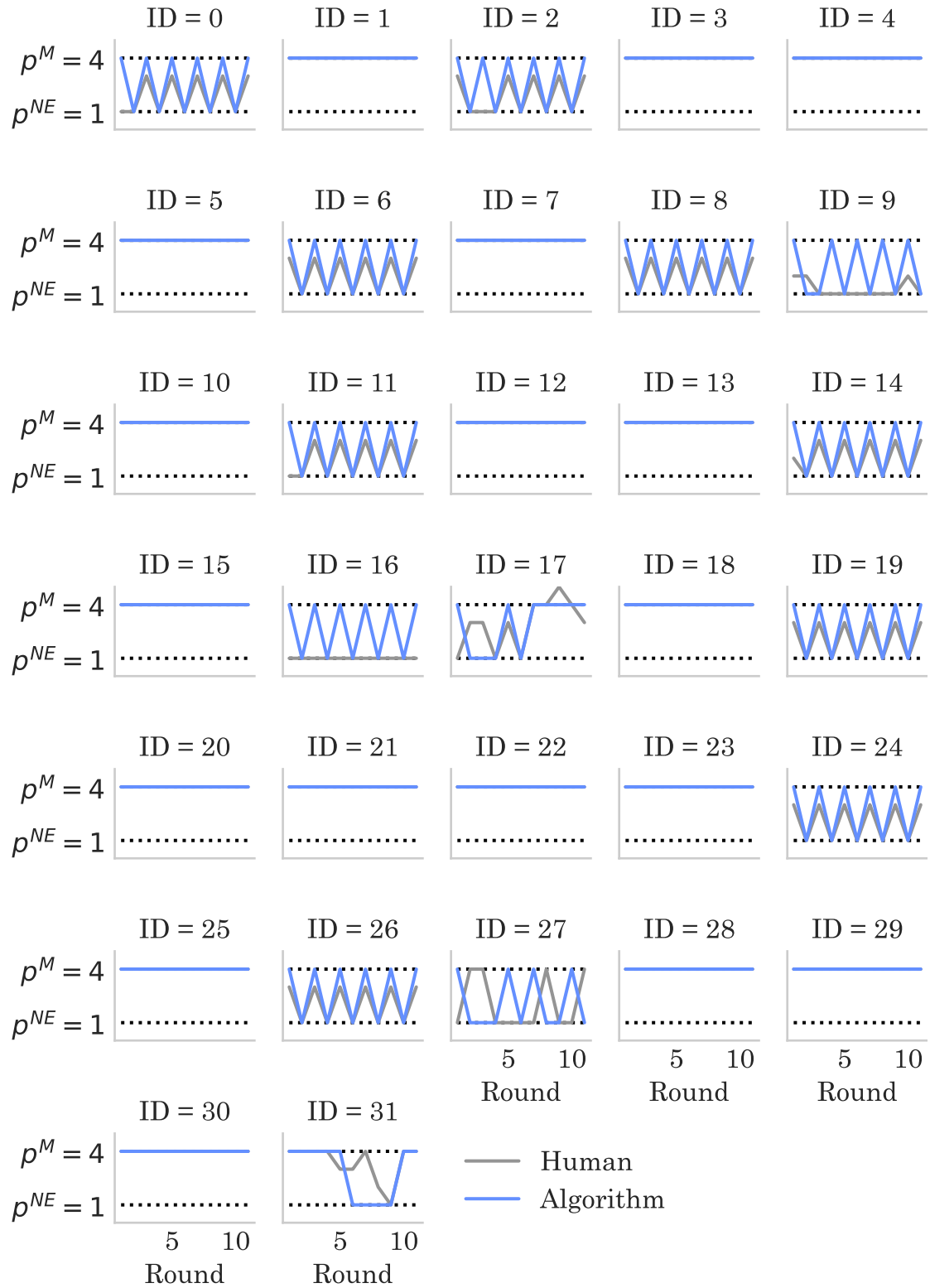


Figure 1.10: Prices for each human participant in the last supergame in 1H1A.

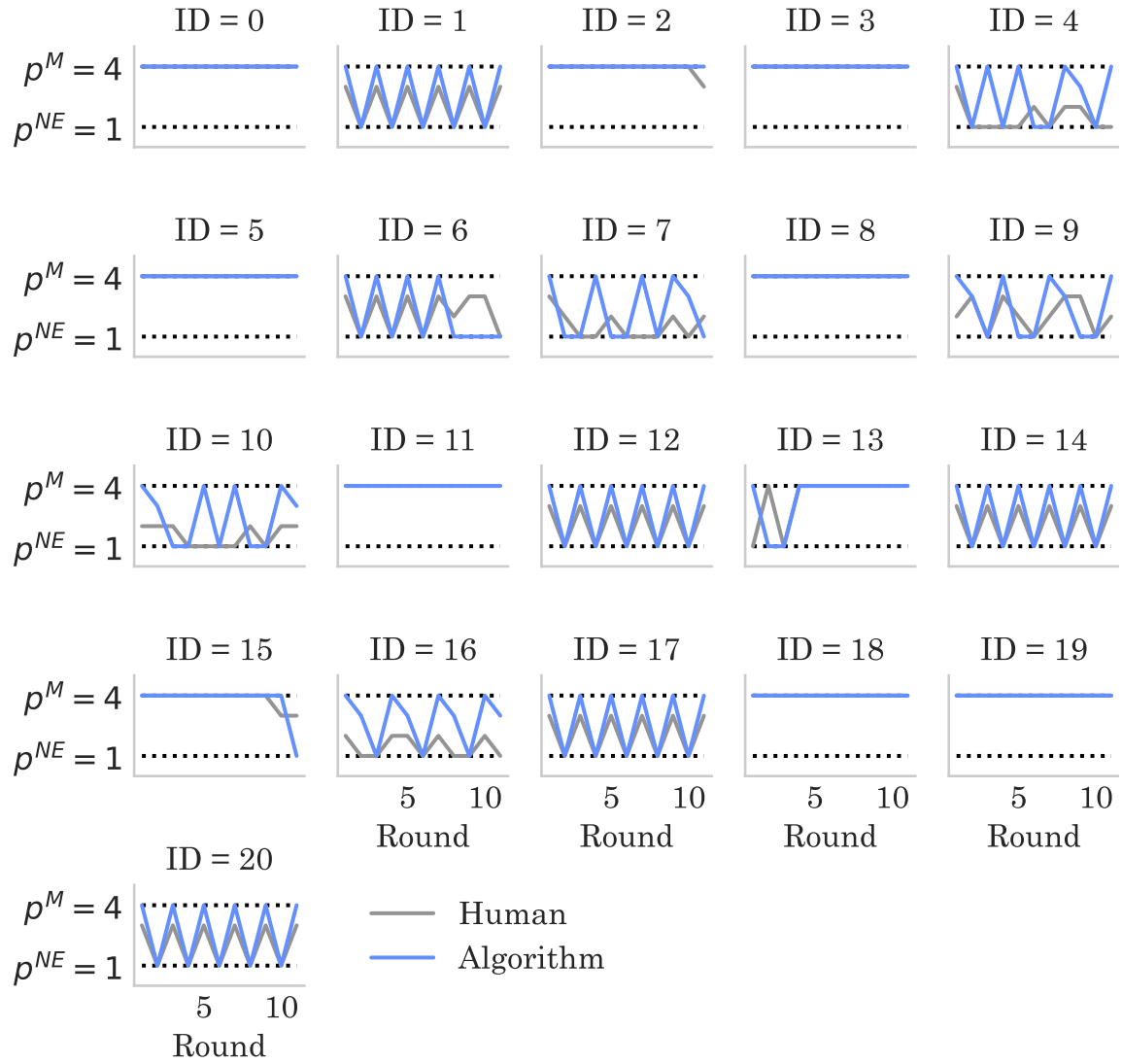


Figure 1.11: Prices for each human participant in the last supergame in 1H2A. Note that both algorithm always play the same price.

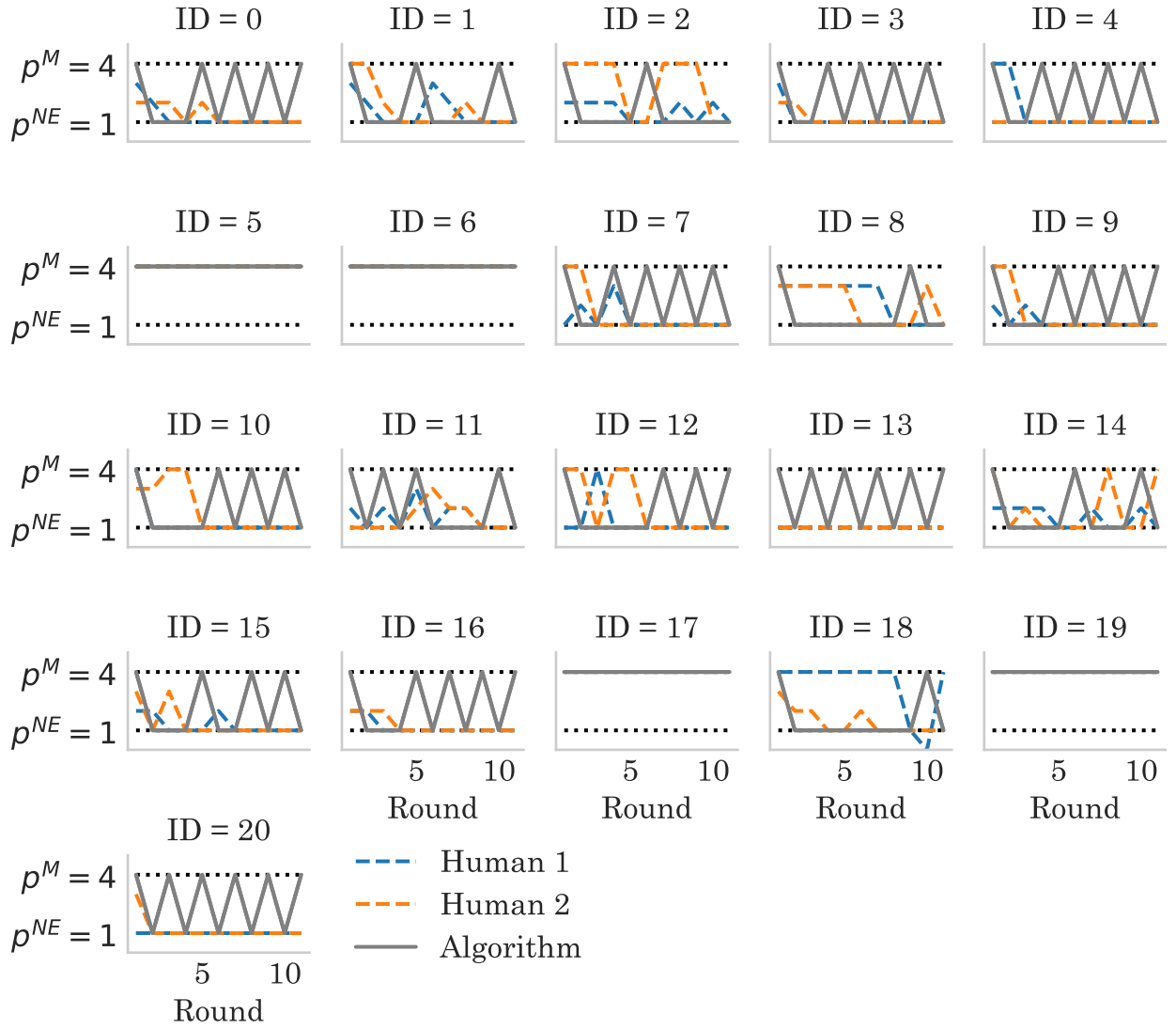


Figure 1.12: Prices for each market in the last supergame in 2H1A.

Chapter 2

Algorithmic Price Recommendations and Collusion: Experimental Evidence

Co-authored with Matthias Hunold

2.1 Introduction

Price recommendations are prevalent on many digital marketplaces. Platforms like Airbnb, Expedia, and eBay provide the sellers on their marketplaces with a recommendation on how to set the price for the product they sell. Those recommendations are typically created by algorithms based on historical market data.¹ Importantly, recommendations are typically non-binding in the sense that sellers can nevertheless freely choose their prices.

Various explanations for the use of price recommendations exist in practice. For instance, price recommendations might reduce information asymmetries between the platform and sellers (see Pavlov and Berman, 2019). This makes them potentially attractive from a business perspective as, for instance, platforms may have better demand information than individual sellers. However, competition authorities are concerned that price recommendation algorithms by a common intermediate can also dampen competition as it might make coordination between sellers easier (Competition & Markets Authority, 2021; Bundeskartellamt and Autorité de la concurrence, 2019). For example, according to reporters of ProRepubblica, price recommendation software allegedly led to coordination effects in the U.S. rental market, especially in regions where few property managers control a large share of the apartments.²

This article examines whether platforms can use algorithms that provide sellers with non-binding price recommendations to make markets more collusive. As digital sales platforms often receive a share of the seller's revenues through commission rates, they can benefit from higher prices if there is seller competition.³ Even if the platform's income does not directly depend on the sellers' prices, using algorithmic pricing software may increase sellers' profits and thus their willingness to pay for joining the platform and using the software. This, in turn, can increase the platform's profits. These arguments

¹For instance, Airbnb uses price recommendation algorithms that utilize historical and geographical data and combine machine learning methods with human intuition. Furthermore, the algorithmic price recommendation changes daily for the upcoming dates for which the accommodation is available. See Hill (2015) for details.

²See Vogell, Coryne & Little, "Technology Rent Going Up? One Company's Algorithm Could Be Why.", <https://www.propublica.org/article/yeildstar-rent-increase-realpage-rent> (accessed on November 25, 2022).

³A natural question is why a platform cannot simply increase the commission rate to induce higher sales prices. We theoretically demonstrate that steering the sales prices with commission rates can be impossible or insufficient for sales platforms for various reasons, so collusive price recommendations may be desirable. See Appendix 2.A for details.

motivate the question of whether recommendation algorithms can indeed raise prices and, if so, merit closer examination by competition authorities and regulators.

We derive two rule-based algorithms from economic theory and behavioral insights to investigate the effects of a platform’s recommendation algorithm on the prices of competing sellers.⁴ These algorithms recommend collusive strategies with different punishment mechanics. The recommendations are non-binding and do not change the game’s strategy space and payoff functions as they do not provide fundamentally new information. In practice, however, collusive outcomes can be challenging to achieve without communication among competitors or another form of coordination (see, for instance, Fonseca and Normann, 2012). We argue that the recommendation algorithms have the potential to facilitate coordination and, thus, collusion.

The theory-based algorithm recommends actions that are consistent with a collusive trigger-strategy and Nash reversion. If a seller undercuts the collusive price level, it recommends competitive prices for several periods until it returns to recommending collusive prices. All players following the recommendations constitutes a subgame perfect Nash equilibrium. Additionally, we consider an algorithm that is motivated by behavioral findings whereby firms often do not use harsh punishment strategies (see, for instance, Wright, 2013). This algorithm recommends brief punishment phases with prices at the level of the deviating price and returns to recommending high prices when the sellers comply. The recommendations in our experiment provide no fundamentally valuable information, such as the state of demand, to sellers. Instead, their only purpose is to coordinate sellers. We abstract from other factors that could make following the recommendations desirable in order to isolate the pure coordination effect of recommendations.

We derive several testable hypotheses based on our theoretical model and test those in laboratory experiments. Subjects resemble competing sellers who repeatedly set prices and receive price recommendations from an algorithm in each round. Across treatments, we vary whether participants receive a recommendation or not as well as the type of recommendation algorithm. We inform the subjects in the experiment that the recommendation algorithm’s objective is to maximize industry profits symmetrically without

⁴Recent studies by Musolff (2022) and Wieting and Sapi (2021) highlight that price algorithms on sales platforms often follow simple rule-based logic. Similarly, even complex reinforcement learning algorithms often converge to strategies that simple rules can describe (see, for instance, Werner, 2021).

favoring a particular seller. Recommending high sales prices is consistent with the incentives of a platform that receives part of the sellers' revenues through commission rates.⁵

The algorithmic price recommendations positively influence individual pricing decisions in the sense that higher recommended sales prices induce sellers to set higher individual prices. The estimated "pass-on rate" from recommended prices to sales prices is between 0.22 and 0.57, depending on the recommendation algorithm. The pass-on rate is higher for the theory-based algorithm that recommends collusive trigger strategies with temporary Nash reversal.

The effects on the realized market prices and profits differ sharply between the distinct recommendation algorithms. We find insightful price patterns for the theory-based algorithm even though the average market prices do not differ from the control treatment without any price recommendation. The substantial heterogeneities can explain the absence of an average treatment effect in the market outcomes. The collusive effect of the algorithm depends on the seller's characteristics. In markets where sellers have low levels of negative reciprocity, the recommendation algorithm decreases market prices. Thus, if participants are usually not willing to punish unfair behavior, the recommendation leads to lower market prices. Furthermore, if sellers are relatively impatient, the recommendations make markets more collusive. In other words, the recommendations increase market prices in groups of sellers which are usually too impatient for collusive strategies to be sustainable.

The behaviourally motivated algorithm also recommends the monopoly price but differs in the reaction to the deviation of a seller. For this algorithm, we find lower market prices and profits than without any recommendation. Participants repeatedly deviate downwards from the recommendation, which triggers a downward spiral that leads to lower prices. We find no evidence that this algorithm fosters collusion for any subgroup. Hence, the algorithm makes markets more competitive. It is particularly interesting against the backdrop of observations where humans prefer soft punishments for deviations from collusion in experiments (see, for instance, Wright, 2013).

Related literature Our article relates to the literature on the collusive effects of algorithmic pricing. There exists evidence that algorithms can foster collusion and lead to anti-competitive prices (Klein, 2021; Calvano et al., 2020a; Hansen et al., 2021; Brown

⁵See Appendix 2.A for details.

and MacKay, 2022; Harrington, 2022). Johnson et al. (2022) focus on tacit collusion among self-learning algorithms on sales platforms and discuss how the platform’s design choices influence it. Normann and Sternberg (2023) and Werner (2021) show experimentally that algorithms may raise market prices even above the price level usually observed in human markets. We differ from this approach as we consider algorithms that only give recommendations but do not compete with the other firms in the market.

We also relate to the literature on recommended retail prices. These are pricing recommendations that a manufacturer provides to its retailers. In theory, they can act as a coordination device (Faber and Janssen, 2019; Buehler and Gärtner, 2013) and can make markets more collusive (Foros and Steen, 2013). Furthermore, they can also influence demand by setting a reference point for the consumers (Bruttel, 2018). The price recommendations in our setup resemble those employed by websites like Airbnb or Expedia. On the other hand, manufacturers use recommended retail prices in a supplier-retailer relationship. Also, recommended retail prices in conventional markets mostly stay the same and are traditionally distributed in a printed format, whereas the digital price recommendations in online platforms may change rapidly. Lastly, the recommendations on platforms are unobservable to consumers. Hence, they cannot influence demand directly but only through the pricing decision of the sellers.

Various papers study experimentally the effect of price announcements on collusion (e.g., Holt and Davis, 1990; Harstad et al., 1998; Harrington et al., 2016). Here, participants can announce prices and observe the announcement of the competitors before making the actual pricing decision. While price announcements can temporarily foster collusion, the effect usually fades, and prices decline to the level without any announcements. This reduced form of communication can be considered a recommendation by a firm in the market to its direct competitors. Our approach is distinct, as recommendations come from an algorithm that does not compete with the firms in the market.

Furthermore, recommendations and requests influence the decision of participants in various experimental games. They can increase contributions to public goods (Silverman et al., 2014; Croson and Marks, 2001), reduce tax evasion (Cadsby et al., 2006), and facilitate coordination in games with correlated equilibria (Duffy and Feltovich, 2010). Schotter and Sopher (2003) show that intergenerational advice provided by previous populations of experimental subjects can help to coordinate behavior. The result is robust to different

games and experimental setups (see Schotter, 2003, for a literature review). Our approach is different as an algorithm instead of previous subjects provides the recommendations.

Sonntag and Zizzo (2015) consider static quantity requests in a Cournot market game. They vary the degree of authority with which the requests are communicated to the participants across treatments. They find that this type of authoritarian recommendation can lower quantities and, thus, make markets more collusive. We consider a setup similar to theirs. However, we concentrate on neutral recommendations by an algorithm without explicitly framing the recommendation as a request. Furthermore, we go beyond static quantity recommendations and focus on dynamic recommendation algorithms that depend on the history of the game.⁶

There are few articles that directly consider price recommendations in platform markets. Pavlov and Berman (2019) consider their effects in a cheap-talk model where the platforms possess superior information about demand. They find that recommendations can be desirable compared to centralized pricing, especially if the variance of the aggregate demand is large. Lefez (2021) focuses on how platforms use price recommendations to disclose information to sellers. The potential collusive effect of price recommendations by offering a coordination device is not explicitly discussed in either paper.

The remainder of the article is structured as follows. Section 2.2 discusses the theoretical framework and the rule-based algorithms we consider. Furthermore, we derive our hypothesis. Then, we introduce the experimental design in Section 2.3 and present the results in Section 2.4. We discuss the implications of our results in Section 2.5. In Appendix 2.A, we demonstrate why a monopoly platform can benefit more from making collusive price recommendations for sellers than from only adjusting its commission rates. Appendix B contains the instructions and survey questions. We document various robustness checks and further algorithm variations in Appendix C.

2.2 Theoretical framework and predictions

We first set up a stylized pricing game with n sellers and solve for equilibria of the one-shot game and the infinitely repeated game. The game describes the experimental setup that we introduce in Section 2.3. We then argue that a recommendation algorithm can

⁶We also conducted a static recommendation treatment which we document in Appendix 2.C.3.

induce different Nash equilibria by acting as a coordination device. Finally, we motivate and explain an alternative algorithm that recommends a softer punishment scheme.

2.2.1 Setting and Nash equilibria

We consider an infinitely repeated Bertrand game with $n \geq 2$ symmetric sellers denoted by A, B, and so on. Each seller aims to maximize its profit and discounts future profit flows with a discount factor of δ .

In each period, each seller chooses its price from the integers in the set $P = \{p^N, p^N + 1, \dots, p^M\} \subset \mathbb{Z}^+$. There are k consumers who are willing to buy one unit of the good each and are willing to pay p^M per unit. The seller with the lowest price in a given period supplies the entire market. If multiple sellers have the lowest price, they share the market equally.

The sellers have no costs and no capacity constraints. Note that abstracting from costs does not change the insights from this analysis. We would get qualitatively the same results if we explicitly modeled costs which, in reality, may include commission payments. What matters for the analysis is that there is a range of prices between the relatively low competitive price level and a collusive price at which all firms make strictly higher profits.⁷

Nash equilibrium of the stage game Suppose that, except for seller A, all sellers set prices larger than p^N and at least at a level of p . Notice that seller A makes zero profits for any price higher than p , whereas setting a price of p yields a profit of $p \cdot k/n$. On the other hand, a deviation to $p - 1$ yields a profit of $(p - 1) \cdot k$. Undercutting the lowest price of a competitor, p , by one unit is the best response if

$$\begin{aligned} (p - 1) \cdot k &> p \cdot k/n \\ \implies (p - 1) &> p \cdot 1/n \\ \implies p \cdot (1 - 1/n) &> 1 \\ \implies p &> n/(n - 1). \end{aligned}$$

⁷See also Appendix 2.A for our analysis of commission payments and, in particular, equation (2.2) which shows how commission payments affect competitive prices.

We define p^N as the integer weakly below $n/(n-1)$. At this price, no firm has an incentive to undercut, such that each firm setting a price of p^N and making a profit of $p^N \cdot k/n$ is a Nash equilibrium. For $n = 3$, there is a strict incentive to undercut any price larger than 1.5, such that $p^N = 1$. As $n/(n-1)$ is decreasing in n , it follows that $p^N = 1$ for any market with $n > 3$. For $n = 2$, both a symmetric price of 1 and a symmetric price of 2 constitute a Nash equilibrium.⁸

Collusive equilibrium of the repeated game We now construct a collusive subgame perfect Nash equilibrium of the infinitely repeated game with trigger strategies. In line with the Folk Theorem, multiple collusive equilibria potentially exist. Variations are possible in the collusive price level and the punishment scheme. For instance, any price above the competitive price can potentially be supported as a collusive outcome. We focus on the highest and most profitable collusive price of p^M , which appears natural here. Among the equilibria with Nash-reversion, we focus on the equilibrium with the shortest possible punishment length. As we explain below, behavioral evidence indicates that punishments are often relatively soft.

Suppose the collusive strategy is as follows:

- If the regime is collusive in the current period, set a price of p^M .
- If the regime is punitive in the current period, set a price of p^N .
- In period one, start in the collusive regime.
- If the regime was collusive in the previous period and everyone set a price of p^M , continue with the collusive regime in the current period.
- If, in the previous period, the regime was collusive, but someone set a price below p^M , switch to the punishment regime for T periods and switch back to the collusive regime afterward.

This yields the stability condition

$$\pi^M \cdot (1 + \delta + \delta^2 + \dots) \geq \pi^D + \sum_{t=1}^T \delta^t \pi^N + \pi^M \cdot (\delta^{T+1} + \delta^{T+2} \dots),$$

⁸For $n = 2$, there is a strict incentive to undercut any (integer) price larger than 2, such that $p^N = 2$. A symmetric price of 1 is also a Nash equilibrium, but there is no strict incentive to undercut a symmetric price of 2 either as $1 \cdot k = 2 \cdot k/2$.

where π^M is the collusive period profit, π^D the deviation profit and π^N the static Nash profit as the punishment profit.

Rearranging yields

$$\begin{aligned}\pi^M \cdot (1 + \delta + \delta^2 + \dots + \delta^T) &\geq \pi^D + \pi^N \cdot \sum_{t=1}^T \delta^t \\ \Leftrightarrow \sum_{t=1}^T \delta^t &\geq \frac{\pi^D - \pi^M}{\pi^M - \pi^N}.\end{aligned}$$

Parameter values in the experiment In the experiment, we have $p^M = 10$, $k=30$, $n=3$, and consequently $p^N = 1$. We use a value of 0.95 for the discount factor δ as this equals the continuation probability in our experiment, which we introduce in Section 2.3. To determine the shortest punishment length T that makes collusion stable for these parameters, we plug in the values for the profits:

$$\pi^M = 10 \cdot 30/3 = 100;$$

$$\pi^D = 9 \cdot 30 = 270;$$

$$\pi^N = 1 \cdot 30/3 = 10.$$

This yields

$$\sum_{t=1}^T \delta^t \geq \frac{270 - 100}{100 - 10} \approx 1.89.$$

As $\delta = .95$; $\delta + \delta^2 = 1.85$; $\delta + \delta^2 + \delta^3 = 2.59$, three punishment periods are necessary and sufficient for the stability condition to hold, which constitutes a subgame perfect Nash equilibrium in trigger strategies of the infinitely repeated stage game.

2.2.2 Algorithm that recommends Nash equilibrium actions

An algorithm that gives non-binding recommendations to all market participants does not change the game's action space and payoff functions. The recommendations do not provide fundamentally new information either and are non-binding.

However, collusion can be challenging to attain without a coordination device (Engel, 2007; Fonseca and Normann, 2012). Players must agree on a common collusive strategy. Within our setup, the collusive price is not necessarily a price of p^M but could be any

price above p^N . Moreover, collusion also depends on a shared understanding of how to punish deviations from the collusive price. It includes a punishment price but also an understanding of how many periods this price is set before, possibly, the players return to a collusive price. Recommendations can act as a coordination device that addresses all these issues. The idea behind a recommendation algorithm is that sellers may expect other sellers to behave according to the recommendation. It makes it incentive-compatible to do the same.

The following algorithm, labeled RECTHEORY, recommends prices according to the trigger strategy derived above.

Algorithm 1 (RECTHEORY).

- If the regime is collusive in the current period, recommend a price of p^M .
- If the regime is punitive in the current period, recommend a price of p^N .
- In period one, start in the collusive regime.
- Afterwards, if the regime was collusive in the previous period and everyone set a price of p^M , continue with the collusive regime in the current period.
- If in the previous period the regime was collusive but someone set a lower price than p^M , switch to the punishment regime for T periods and switch back to the collusive regime afterward.

It is sensible for the competing sellers to follow the recommendations, provided it is individually rational. We inform the subjects in the experiment that the recommendation algorithm's objective is to maximize industry profits symmetrically, that is, without favoring a particular seller. Inducing high sales prices is consistent with the incentives of a platform that receives part of the sellers' revenues through commission payments. If a seller expects the other two sellers to follow the algorithm's recommendations, then it is best off in doing the same, as this constitutes a subgame perfect Nash equilibrium. A deviation from RECTHEORY is not profitable provided that the other sellers follow the recommendation and play the static Nash price of p^N in the punishment periods, which is again a mutually best response.

We test the following hypotheses in the experiment based on those considerations.

Hypothesis 1. *Recommendations positively influence individual prices. A higher recommended price leads to higher individual prices.*

Hypothesis 1 states that firms' prices are increasing in the recommendation. As the recommendation may act as a coordination device, we expect that firms factor it into their pricing decision, and we hypothesize that higher recommendations lead to higher individual prices. It is a minimal requirement for any sensible algorithm to have a collusive effect.

Hypothesis 2. *The RECTHEORY recommendation algorithm leads to higher market prices than the absence of a recommendation algorithm.*

Hypothesis 2 builds on the rationale that RECTHEORY acts as a coordination device among the firms and thereby indeed facilitates collusion.

2.2.3 Behaviourally motivated soft punishment algorithms

Empirical and experimental evidence indicates that punishment is often less harsh than in theory models with trigger or even grim-trigger strategies. For instance, Wright (2013) finds that only a small fraction of subjects in market experiments use optimal or grim punishment strategies. Most punishment strategies are softer and more gradual. It concerns both the punishment length, as well as by how much prices are reduced in a punishment phase. Similarly, Dal Bó and Fréchette (2019) show that humans often use tit-for-tat strategies in the iterated prisoners' dilemma, which is strategically similar to our stylized market environment.

To reflect these practices, we set up a behaviourally motivated recommendation algorithm. It works as follows:

Algorithm 2 (RECSOFT).

- Start with a recommendation of the monopoly price of p^M and continue with this recommendation in future periods as long as all sellers adhere to the recommendation.
- In case of a deviation, recommend a punitive price equal to the lowest price from the previous period (e.g. $\min(10,10,9)=9$).

- If all sellers play the same price in a given period, recommend the monopoly price of p^M in the next period.

In line with the behavioral insights cited above, such a recommendation mechanism may be superior to the algorithm implementing a subgame perfect Nash equilibrium with trigger strategies. The recommendation is similar to a tit-for-tat algorithm as it mimics the firms' decisions in the previous period. However, it also proactively tries to increase prices after firms agree on a joint price level.

Following the recommendation might be behaviourally attractive as no harsh punishment needs to be implemented. With k -level reasoning, for instance, a seller might rationalize that other sellers prefer to punish if it bears little costs and it yields an expected price soon after. Suppose sellers anticipate punishment under the current soft punishment algorithm. In that case, it may deter them from departing from the collusive price.⁹ Furthermore, if sellers deviated in the past, the algorithm promotes cooperation as it again recommends the monopoly price once sellers agree on a joint price level. Since collusion at the monopoly price is the long-run objective of the algorithm, we argue that

Hypothesis 3. *The RECSOFT recommendation algorithm leads to higher market prices than the absence of a recommendation algorithm.*

It is noteworthy that following these recommendations does not constitute a subgame perfect Nash equilibrium. To see this, suppose that all sellers follow the recommendations throughout the game. If seller A follows the recommendations, the per-period profit is $p^M \cdot k/n$ in each period, yielding a profit stream of

$$p^M \cdot k/n \cdot (1 + \delta + \delta^2 + \delta^3 + \dots).$$

Consider a one-shot deviation of setting a price of $p_A < p^M$ while the algorithm recommends a price of p^M . The profit in the deviation period equals $p_A \cdot k$. The algorithm recommends a price of p_A in the next period. All sellers that follow the recommendation receive a profit of $p_A \cdot k/n$. Afterward, the algorithm reverts to the monopoly price of p^M . Thus, the deviating seller obtains a deviation profit of $p_A \cdot k$ for one period and a

⁹We also consider a recommendation algorithm without any punishment in the Appendix 2.C.3.

punishment profit of $p_A \cdot k/n$ for another period. Hence, the profit stream is

$$p_A \cdot k + p_A \cdot k/n \cdot \delta + p^M \cdot k/n \cdot (\delta^2 + \delta^3 + \dots),$$

which is highest for the highest feasible deviation price of $p_A = p^M - 1$. The difference between the deviation profit stream and the collusion profit stream is

$$k \cdot (p_M \cdot (1 - 1/n) - 1 - 1/n \cdot \delta).$$

Thus, deviating from the recommendation is profitable if

$$p^M > \frac{n + \delta}{n - 1}.$$

For $n \geq 2$ and $\delta < 1$ the condition holds for any $p^M > 2$. Thus, following the recommendation does not constitute a subgame perfect Nash equilibrium for the parameters used in the experiment as $p^M = 10$.

Following the soft recommendation algorithm may nevertheless be more attractive than the recommendation involving Nash reversion in punishment phases. It depends on the willingness of the sellers to implement drastic and longer-lasting punishments and their belief about the behavior of the other market participants. Nevertheless, sellers might find the soft punishment not harsh enough. Which recommendation algorithm performs better thus remains an ex-ante open question. We, therefore, do not postulate a hypothesis in this regard.

2.3 Experimental design

To experimentally investigate the collusive effect of price recommendations, we consider a market setup that mimics the theoretical framework in Section 2.2. Each of the $n = 3$ sellers in a market is represented by a participant. The market size is chosen such that tacit collusion is unlikely without any recommendation (Huck et al., 2004). The demand side consists of $k = 30$ computerized consumers. The participants play a repeated game. In each round of the game, all participants choose their prices independently. There is no direct communication between the participants. Across treatments, we vary whether

participants receive a price recommendation and which type of algorithm provides this recommendation. Each participant in a market receives the same price recommendation. After each participant selects a price, the participants receive information about the pricing decision of the other participants and their payoff in the given round. Furthermore, the recommended price is again shown to the respective treatment participants.

2.3.1 Treatments

There exists a BASELINE treatment in which we do not provide any price recommendation to the participants. Furthermore, we consider two treatments with rule-based price recommendation algorithms that are motivated by our theoretical considerations (Section 2.2). In the RECTHEORY treatment, the initial price recommendation is the monopoly price of $p^M = 10$. Any deviation from the recommendation by any participants triggers a punishment phase in which the stage game Nash equilibrium p^N is recommended. The punishment phase lasts for three periods. Afterward, the algorithm recommends the monopoly price again. Following the analysis in Section 2.2, RECTHEORY recommends actions that constitute a subgame perfect Nash equilibrium. In the RECSOFT treatment, the algorithm also recommends a price of $p^M = 10$ in the initial round. However, after a deviation, the recommended price is the lowest price from the previous period. If all participants choose the same price in a given round, the algorithm recommends the monopoly price again. In addition to the main recommendation algorithms (RECTHEORY and RECSOFT), we consider two additional mechanisms as a robustness check. In RECSTATIC, the algorithm provides a static price recommendation at the monopoly price similar to Sonntag and Zizzo (2015). Additionally, we analyze an algorithm similar to RECTHEORY but with a shorter punishment phase. Both additional algorithms do not foster collusion compared to BASELINE, and we only discuss them in the Appendix.

We focus on rule-based algorithms as they are highly tractable and allow us to derive clear, theoretical guided hypotheses that we developed Section 2.2. Furthermore, in digital platform markets, many algorithms are simple as well. Wieting and Sapi (2021) and Musolff (2022) show that real-world pricing algorithms are often rule-based and follow straightforward conditional processes. Moreover, although alternative methods like reinforcement learning algorithms have more complex routines to learn a pricing strategy, they eventually often converge to strategies that simple rules can describe (Werner, 2021;

Klein, 2021).¹⁰ Hence, our focus on those algorithms is attractive from a methodological perspective but also realistic regarding the tools used in actual markets.

2.3.2 Procedure

The experiments were conducted between February 2020 and August 2021 in the DICE Lab of the University of Duesseldorf. We used ORSEE (Greiner, 2015) to recruit the subject for the experiments. The experiment was programmed in oTree (Chen et al., 2016a). We utilized a between-subject design, and each subject only participated once.

At the beginning of each experiment session, participants were randomly assigned to a computer in the lab and could read the instructions on the computer screen. Additionally, the participants received a printed version of the instructions. The instructions were the same for each subject. A translated version of it is in Appendix 2.B.1. After the subjects read the instructions, they answered several control questions to ensure they understood the setup.¹¹ In case a participant failed to answer all control questions correctly, the software asked the participant to reread the instructions and allowed the participant to ask the experimenter clarifying questions in private.

In RECTHEORY and RECSOFT, the instructions describe the objective of the algorithms to the participants. To be precise, we explain that the recommendation algorithm aims to increase long-term industry profits. One control question specifically assesses whether participants comprehend the design purpose of the algorithm. Thus, we expect that all participants have the same belief about the algorithm’s objective.

Furthermore, the instructions emphasize that the price recommendation is non-binding so that each subject can choose a different price. This approach is motivated by the price suggestions that sellers receive in popular online marketplaces like Airbnb.

To mimic an infinitely repeated game as outlined in Section 2.2, each round of the game has a continuation probability of 95%. Thus, with a probability of 5% each game terminates after a given round. Within this setup, the continuation probability is equivalent to the discount rate of $\delta = 0.95$ (Roth and Murnighan, 1978). The game is repeated

¹⁰The Q-learning algorithms in Klein (2021) punish for a certain number of periods before reverting to the monopoly price. In Werner (2021), they learn one-period punishment strategies similar to a win-stay lose-shift strategy.

¹¹All control questions are in Appendix 2.B.2.

for three supergames to observe possible learning effects.¹² Within each supergame, the group composition is fixed. We use a perfect stranger matching scheme across supergames. Hence, the participants know they will meet each participant only once during the entire experiment. It rules out possible reputation effects that might arise otherwise. At the end of the experiment, the participants answered different survey questions that are listed in Appendix 2.B.2.

Table 2.1: Number of observations by treatment

Treatment	Number of participants	Number of independent observations		
		Supergame 1	Supergame 2	Supergame 3
BASELINE	54	18	6	6
RECSOFT	54	18	6	6
RECTHEORY	54	18	6	6

Note: The number of independent observations in later supergames is determined by the matching group size which always consist of nine participants.

In total, we distributed 162 participants evenly distributed across the three main treatments.¹³ This corresponds to 18 independent observations in the first supergame, as a market always has three sellers. In the later supergames, there are six independent observations by treatment.¹⁴ Table 2.1 contains an overview. We used an experimental currency unit (ECU) with an exchange rate of 100 ECU = EUR 1. On average, the participants received a payoff of EUR 10.73 plus a show-up fee of EUR 4.¹⁵ The average session length was 45 minutes.

2.4 Results

In this section, we discuss the experiment's results and test the hypotheses we derived in Section 2.2.

¹²The exact number of rounds was pre-drawn with a random number generator to allow for the same supergame length across different experimental sessions. The round numbers are 27 (Supergame 1), 8 (Supergame 2), and 18 (Supergame 3).

¹³For details on the additional control treatments see Appendix 2.C.3.

¹⁴The number of independent observations in later supergames is determined by the matching group size. A matching group consists of nine participants.

¹⁵During the COVID-19 pandemic, we paid each participant an additional EUR 4. This bonus was announced after the end of the session. Thus, it does not influence the behavior in the experiment itself.

2.4.1 The influence of price recommendations on individual prices

Hypothesis 1 states that price recommendations influence individual prices as participants base their pricing decision on the recommendations. To test this hypothesis, we regress the individual prices (p_t^i) on the recommended prices (p_t^R). The results of the linear regressions are in Table 2.2.

Table 2.2: Individual prices explained by the recommendation in a linear regression

Dependent Variable: Model:	Individual price (p_t^i)			
	(1)	(2)	(3)	(4)
<i>Variables</i>				
(Intercept)	2.77*** (0.406)	-0.201 (0.153)		
p_t^R	0.376*** (0.075)	0.203*** (0.026)	0.385*** (0.083)	0.224*** (0.038)
p_{t-1}^i		0.554*** (0.029)		
p_{t-2}^i		0.223*** (0.013)		
RECTHEORY				0.348 (0.713)
$p_t^R \times \text{RECTHEORY}$				0.346*** (0.078)
Further controls:			Yes	Yes
<i>Fixed-effects</i>				
Round			Yes	Yes
Supergame			Yes	Yes
Observations	5,724	5,076	5,724	5,724
<i>Clustered (Matching group) standard-errors in parentheses</i>				
<i>Signif. Codes: ***: 0.01, **: 0.05, *: 0.1</i>				

In all four columns, price recommendations positively and significantly affect individual prices. The effect is maintained when we control for lagged prices (column 2) and time-fixed effects (column 3). Furthermore, the effect size is more extensive for RECTHEORY than for RECSOFT (column 4). In specifications 3 and 4, we furthermore control for a set of individual-specific control variables.¹⁶ We conclude that the recommendations positively influence the prices. This is in support of Hypothesis 1.

¹⁶These include economic preferences and measures the socioeconomic status. We provide a list in the Appendix 2.B.2.

Result 1. *Sellers condition their prices on the recommendation of the algorithms. Price recommendations positively influence the individual sales prices of the participants.*

In all regression specifications, the coefficient of the price recommendation is below one. It indicates that the price recommendation only translates partially into the individual price. Increasing the price recommendation by one only increases the individual price by 0.20 to 0.57, depending on the model specification and treatment. Thus, albeit prices change with the recommendations, it appears that, on average, participants do not fully follow the recommendation.

2.4.2 Collusive effects of price recommendations

Building on the finding that subjects use the algorithms' recommendations for their pricing decisions, we now investigate whether the recommendations effectively foster collusion. Therefore, we compare the mean market prices in the treatments featuring recommendations with outcomes in the baseline treatment of no price recommendations. Note that the market price has a 1:1 relation with industry profits, so an analysis of the market price is equivalent to an analysis of the profits.

According to Hypotheses 2 and 3, price recommendations foster collusion as they provide a common reference point and simplify coordination on common punishment strategies after the deviation of a firm.

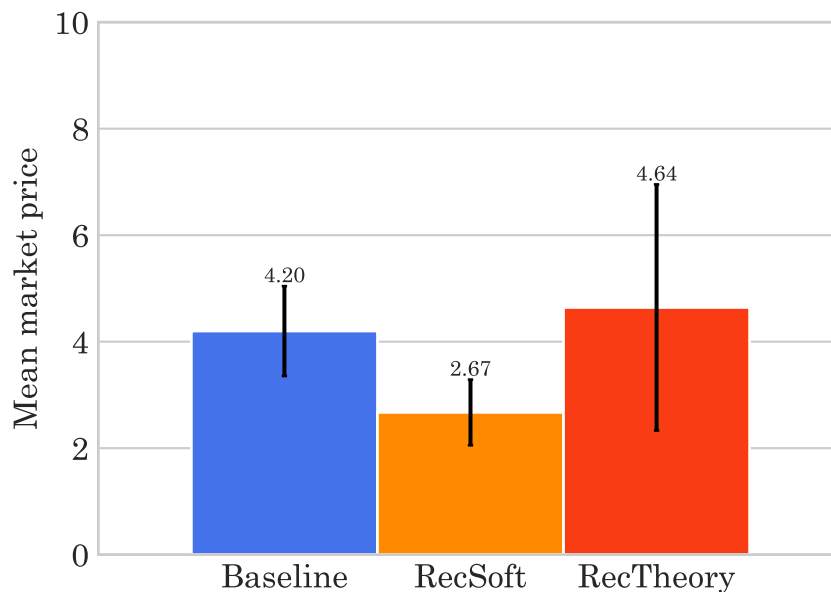


Figure 2.1: Market price for the main treatments. The error bars represent 95% confidence intervals.

Figure 2.1 shows the mean market prices by treatment pooled across supergames.¹⁷ The average market prices in BASELINE and RECTHEORY are similar. There are no statistically significant differences (p-value= 0.818, two-sided Mann–Whitney U test). Thus, we find no evidence that, on average, the RECTHEORY recommendation algorithm fosters tacit collusion.

One of our initial conjectures was that the RECTHEORY recommendation algorithm, while constituting a subgame perfect Nash equilibrium, might feature too harsh punishments from the perspective of human players. We, therefore, designed the softer recommendation algorithm RECSOFT. On balance, this algorithm, however, does not foster tacit collusion either. In fact, the market prices are on average *smaller* than in BASELINE (p-value< 0.05, two-sided Mann–Whitney U test). In other words, the algorithm makes competitive market outcomes more likely even though the initial design objective was to make markets more collusive. It is in contrast to the consideration provided in Section 2.2 and to Hypotheses 2 and 3. Furthermore, average market prices in RECSOFT are also smaller than in RECTHEORY (p-value< 0.1, two-sided Mann–Whitney U test). Thus, the game theory based algorithm is preferred if an upstream firm wants to use a price recommendation algorithm to foster collusion in the downstream market.

Table 2.3 displays the results from a linear regression of the market price on the different treatment indicators. The average market prices in RECTHEORY are higher than in BASELINE, but the standard errors are relatively large, so the differences are not statistically significant. In line with the results from the non-parametric test, the market prices are lower in RECSOFT than in BASELINE. The effect is robust to the inclusion of time-fixed effects and to the use of different aggregated survey measures.¹⁸

Result 2. *RECSOFT leads to statistically significantly lower prices than BASELINE. Price recommendations in RECTHEORY do not foster tacit collusion.*

Result 2 summarizes the findings about the average treatment effects. While we find support for the hypothesis that recommendations influence individual prices (Result 1), we find no evidence that, on average, price recommendations foster tacit collusion. On the contrary, price recommendations can make markets more competitive and lower

¹⁷The results by supergame do not differ substantially and are provided in Table 2.12 in the appendix. Furthermore, we provide an overview of the development of market prices across time in Figure 2.2.

¹⁸The survey measures were elicited on an individual level and listed in Appendix 2.B.2. We aggregate them on the group level by calculating the mean across all group members.

Table 2.3: Linear regression for the treatment effects

Dependent Variable:	Market price		
Model:	(1)	(2)	(3)
<i>Variables</i>			
(Intercept)	4.20*** (0.308)		
RECSoft	-1.53*** (0.382)	-1.53*** (0.384)	-1.45** (0.530)
RECTheory	0.442 (0.899)	0.442 (0.904)	0.823 (1.00)
Further controls			Yes
<i>Fixed-effects</i>			
Round		Yes	Yes
Supergame		Yes	Yes
Observations	2,862	2,862	2,862

Clustered (Matching group) standard-errors in parentheses
*Signif. Codes: ***: 0.01, **: 0.05, *: 0.1*

market prices relative to a baseline without any recommendations. Note that lowering the sales prices might be a desirable strategy for a sales platform if double marginalization is an issue (see Appendix 2.A for a theoretical illustration). There is empirical evidence whereby the sales platform Amazon appears to make offers with low prices more prominent in certain circumstances (Chen et al., 2016b; Hunold et al., 2022). It could also be part of a dynamic business strategy that attempts to invest in a large consumer base first to possibly charge higher prices on many locked-in consumers later. In the following section, we explore the mechanism for the price-decreasing effects of the soft recommendation algorithm. Furthermore, we discuss heterogeneous treatment effects for the RECTheory treatment.

2.4.3 Heterogeneity and mechanisms

The average treatment effects are in sharp contrast to our theoretical predictions. Price recommendations do not raise market prices in the mean and can even have pro-competitive effects. In the following, we discuss potential mechanisms that can explain those findings. We first consider heterogeneity in market outcomes across treatments and highlight

specific stylized facts that drive the heterogeneity. Then, we examine price patterns that arise in the different treatments.

Heterogeneous response to recommendations

There are substantial differences in market outcomes in RECTHEORY across matching groups. While the variance in average market prices in BASELINE ($\sigma^2 = 0.65$) and RECSOFT ($\sigma^2 = 0.34$) is small, there exists a large variation in RECTHEORY ($\sigma^2 = 4.85$). Those differences in variances are statistically significant ($p < 0.05$, two separate Bartlett-tests).¹⁹ This indicates that the recommendation algorithm RECTHEORY, which recommends strategies that constitute a subgame perfect Nash equilibrium, fosters more heterogeneous market outcomes.

To study the origin of the differences in variances, we show the maximal, minimal, and median average market price across matching groups for each treatment in Table 2.4. In line with the previous analysis, the median market price in RECSOFT is small, and the maximal price is even below the median of the other treatments. Interestingly, although the median prices in BASELINE and RECTHEORY are similar, market prices are more spread out in RECTHEORY than in BASELINE. The recommendations in RECTHEORY make specific markets more collusive, whereas they make others more competitive.

Table 2.4: Market price statistics by treatment

	BASELINE	RECTHEORY	RECSOFT
$\bar{p}_{max}$	4.94	7.96	3.5
$\bar{p}_{median}$	4.42	4.71	2.77
$\bar{p}_{min}$	2.62	1.45	1.71

We confirm this by dividing the observations for each treatment into subgroups that are above (HIGH) and below (LOW) the treatment-specific median market price. We observe that the average market prices for the RECTHEORY-HIGH subgroup are statistically significantly higher than in BASELINE-HIGH, although only at the 10% level (two-sided MWU test). Also, the market prices in BASELINE-LOW are higher than in RECTHEORY-

¹⁹As in the previous analysis, we aggregate the market prices at the matching group level. Thereby, we account for dependencies that arise by rematching participants at the end of each supergame. It allows for correct statistical inference. We provide an overview of the number of independent observations in Table 2.1.

Low. Nevertheless, those differences are not statistically significant, likely due to a lack of statistical power because of the sample split ($p=0.4$, two-sided MWU test).

Result 3. *The variance in market outcomes is larger in RECTHEORY compared to RECSOFT and BASELINE.*

Relationship between seller preferences and the effect of recommendations

To understand the origins of the heterogeneity in market outcomes, we regress the market prices on the different economic preference measures and interact them with the treatment variables. We focus on two variables that are critical for collusion to be sustainable from a theoretical perspective. First, we consider negative reciprocity. In the context of collusion in a market game, it is a natural measure to understand the willingness to punish deviations from a certain price level. Secondly, we analyze how time preferences interact with our treatments. Firms must be sufficiently patient for collusive strategies to be sustainable, as they have to value the long-run profits more than the short-term gain from deviating. Importantly, any heterogeneity is not driven by a lack of randomization but rather by differences in response to the treatment, conditional on distinct levels of those social preferences. We provide randomization checks in Table 2.7 in the Appendix.

We elicited the economic preferences on an individual subject level at the end of the experiment using the validated survey questions by Falk et al. (2022).²⁰ We apply a min-max normalization to all economic preferences on the individual level. Thus, all measures are between zero and one. Furthermore, we average them on the market level for the subsequent analysis.

Differences in negative reciprocity lead to vastly different market outcomes across treatments (see Table 2.5). In the BASELINE treatment without any price recommendations, higher degrees of negative reciprocity lead to lower market prices, as indicated by the negative coefficient of NEG. REC. in model specification 2. In other words, markets tend to exhibit lower market prices if the participants are more inclined to punish each other when they feel maltreated. For the treatments with price recommendations, this pattern is different. While the price level is lower for small levels of negative reciprocity in RECTHEORY and RECSOFT, as indicated by the negative coefficients of RECTHEORY and

²⁰Next to negative reciprocity and time preferences, the survey also includes positive reciprocity, time preferences, risk aversion, and measures of altruism and trust. We report the results regarding those variables in Appendix 2.C.2.

Table 2.5: Market price explained by negative reciprocity and treatments

Dependent Variable: Model:	Market price		
	(1)	(2)	(3)
<i>Variables</i>			
(Intercept)	2.63*** (0.810)	6.38*** (0.829)	
NEG. REC.	2.03 (1.40)	-3.40** (1.54)	-3.40** (1.55)
RECTHEORY		-2.44** (0.936)	-2.44** (0.941)
RECSOFT		-5.51*** (1.22)	-5.51*** (1.22)
NEG. REC. $\times$ RECTHEORY		4.54** (2.12)	4.54** (2.13)
NEG. REC. $\times$ RECSOFT		6.85*** (2.04)	6.85*** (2.05)
<i>Fixed-effects</i>			
Round			Yes
Supergame			Yes
Observations	2,862	2,862	2,862
<i>Clustered (Matching group) standard-errors in parentheses</i>			
<i>Signif. Codes: ***: 0.01, **: 0.05, *: 0.1</i>			

RECSOFT, the coefficients of the interaction terms with negative reciprocity are positive and statistically significant. Thus, as the degree of negative reciprocity increases, market prices in BASELINE become similar to the outcomes in RECTHEORY and RECSOFT.²¹ We interpret negative reciprocity as a willingness to punish deviations in this context. Thus, the recommendations harm collusion in markets with sellers that are usually unwilling to punish. Possibly, the recommendations lead to harsh punishments that would not have happened without them. If participants are unable to recover from the punishment, the recommendations reduce the market prices below the level that is observed in markets without recommendations but with similarly low levels of negative reciprocity. Those heterogeneous treatment effects can explain lower prices than in BASELINE for the treatments with a price recommendation.

²¹The average marginal effect of the treatment dummies is not statistically significant at the 10%-level if NEG. REC. is equal to one (Model specification 2 in Table 2.5). Note that one is the maximal value that NEG. REC. can take due to the normalization we apply.

Table 2.6: Market price explained by time preferences and treatments

Dependent Variable: Model:	Market price		
	(1)	(2)	(3)
<i>Variables</i>			
(Intercept)	0.595 (1.61)	-1.99 (2.52)	
TIME	4.57* (2.22)	9.06** (3.56)	9.06** (3.57)
RECTHEORY		8.25* (4.18)	8.25* (4.20)
RECSOFT		0.154 (3.47)	0.154 (3.48)
TIME $\times$ RECTHEORY		-11.2** (5.21)	-11.2** (5.24)
TIME $\times$ RECSOFT		-2.65 (5.01)	-2.65 (5.03)
<i>Fixed-effects</i>			
Round			Yes
Supergame			Yes
Observations	2,862	2,862	2,862
<i>Clustered (Matching group) standard-errors in parentheses</i>			
<i>Signif. Codes: ***: 0.01, **: 0.05, *: 0.1</i>			

Furthermore, the differences in time preferences amongst sellers lead to distinct market outcomes (see Table 2.6 where a higher level of TIME corresponds to more patience). In BASELINE and RECSOFT, the prices are more collusive for markets with more patient participants. This makes intuitive sense. For collusion to be sustainable, participants must value the long-run profits more than any short-term gains from possible deviations. This is arguably the case for groups of sellers who are more patient. For RECTHEORY, on the other hand, market prices are higher than in BASELINE if market participants are impatient. In other words, the recommendations foster collusion in situations where participants tend to deviate more due to their lack of patience. As the effect of TIME is negative in this treatment, the effect wears off for more patient participants, and market prices become similar to BASELINE for values of TIME close to one. For large values of TIME, the recommendation even has a negative effect on market prices compared to

BASELINE.²² Evidently, the recommendations lead to lower prices if sellers are particularly patient. It is possible that participants who are particularly patient would not have punished in the first place without the recommendation. Small deviations may lead to harsher punishment than usual. Result 4 summarizes our findings regarding negative reciprocity and time preferences.

Result 4. *Variations in economic preferences of negative reciprocity and patience can explain the heterogeneous market outcomes. Low negative reciprocity among market participants who receive price recommendations reduces collusion. The recommendation algorithm RECTHEORY makes markets more collusive if the sellers are impatient.*

Especially the markets with the RECTHEORY algorithm, which is motivated by our game theoretical considerations, outcomes depend on the sellers' degrees of negative reciprocity and patience. Those differences make intuitive sense and explain the considerable heterogeneity in market outcomes discussed in Result 3.

The result emphasizes that algorithms can be pro-collusive for particular subgroups, even though we do not find statistically significant pro-collusive effects on average. Hence, if platforms understand their users and target the recommendation to the specific population of sellers, the algorithm could increase the price level. As online sales platforms gather more and more data about their users, recommendations are more likely to be tailored to specific markets. Our results suggest that this could lead to an increased risk of collusion.

Price patterns across treatments

In Figure 2.2, we plot the market prices for each treatment by supergame and round. In the initial round, market prices in RECSOFT and RECTHEORY are higher than in BASELINE following the recommendation of $p_{t=1}^R = 10$ (p-value=0.052 & $p < 0.05$, two-sided Mann–Whitney U tests).²³ Yet, there are deviations from the recommendation in 86.1% of all markets in the first round. As a result, the treatment-specific punishment mechanisms take effect in the subsequent round.

²²The average marginal effect of RECTHEORY is negative and significant at the 5% level for TIME being equal to one.

²³Market prices in RECTHEORY and RECSOFT are similar in the first round following the same initial recommendation ($p=0.25$, two-sided Mann–Whitney U test), which confirms that randomization into treatments worked.

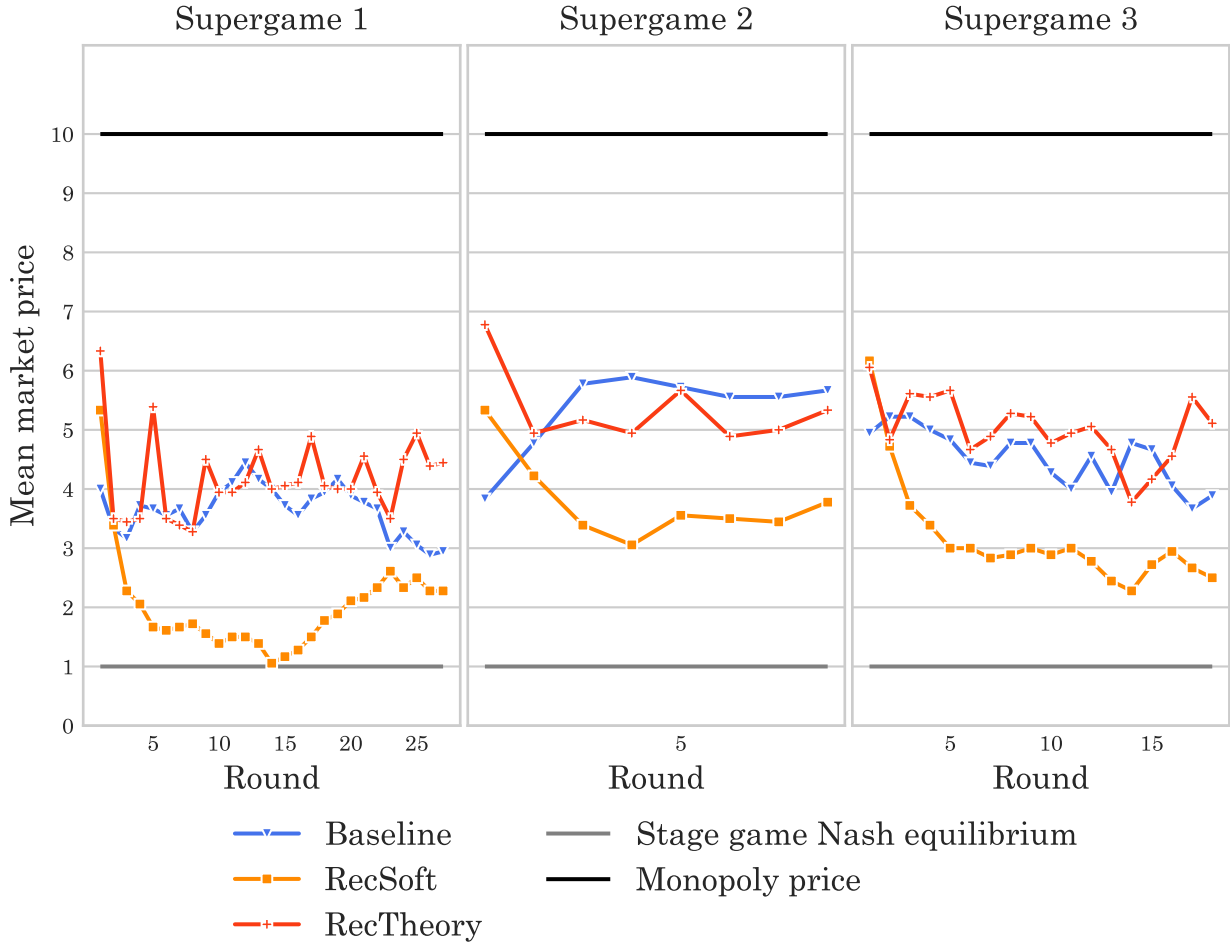


Figure 2.2: Market price for each treatment by supergame and round.

Let us focus first on the pattern of RECTHEORY in Figure 2.2. In response to deviations in the first round, the market prices drop for the following three rounds. At the end of this punishment phase, the prices increase sharply as the algorithm reverts to recommending the monopoly price. However, the prices do not stabilize completely at this level. In the following rounds, there are reoccurring deviations after a recommendation at the monopoly price. This results in clearly visible spikes in the price pattern. In the second and third super games, the spikes become less frequent, and the price patterns are more similar to BASELINE.

The recurring deviations in RECTHEORY are almost entirely driven by matching groups with below median market prices (RECTHEORY-LOW) as discussed in Section 2.4.3. It becomes clear when assessing the market price patterns for RECTHEORY for both subgroups separately in Figure 2.3.²⁴ Whereas there are deviations in both subgroups in the first round, the market prices stabilize in RECTHEORY-HIGH after the

²⁴For the respective analysis for RECSoft see Figure 2.5.

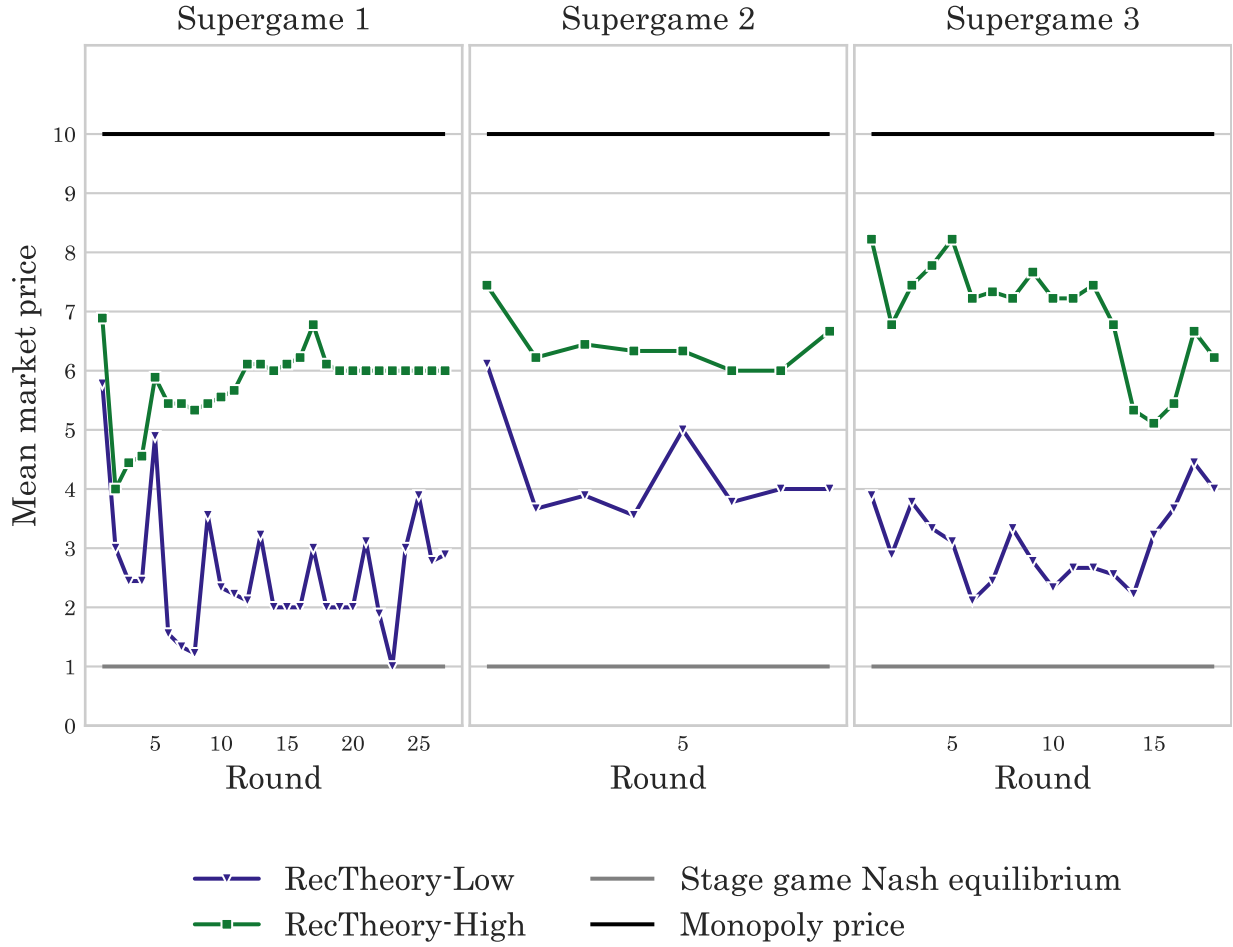


Figure 2.3: Market price in RECTHEORY for matching groups above (High) and below (Low) the median market price by supergame and round.

initial punishment phase. In RECTHEORY-LOW, the share of markets with deviations from the collusive recommendations is significantly higher after the first round, which results in price spikes (p-value<0.05, two-sided Mann–Whitney U test).²⁵

Matching groups with below-median market prices show repeated deviation patterns in the first supergame. They do not recover from this experience as average market prices remain lower throughout the rest of the experiment.²⁶ Hence, we find suggestive evidence that the recommendation in RECTHEORY works as expected for specific subgroups. However, other participants repeatedly deviate from the recommendation, which leads to lower market prices than in BASELINE for this subgroup.

²⁵We test this by restricting the data to the first supergame and to rounds in which the monopoly price was recommended. Then, we calculate for each market in RECTHEORY-HIGH and RECTHEORY-LOW the share of rounds in which at least one participant deviated from the recommendation. We test for differences in this variable across the two subgroups. Rematching only occurs after the first supergame, so each market constitutes an independent observation, allowing correct inference.

²⁶Similarly, Dal Bó and Fréchette (2018) show that participants' initial experience in the infinitely repeated Prisoners Dilemma is essential for their cooperation behavior in subsequent supergames.

For RECSOFT, the price patterns in Figure 2.2 are also interesting. We designed this recommendation algorithm to be forgiving to slight deviations as it does not immediately punish at the stage game Nash equilibrium price of $p^N = 1$ (see Section 2.2). We expected the punishment to be softer and, possibly, short compared to RECTHEORY. Yet, the data does not support this claim. After an initial deviation from the recommended price, the following recommendation is usually above the stage game Nash equilibrium ($\bar{p}_{t=2}^R = 5.33$). Hence, in contrast to RECTHEORY, there are again profitable deviations from the recommendation in the following round. Participants repeatedly deviate from the recommendation. This triggers a downward spiral as the recommendation for the next period is again the deviation price. There are, on average, 5.44 rounds with a recommendation below the monopoly price after the first deviation in the first supergame. This initial punishment period is significantly longer than in RECTHEORY, which always punishes for three periods (p-value < 0.05, one-sided one-sample t-test). Due to those frequent deviations from the recommendation, market prices deteriorate in the first rounds and only recover insufficiently in the subsequent rounds. As a result, the average prices in the treatment RECSOFT are low, and markets are even more competitive than in BASELINE.

Result 5. *Repeated deviations from the recommendation in RECTHEORY lead to lower market prices for specific markets. The recommendation in RECSOFT offers repeated deviation opportunities that drive the market prices down.*

Result 5 again emphasizes the adverse effects that recommendations can have for a platform if they are not designed appropriately. Furthermore, it suggests that platforms can use recommendations to decrease sellers' prices. In specific scenarios, this can be attractive, for instance, to avoid excessive double marginalization. Platforms could specifically design a recommendation algorithm to foster competition among the sellers. Our results suggest those price recommendations are feasible by using recommendation patterns as in RECSOFT.

2.5 Concluding remarks

Price recommendations are a vital feature of many digital. Companies like Airbnb and Expedia give non-binding price recommendations to the sellers operating on their plat-

forms. Furthermore, also in traditional markets, firms use recommendation algorithms to optimize their pricing decisions. While those pricing algorithms may be beneficial for the firms using them, competition authorities are concerned that those recommendations can dampen competition by helping firms to coordinate on non-competitive prices (Competition & Markets Authority, 2021; Bundeskartellamt and Autorité de la concurrence, 2019). For instance, a recent report on the U.S. rental market suggests that such algorithms may lead to higher rental prices by enabling coordination between landlords.²⁷

We derive two rule-based recommendation algorithms and study their effects on seller collusion in a stylized Bertrand market environment. Both algorithms have the objective foster collusion compared to a baseline without any recommendation. The recommendation of the RECTHEORY algorithm uses harsh punishment phases after deviations from the recommended price and aims at implementing a subgame perfect Nash equilibrium. Motivated by experimental evidence, we also design a recommendation algorithm (RECSOFT) that recommends softer punishments after a seller deviates from the collusive price. We test both algorithms in a laboratory experiment in which each participant represents a seller.

We find clear evidence that the recommendations influence the sales prices in the sense that higher recommended sales prices induce sellers to set higher individual prices. The estimated "pass-on rate" from recommended prices to sales prices is between 0.22 and 0.57, depending on the recommendation algorithm. This pass-on rate is higher for RECTHEORY. However, the effects on the realized market prices differ sharply between the different recommendation algorithms.

The algorithm RECTHEORY, which recommends collusive trigger strategies with temporary Nash reversal, does not lead to higher prices on average. However, we find extensive and interesting heterogeneities in market outcomes. The collusive effects depend on the seller's characteristics. RECTHEORY lowers market prices in markets with low levels of negative reciprocity among sellers. Moreover, we find evidence that the recommendation can make markets more collusive if sellers are too impatient to commit to a collusive pricing strategy.

²⁷See Vogell, Coryne & Little, "Technology Rent Going Up? One Company's Algorithm Could Be Why.", <https://www.propublica.org/article/yieldstar-rent-increase-realtor-rent>, (accessed on November 25, 2022).

For the behaviourally motivated algorithm RECSOFT, which can recommend brief punishment phases with moderate price levels, we find lower market prices and profits compared to the case without any recommendation. Participants frequently deviate from the recommendation, which starts a downward spiral that lowers prices. Similarly to RECTHEORY, market prices are lower than in BASELINE for markets with sellers that have low negative reciprocity. There is no evidence that RECSOFT facilitates collusion for any subgroup. Yet, it can be used to foster competitive outcomes and lower market prices for consumers.

We view our research as one of the first steps in understanding the effects of price recommendation algorithms on seller competition in digital sales platforms and other marketplaces. We demonstrate theoretically and experimentally that recommendations can benefit a platform by influencing sellers' pricing in the platform's favor. Our experimental evidence indicates that recommendation algorithms can facilitate seller collusion if designed appropriately and if the sellers are rather impatient. Thus, recommendations may in particular foster collusion and harm consumers if sales platforms understand the sellers' characteristics and target the recommendation based on these characteristics.²⁸

In other cases, we find that recommendation algorithms may have no price effects or even decrease prices, despite being designed and intended to facilitate collusion. The finding is consistent with the theoretical insight that all players following this behaviorally motivated algorithm does not constitute a subgame perfect Nash equilibrium. One interpretation of the price-reducing effects of the algorithm with soft punishments is that platforms may be able to use recommendation algorithms to make the sellers' offers more competitive. Under certain circumstances, such as excessive double marginalization or a dynamic pricing strategy, this could be in the interest of a sales platform. A caveat applies as we told our experiment participants that the algorithm would aim at increasing prices and profits, in line with our expectations. On average, the opposite, however, turned out to be the case for this algorithm. Over time, sellers may thus lose trust in following the algorithm's recommendations. More research in this regard would be desirable.

While the results are, on balance, not alarming regarding the collusive risks of recommendation algorithms, we do provide reasons for potential concern. It is important to

²⁸For example, accommodation platforms may gather more and more data about their hosts and guests over time and thus could condition their recommendations on seller characteristics in specific local markets to make them more effective.

note that, in our experiment, the only purpose of the recommendations was to coordinate sellers. In practice, recommendations can provide additional information on demand or help with pricing more generally, which could make sellers more likely to follow them. We chose to abstract from these factors in order to isolate the pure coordination effect, but we suspect that the collusive potential of recommendations may be higher when there are other reasons for sellers to follow them. Therefore, we believe our experiment is relatively conservative in terms of demonstrating collusive effects. We consider it fruitful for future research to study the collusive effects of recommendations that incorporate these additional factors.

Appendix

Appendix 2.A Additional theoretical results

High commission rates versus collusive recommendations Let us demonstrate why a monopoly platform can achieve higher profits through collusive price recommendations for sellers than when just charging a high commission rate to the sellers. This comparison is relevant as many online sales platforms use commission rates and a natural question is whether a sales platform needs recommendations to achieve the desired seller price level. This analysis adds to the developing literature on seller collusion on online platforms, which focuses on settings where commission rates are sufficient for achieving high prices (Schlütter, 2022; Teh, 2022).²⁹ Teh (2022) has a related finding whereby it can be optimal for a platform to increase the seller margins through platform design, such as entry regulations, if that is value-generating. He does, however, not consider seller collusion and uses a different modeling approach.

In Section 2.2, we use a stylized model that abstracts from the platform and costs on the sellers' side to focus on the collusive algorithm. In this extension, we instead focus on contracting between the platform and sellers and how the optimal price level depends on the commission rates. There are $n \geq 2$ symmetric sellers who sell differentiated products with marginal costs c . Each seller i makes a profit of

$$\pi_i = (p_i - r \cdot p_i - c)q_i(p_i, p_{-i})$$

when selling on the platform which charges a commission rate of r . Demand q_i has the usual properties and, in particular, decreases in the own price p_i and increases in the price(s) p_{-i} of the competitors. A seller's (opportunity) costs of being active on the

²⁹Schlütter (2022) primarily studies price parity clauses in a market where sellers can alternatively sell only via their direct sales channel.

platform are $I \geq 0$, such that the seller participates if and only if

$$\pi_i \geq I. \quad (2.1)$$

We focus on situations where the platform wants to ensure that all sellers participate, so that condition (2.1) holds for all i .

A seller's first order condition with respect to its sales price is

$$(p_i - r \cdot p_i - c) \frac{\partial q_i(p_i, p_{-i})}{\partial p_i} + (1 - r)q_i(p_i, p_{-i}) = 0$$

and can be written as

$$(p_i - \frac{c}{(1 - r)}) \frac{\partial q_i(p_i, p_{-i})}{\partial p_i} + q_i(p_i, p_{-i}) = 0. \quad (2.2)$$

Let $p^*(r, c)$ denote the symmetric Nash equilibrium price that solves the above equation when all n sellers compete.

At the competitive sales price, the platform makes a profit of

$$\Pi(r) = r \cdot p^*(r, c) \cdot n \cdot q_i(p^*, p^*).$$

Platform profit maximization when sellers compete For $c = 0$ the platform cannot influence the price level with r as it disappears in the first order condition (2.2).

For $c > 0$ the implicit function theorem on the first order condition (2.2) for the case of symmetric sales prices $p_i = p_{-i} = p^*$ yields $\partial p^*/\partial r > 0$ under the standard assumptions of a strictly concave seller profit π_i in p_i and decreasing demand ($\partial q_i/\partial p_i < 0$). The platform can thus raise the price level as long as selling remains profitable for the sellers.

Suppose that the sellers make lower profits if their common input costs increase. This is consistent with economic intuition and holds under standard demand assumptions. A sufficient condition for this is that the sellers' profit margin decreases as the costs increase: $\partial p^*/\partial k < 1$ with $k = c/(1 - r)$. This is, for instance, the case with linear demand.

For illustration, suppose that $r = 0$ and that the resulting seller profits equal zero:

$$(p^* - c)q_i(p^*, p^*) - I = 0.$$

It is thus not feasible for the platform to charge a positive commission rate as the sellers would not break even. This argument generalizes to the case where break-even occurs at a positive commission rate that yields a price level $\hat{p}$, which is below the level which maximizes the industry profit. The platform is then restricted in the setting of the commission rate and thus cannot maximize the industry profit.³⁰

Conversely, it might be that the platform achieves the profit-maximizing industry price at a commission rate where the sellers make positive profits ($\pi_i > I$). This occurs if I is small enough. The platform then leaves more profits than necessary for participation to the sellers. It would thus be optimal for the platform to charge a higher commission rate while keeping the sales prices constant. This relates to the problem of double marginalization.

Platform profit maximization when sellers collude For simplicity, assume that the platform can implement any price level p through recommendations. The platform can thus implement a price p and set r such that

$$(p \cdot (1 - r) - c) \cdot q_i(p, p) = I.$$

The platform can thus implement the industry-maximizing price

$$p^M = \arg \max_p (p - c) \cdot \sum_{i=1}^n q_i(p, p) - n \cdot I$$

and extract through r all seller revenues up to $I + c \cdot q_i(p^M, q^M)$ per seller.

In summary, this analysis shows that a platform can benefit from prices recommendations even if it can change its commission rates for one of the following reasons:

- The sellers have (opportunity) costs of selling on the platform, such that a high commission rate is not acceptable but would be necessary for achieving high sales prices of competing sellers.
- The platform charges a commission rate and the sellers do not have marginal costs other than the commission payment, so that the commission rate does not affect the sellers' pricing.

³⁰Fixed fees might solve the problem. However, in particular, transfers to sellers might not work in practice. For instance, they might incentivize people to register as sellers just to obtain the transfers.

- A high commission rate is optimal to extract the seller's profits but yields too high sales prices (excessive double marginalization). In this case recommendations below the competitive level can be optimal.
- In addition to the above formal analysis, a platform might desire to charge the same commission rate across different markets to maintain a simple transparent policy albeit different seller price levels are optimal.

Appendix 2.B Instructions and survey questions

2.B.1 Instructions

Hello and welcome to our experiment. In the next hour, you will make decisions on a computer. Please read the instructions carefully. All participants will receive the same instructions. You will also find a printed copy of these instructions at your seat. You will remain completely anonymous to us and to the other experiment participants. We will not save any data associated with your name.

Particularly important: Do not talk to your neighbors, do not use your cell phone, and keep quiet throughout the experiment. If you have any questions, please let us know. We will then come to your site and help.

In this experiment, you will repeatedly make pricing decisions. These allow you to earn real money. How much you earn depends on your decisions and on those of your fellow players. **Regardless, you will receive 4.00 euros for participating.**

In the experiment, we use a fictional monetary unit called ECU. After the experiment, the ECU will be converted to euros and paid to you. **Here, 100 ECU equal one euro.** The euro amounts are rounded to the first decimal place.

Example:

Participant A earned 465 ECU in the experiment. Converted, this is equal to 4.65 euros. Rounded to the first decimal place, Participant A is paid 4.70 euros.

Explanations:

In this game, you represent a company in a virtual product market. In the market, two other companies sell the same product as you do. These companies are represented by two other experiment participants. The game has several rounds. You will meet the same companies (i.e. experiment participants) in each round of the game.

All companies decide again **independently and simultaneously** in each round, for how many ECU you want to sell your product. You can sell your product for a price of 1, 2, ... or 10 ECU to sell (whole units only). **There are no costs of production.** Your profit is the product of price and the number of units sold. In formal terms:

$$\text{Profit} = \text{price} \times \text{units sold.}$$

The market has 30 identical customers. Each customer wants to buy **one unit** of the product as cheaply as possible in each round of a game. Each customer is willing to spend up to 10 ECU for that unit of the product.

The company with the **lowest price in the respective round** sells its products. **So the lowest price is the market price in that round.** Firms with a **price greater than the market price do not sell any products in that round and therefore receive a profit of zero.** If two or all three firms want to sell their product for the same market price, the demand is split evenly between the two or three firms.

Examples **Example 1** Firm A sets a price of 4, Firm B sets a price of 4, Firm C sets a price of 6. Thus, Firms A and B together have set the lowest price. Firms A and B both sell the same amount of products, both firms have 15 customers each and thus get the same profit of 60 ECU. Firm C sells nothing and has a profit of 0.

	Firm A	Firm B	Firm C
Prices	4	4	6
Profits	60	60	0

Example 2: Firm A sets a price of 7, Firm B sets a price of 7, Firm C sets a price of 7. Thus, Firms A, B and C together have set the lowest price. They all sell the same amount of products (10 each) and thus get the same profit of 70 ECU.

	Firm A	Firm B	Firm C
Prices	7	7	7
Profits	70	70	70

Example 3: Firm A sets a price of 1, firm B sets a price of 4, firm C sets a price of 10. Thus, firm A has set the lowest price. Firm A is the only one that sells the product at a price of 1 to all 30 customers and thus gets a profit of 30 ECU. Firms B and C both sell nothing and have a profit of 0.

	Firm A	Firm B	Firm C
Prices	1	4	10
Profits	30	0	0

Price recommendations

Before you choose your price in each round, you receive a specific **price recommendation** from a computer algorithm. **All three firms in the market receive the same price recommendation.**

The algorithm aims to **maximize the total profits of all firms across all rounds**. Therefore, you will be given a recommendation that will allow all firms to make the highest possible profit in the long run. This means that the algorithm does **not** necessarily recommend a price that achieves the highest possible profit in a single round. **It recommends prices that achieve a high total profit over the entire game.**

The algorithm itself is not a market participant and cannot generate profits, **it only serves as information for all participants.**

Note: The recommended price is only a **proposal**. You are free to set any other price than the recommended one.

Duration of the experiment

After each round, all firms are informed about the chosen prices of all three firms and their own profits. In the next round, each firm has again the opportunity to choose their price. You interact with the same participants in each round within a game.

After each round, a random mechanism decides whether another round is played or the game ends. The probability that another round will be played is 95%. The game therefore

ends after each round with a probability of 5%.

In other words, the computer throws a virtual dice with 20 sides before each possible further round. The result decides whether another round is played or not. With the number 20, the game is over, with all other numbers, another round is played.

Note:

You play the described game a total of three times. After each game, you will be put together with new participants to form a new market. This means that in each of the three games you interact with other participants.

After all games are finished, it will be randomly decided which of the three games will be paid out. You will receive your payoff after the experiment. You will also receive an additional 4.00 euros for participating in this experiment.

As a help we display a virtual calculator, with which you can calculate your profits in each round. **Comprehension Questions**

Question 1: How many consumers are in the market who want to buy the product?

- 25
- 35
- 30
- 40

Question 2: What is the probability of playing another round after completing one?

- 95%
- 5%
- 50%

Question 3: You are firm A and choose a price of 2, firm B chooses a price of 10, firm C chooses a price of 9. What is your profit in ECU in this round?

Question 4: You are firm A and choose a price of 8, firm B chooses a price of 8,

firm C chooses a price of 8. What is your profit in ECU in this round?

Question 5: You have a profit of 650 ECU, what is your profit in euros?

Questioni 6: What is the objective of the algorithm?

- Maximizing profits for all firms in a single round
- Maximizing total profits for all firms across all rounds
- Maximizing total profits for a single firm across all rounds
- Maximizing profits of a single firm in a single round

2.B.2 Survey questions

Gender: What is your gender?

- Male
- Female
- Diverse
- No specification

Experiments: In how many economic experiments have you (approximately) already participated?

GPA (School): What was the final grade of your last school diploma (1.0 - 4.0)?

Math Grade: What was your last math grade (1.0 - 6.0)?

Budget: How much money do you have available each month (after deducting fixed costs such as rent, insurance, etc.)?

Spending: How much money do you spend each month (after deducting fixed costs such as rent, insurance, etc.)?

Risk: Are you generally a person who is willing to take risks or do you try to avoid risks? Please indicate your answer on a scale of 0 to 10, where 0 means not willing to take risks at all and 10 means very willing to take risks.

Time: Compared to others, are you generally willing to give up something today in order to benefit from it in the future, or are you unwilling to do so compared to others? Please indicate your answer on a scale of 0 to 10, where 0 means not willing to give up at all and 10 means very willing to give up something.

Trust: As long as I am not convinced of the opposite, I always assume that other people only have the best in mind. How strongly do you agree with this statement? Please indicate your answer on a scale of 0 to 10, where 0 means not true at all and 10 means very true.

Neg. Rec.: Are you someone who is generally willing to punish unfair behavior, even if it comes at a cost for you, or are you unwilling to do so? Please indicate your answer on a scale of 0 to 10, where 0 means not willing to punish at all and 10 means very willing to punish.

Pos. Rec.: If someone does me a favor, I'm willing to return it. How strongly do you agree with this statement? Please indicate your answer on a scale of 0 to 10, where 0 means not true at all and 10 means very true.

Altruism: Imagine the following situation: You won 1,000 € in a prize competition. How much would you donate to charity in your current situation?

Appendix 2.C Further results

2.C.1 Randomization checks

Table 2.7 provides the average outcome for different survey measures for the main treatments. Furthermore, we test for differences in those measures across treatments using Kruskal-Wallis tests. A complete list of the different survey questions we ask the participants is provided in Appendix 2.B.2. There are few negligible differences in the control variables across treatments. Only the budget participants have each month differs between treatment at the 5%-level. Importantly, controlling for those survey measures does not influence the main outcomes (see Table 2.2 and 2.3). Thus, we conclude that randomization into treatments worked as expected.

Table 2.7: Survey measures by treatment

	Risk	Time	Trust	Neg. Rec.	Pos. Rec	Altruism	Woman
BASILINE	0.53	0.68	0.39	0.64	0.92	0.12	0.48
RECSOFT	0.49	0.70	0.41	0.52	0.88	0.14	0.50
RECTHEORY	0.54	0.74	0.49	0.62	0.84	0.10	0.54
P-values	0.73	0.35	0.18	0.09	0.19	0.84	0.84

	Experiments	Math Grade	GPA (School)	Budget	Spending
BASILINE	7.24	1.99	1.99	414.98	300.15
RECSOFT	11.52	2.47	2.31	378.61	275.63
RECTHEORY	10.69	2.12	2.05	523.89	339.87
P-values	0.38	0.07	0.05	0.04	0.39

Note: The preferences measures are based on the survey questions by Falk et al. (2022) and scaled between zero and one. The p-values are based on Kruskal-Wallis-tests.

2.C.2 Economic preferences and recommendations

In Section 2.4, we discuss the influence of negative reciprocity on market prices when participants receive a price recommendation (see Table 2.5). Here, we provide the same analysis for the other economic preferences measures. All measures have been normalized to be between zero and one. Furthermore, as the measures have been elicited on the individual level, we aggregated them by calculating the group specific mean.

Altruism and trust do not influence market prices. Interestingly, there is also no significant effect of positive reciprocity on market outcomes. In other words, while differences

in negative reciprocity lead to vastly different prices within and between treatment, it is not the case for positive reciprocity.

Table 2.8: Market price explained by altruism and treatments

Dependent Variable: Model:	Market price		
	(1)	(2)	(3)
<i>Variables</i>			
(Intercept)	4.00*** (0.575)	4.18*** (1.11)	
ALTRUISM	-1.34 (3.34)	0.191 (8.56)	0.191 (8.61)
RECTHEORY		-0.209 (1.65)	-0.209 (1.65)
RECSOFT		-1.10 (1.14)	-1.10 (1.14)
ALTRUISM $\times$ RECTHEORY		6.43 (10.9)	6.43 (10.9)
ALTRUISM $\times$ RECSOFT		-3.14 (8.65)	-3.14 (8.70)
<i>Fixed-effects</i>			
Round			Yes
Supergame			Yes
Observations	2,862	2,862	2,862
<i>Clustered (Matching group) standard-errors in parentheses</i>			
<i>Signif. Codes: ***: 0.01, **: 0.05, *: 0.1</i>			

Table 2.9: Market price explained by positive reciprocity and treatments

Dependent Variable:	Market price		
Model:	(1)	(2)	(3)
<i>Variables</i>			
(Intercept)	0.833 (3.33)	-1.57 (4.97)	
POS. REC.	3.41 (3.62)	6.26 (5.53)	6.26 (5.56)
RECTHEORY		3.82 (5.54)	3.82 (5.57)
RECSOFT		-3.13 (7.41)	-3.13 (7.44)
POS. REC. $\times$ RECTHEORY		-3.43 (6.12)	-3.43 (6.15)
POS. REC. $\times$ RECSOFT		2.07 (8.31)	2.07 (8.36)
<i>Fixed-effects</i>			
Round			Yes
Supergame			Yes
Observations	2,862	2,862	2,862
<i>Clustered (Matching group) standard-errors in parentheses</i>			
<i>Signif. Codes: ***: 0.01, **: 0.05, *: 0.1</i>			

Table 2.10: Market price explained by risk preferences and treatments

Dependent Variable: Model:	Market price		
	(1)	(2)	(3)
<i>Variables</i>			
(Intercept)	2.46** (0.964)	1.84 (1.07)	
RISK	2.65 (1.55)	4.45* (2.16)	4.45* (2.17)
RECTHEORY		4.21 (3.03)	4.21 (3.05)
RECSOFT		-0.654 (1.56)	-0.654 (1.57)
RISK $\times$ RECTHEORY		-7.07 (4.81)	-7.07 (4.83)
RISK $\times$ RECSOFT		-1.45 (3.13)	-1.45 (3.15)
<i>Fixed-effects</i>			
Round			Yes
Supergame			Yes
Observations	2,862	2,862	2,862
<i>Clustered (Matching group) standard-errors in parentheses</i>			
<i>Signif. Codes: ***: 0.01, **: 0.05, *: 0.1</i>			

Table 2.11: Market price explained by trust preferences and treatments

Dependent Variable: Model:	Market price		
	(1)	(2)	(3)
<i>Variables</i>			
(Intercept)	3.32*** (0.863)	4.60** (1.71)	
TRUST	1.20 (1.73)	-1.01 (4.13)	-1.01 (4.15)
RECTHEORY		-0.835 (3.06)	-0.835 (3.07)
RECSOFT		-2.16 (1.77)	-2.16 (1.78)
TRUST $\times$ RECTHEORY		2.79 (5.58)	2.79 (5.61)
TRUST $\times$ RECSOFT		1.58 (4.38)	1.58 (4.40)
<i>Fixed-effects</i>			
Round			Yes
Supergame			Yes
<i>Fit statistics</i>			
Observations	2,862	2,862	2,862
R ²	0.00319	0.05525	0.09202
Within R ²			0.05736
<i>Clustered (Matching group) standard-errors in parentheses</i>			
<i>Signif. Codes: ***: 0.01, **: 0.05, *: 0.1</i>			

2.C.3 Additional control treatments

We consider two additional control treatments. In RECSOFT, participants receive a static price recommendation at the monopoly in each period. Also, after deviations from the recommended price, the algorithm recommends the monopoly price, and there is no punishment mechanism. Furthermore, we consider the RECNASH algorithms. Similar to RECTHEORY, after any deviation from the monopoly price, the stage game Nash equilibrium is recommended in the subsequent period. Yet, the algorithm reverts back to the monopoly after one punishment round and is thus, in contrast to RECTHEORY, not incentive compatible. We consider for both algorithms only 36 subjects, which yields considerably less power than in the main treatments. The results are provided in Figure 2.4. Both treatments yield similar market prices as in BASELINE and RECTHEORY.

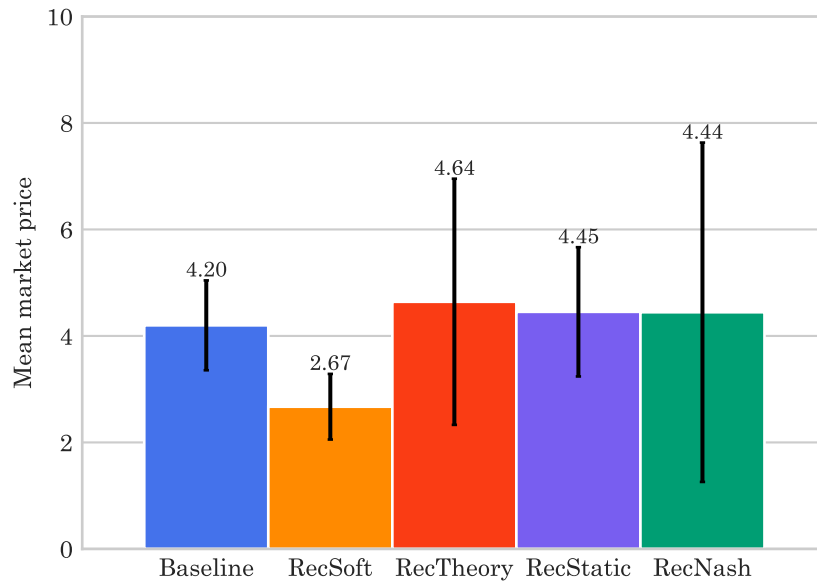


Figure 2.4: Market price for all additional algorithms. The error bars represent 95% confidence intervals.

2.C.4 Further figures and tables

Table 2.12: Linear regression with the treatment effects by supergame

Dependent Variable:	Market price		
	(Supergame 1)	(Supergame 2)	(Supergame 3)
<i>Variables</i>			
RECSoft	-1.63** (0.572)	-1.56* (0.881)	-1.36* (0.740)
RECTheory	0.539 (0.990)	-0.007 (1.27)	0.497 (1.26)
<i>Fixed-effects</i>			
Round	Yes	Yes	Yes
<i>Fit statistics</i>			
Observations	1,458	432	972
R ²	0.08748	0.04145	0.05861
Within R ²	0.07274	0.03805	0.04175

Clustered (Matching group) standard-errors in parentheses

*Signif. Codes: ***: 0.01, **: 0.05, *: 0.1*

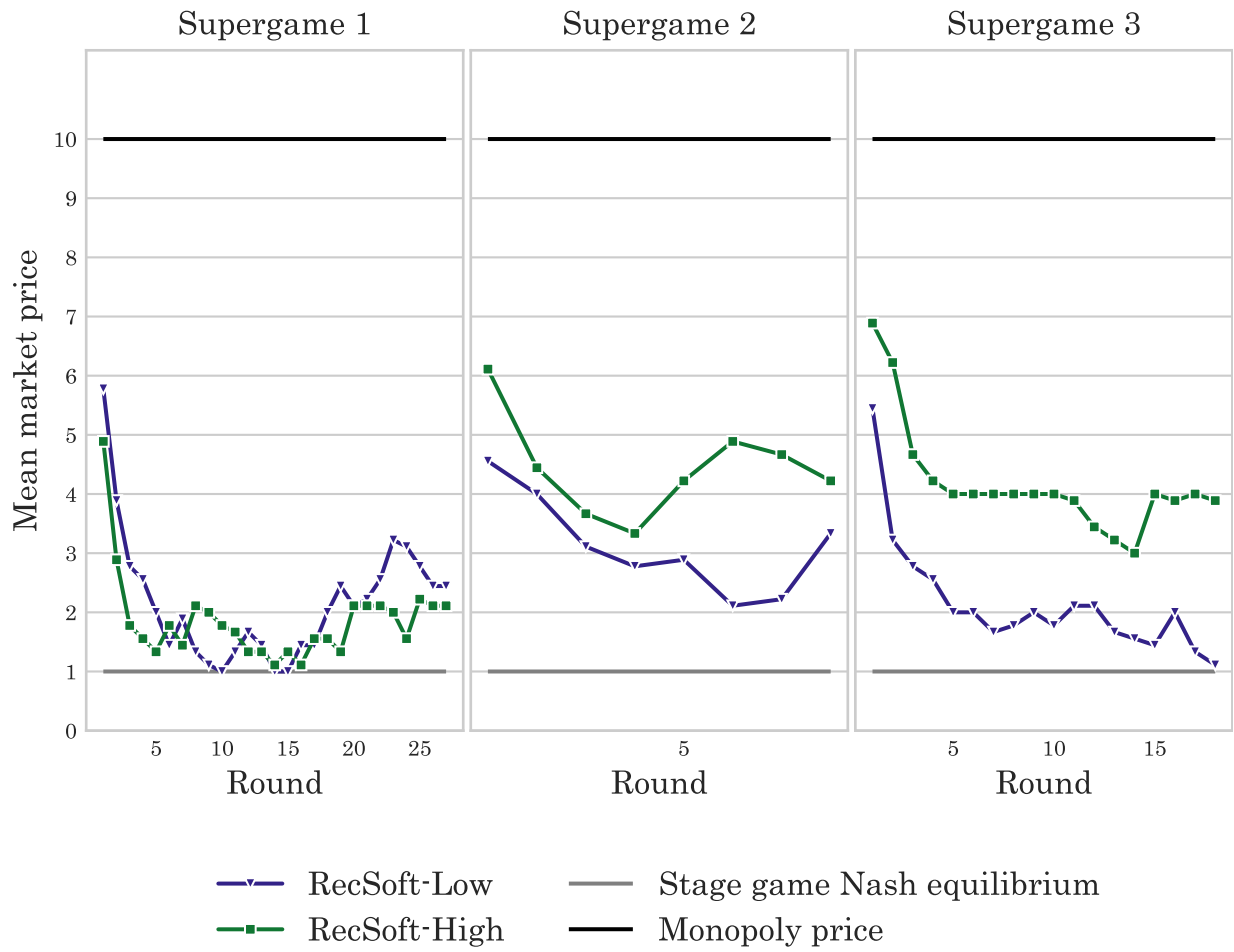


Figure 2.5: Market price in RECSoft for matching groups above (High) and below (Low) the median market price by supergame and round.

Declaration of Contribution

Hereby I, Tobias Felix Werner, declare that the chapter “Algorithmic Price Recommendations and Collusion: Experimental Evidence” is coauthored by Matthias Hunold. Both authors contributed equally to the chapter.

Signature of coauthor (Matthias Hunold):



Prof. Dr. Matthias Hunold

Chapter 3

Willingness to Volunteer Among Remote Workers is Insensitive to the Team Size

Co-authored with Adrian Hillenbrand and Fabian Winter¹

¹This chapter is an updated version of my master's thesis, which I wrote while studying at the Rheinische-Friedrich-Wilhelms-Universität Bonn. It includes new data and has been significantly extended.

3.1 Introduction

Volunteering is an important feature of the fundamental organization of firms. In various situations, tasks and resources are not allocated among employees by some supervisor, but rather employees have to solve the allocation process by themselves. In fact, especially with the rise of remote work, many tasks are organized more informally without a clear hierarchy and require the own initiative of team members. While about 4% of the workforce in Europe's biggest economy Germany worked from home in the years prior to 2020, up to 27% worked remotely in 2021 (Hans Böckler Stiftung, 2021). The US gives a similar picture, with 7.6% of exclusively remote working before the pandemic and peaking at 31.4% in mid-2020 (Bick et al., 2022). In these situations, volunteering is a natural allocation mechanism, which we want to examine in this chapter. Naturally, completing the task requires time and effort, so each team member would prefer another person to finish it. If the task is completed, the product advances, which yields a higher reputation to the team in the organization and may improve the firm's overall performance in the market. This, in turn, benefits the whole product team as it may improve wages or the job prospects of all team members.

While the described allocation process might seem inefficient due to the challenges of coordination and free-riding incentives, situations like this still frequently occur – and apparently for good reasons – in organizational contexts. For example, many duties in academia are allocated based on voluntary decisions (Babcock et al., 2017), just as the development of open-source software projects (Johnson, 2004), the contribution to network technology (Lee et al., 2007), the creation of online knowledge platforms like Wikipedia (Zhang and Zhu, 2011), or modern work allocation mechanisms like the so-called agile project methods, which are commonly used in software development (Hoda et al., 2018). Arguably, the substantial increase in remote work in the last years might have led to an increase in these allocation mechanisms in the work environment.

The volunteering mechanism does have attractive qualities. Not least, it may reduce the organizational overhead required to organize task allocations. Yet, volunteering in an organizational context is usually not an altruistic act towards others but often the individually profit-maximizing response to an organizational problem (Murnighan et al., 1993; Kim and Murnighan, 1997). Economic (game) theory thus suggests that the volunteering mechanism also introduces two significant obstacles. For one, it creates a coordination

problem that has to be resolved by the workers. More importantly, however, the individual incentive structure of the firm may give rise to a social dilemma: while working is costly, but the exact costs are usually unknown, all team members may enjoy the fruits of labor. This, in turn, leads to the famous Volunteer's Dilemma (Diekmann, 1985).

The strategic analysis of the Volunteer's Dilemma closely resembles several relevant factors for the success of volunteering choices in organizations. First, companies and the teams within the company may be of various sizes, which in turn influences the degree of volunteering. This has been robustly shown in various lab and field experiments on the topic with a mostly negative effect of group size on individual volunteering (Diekmann, 1986; Franzen, 1995; Goeree et al., 2017; Kopányi-Peuker, 2019; Latané and Nida, 1981; Przepiorka and Berger, 2016; Barron and Yechiam, 2002; Campos-Mercade, 2021). Second, the individual costs of volunteering might differ for different workers, which also affects individual volunteering choices (Diekmann, 1993; Przepiorka and Berger, 2016). While lab evidence suggests that both factors negatively affect the provision of voluntary work, the prevalence of volunteering in real-world organizations is striking. This begs the question of whether economic theory and experimental results from the lab translate into real-world work environments or whether other factors like peer effects or image concerns might counteract these problems.

In this chapter, we scrutinize these economic arguments against volunteering at the workplace and the existing empirical evidence on volunteering by putting them to the test in real-life workplaces. In a large-scale field experiment with more than 2000 workers, we analyze the prevalence of volunteering at the workplace and the effect of group size on volunteering behavior. Our main treatment manipulation is thus varying the group size of work teams: we compare the willingness to volunteer for a specific task when working alone to working in small (3 workers), medium (30 workers), and large groups of 300 workers.

In our field experiment, we act as an employer in an online labor market and offer a simple rating task to the online workers. Workers are assigned to a team of a certain group size, and after finishing the individual task, we offer each participant the opportunity to continue working on the task. If at least one person in a group volunteers for the additional team task, each team member receives an additional bonus payment, which resembles the Volunteer's Dilemma. Before the beginning of the team task, participants

receive information about their costs and the costs of others as we inform our workers about the time it has taken them to perform the individual task and how long it has taken others. Finally, we elicit workers' beliefs about whether they think another co-worker will volunteer.

Our experimental setting allows us to study the causal effect of differences in team size. By creating a natural yet anonymous work environment, we can exclude reputational concerns as well as personal relationships between workers, which would impede the analysis of volunteering at more traditional workplaces. The online labor market provides a unique ecosystem with tight control over the environment while allowing us to precisely measure individual opportunity costs, beliefs, and other essential variables. We also contribute to the literature on the Volunteer's Dilemma by providing an application in a real work environment. Compared to existing experiments where effort is purely monetary, we can capture more subtle differences using actual work tasks. Importantly, workers in our study do not learn whether others have done the task before. While this differs from many classical work environments where the task would be announced as being completed, this crucial design allows us to measure workers' individual willingness to volunteer, which would not be possible in a work environment where one can only observe whether there was a volunteer or not.

The results of our study stand in stark contrast to the game-theoretical predictions and results from earlier laboratory studies and field settings outside the work context. We find no support for the hypothesis that the group size influences the volunteering decision of our workers. More precisely, in groups of 3, 30, or 300 workers, the volunteering rates range between 51% and 55% and are not statistically different from each other. Also, the costs of volunteering play a negligible role, given our proposed cost measure. We replicate our results and show they are robust to multiple potential factors. Additional control treatments in which workers work on the same task but without strategic interaction, i.e., where each worker decides to volunteer on his own and is paid accordingly, help us to rule out several possible explanations: workers are less likely to volunteer if volunteering is not compensated, suggesting that the effort is indeed costly. Also, they are more likely to volunteer if their payment only depends on their own actions, suggesting that they make a difference between strategic and non-strategic situations. Thus, workers react to

the strategic interaction in the Volunteer’s Dilemma setting but do not react to the team size.

We explain these surprising results based on workers’ beliefs about their co-workers’ volunteering choices. We asked workers how likely it is that at least one other worker in their team also volunteers. We find that those who believe that it is more likely that there is at least one other volunteer are *more* likely to volunteer themselves. Image concerns can explain this seemingly irrational behavior. Subjects that believe that many others probably will volunteer do not want to be perceived (by themselves or us as an employer) as selfish. We show that this form of *conditional volunteering* behavior is the critical driver of volunteering in our work setting.

Overall, our results thus suggest that if a company only cares about the task being completed, volunteering might be justified as an allocation mechanism also in large groups, even though theoretical reasoning would predict otherwise. A key difference in our setting compared to the literature on the Volunteer’s Dilemma is that an employer is present, potentially increasing the importance of image concerns. According to our results, this is true even without actual reputation concerns or other punishment mechanisms for non-volunteering. This also suggests that in more anonymous environments, e.g., due to a shift to more remote work, volunteering rates do not necessarily decline as long as workers have a certain regard for how their work (or not working) is perceived.

The remainder of the chapter is structured as follows. In Section 3.2, we explain the experimental design of the field experiment. Following this, we present our pre-registered hypotheses in Section 3.3.² Our hypotheses are empirically tested in Section 3.4. Section 3.5 discusses our findings and concludes the chapter.

3.2 Experimental design

We study volunteering in a workplace environment through the lens of the Volunteer’s Dilemma (Diekmann, 1985; Darley and Latané, 1968). In the Volunteer’s Dilemma, a certain number of participants can volunteer to supply a public good to all group mem-

²The hypotheses are derived using a formal model of volunteering under cost and population uncertainty by extending the work by Weesie (1994) and Hillenbrand and Winter (2018). Details are provided in Appendix 3.A. The pre-registration of the field experiment and the main hypothesis can be found under <http://dx.doi.org/10.17605/OSF.IO/7RQVH> and <https://doi.org/10.17605/OSF.IO/8TZ57>. Ethical approval was obtained from the German Association for Experimental Economic Research e.V. and can be accessed under <https://gfew.de/ethik/Ft9eR5SK>.

bers. Volunteering is assumed to be costly. A single volunteer in the group is sufficient to produce a benefit for all its members, which no one receives in case no volunteer can be found. Furthermore, the individual benefit from the public good is greater than the costs of volunteering. This gives rise to the dilemma situation of the game: If there were another volunteer in the game with certainty, workers would never volunteer. However, given that the benefit is greater than the costs, workers would prefer to volunteer if all of their colleagues defected. It captures three essential characteristics of volunteering decisions in the workplace. First, volunteering is chosen simultaneously and without communication, which sets a lower bound for our purposes. Second, there is heterogeneity in, and incomplete information about, the costs of volunteering; and third, there might be variations in the group size. Its numerous variations have been examined in several experimental studies, and the main predictions, particularly the diffusion of responsibility effect, are reasonably robust (Diekmann, 1986; Franzen, 1995; Goeree et al., 2017; Kopányi-Peuker, 2019). Also, in different field environments, the predictions from the Volunteer’s Dilemma turn out to be robust (Latané and Nida, 1981; Przepiorka and Berger, 2016; Barron and Yechiam, 2002). Given the vast empirical support for the model (see, e.g., Latané and Nida, 1981; Przepiorka and Berger, 2016; Barron and Yechiam, 2002), one might expect it to be useful in providing predictions in our setting of an online workplace as well.

We use the volunteer dilemma framework to guide our experimental design of the field experiment. In order to establish causal claims, we must maintain a high degree of experimental control. Online labor markets are, therefore, a convenient and particularly useful environment for our experiment. Importantly, it is a regular and natural workplace for our experimental workers, and workers differ in their effort costs. At the same time, it allows us to set group sizes exogenously. This further differentiates online labor markets from classical work environments, where workers are usually connected through personal relationships and a common history across and within teams. These factors would impede the identification of the causal effects of group sizes, making the use of an online labor market crucial for this study.

In a nutshell, the field experiment consists of two stages (see Figure 3.1). In the first part of the job, workers were invited to work individually on a coding task for a fixed payment. Upon joining the job, the workers were randomly matched to one of five treatments in which we varied the group size and the incentive structure. After completing

the first individual task, workers were informed about the second stage. We asked them whether they would like to volunteer for a second coding round, just like the one they had done before, but with a different payment scheme. This second stage implements the Volunteer’s Dilemma: Only if at least one worker in the group task volunteered were all group members paid an additional bonus. Finally, we elicited the workers’ beliefs about their team members’ volunteering decision, and they had to answer a short questionnaire. All workers received the payoffs some days after the experiment.

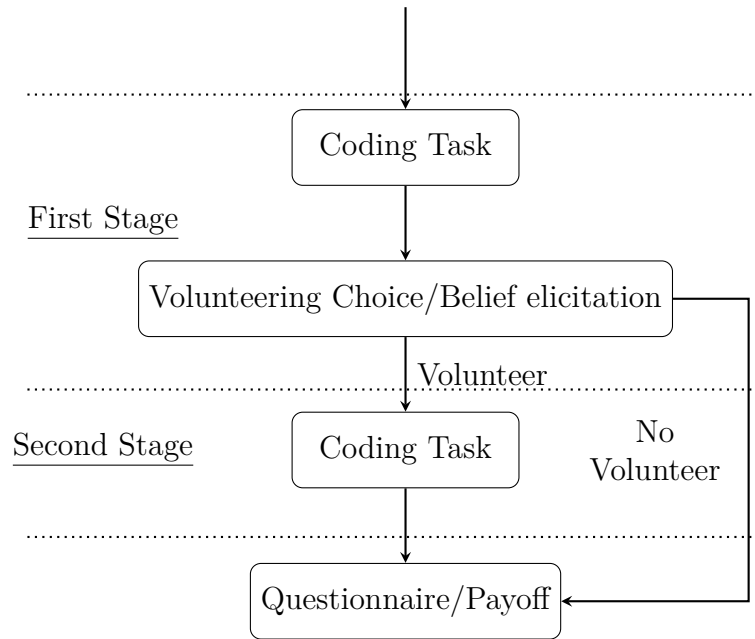


Figure 3.1: Structure of the experiment.

3.2.1 Workers and the online labor market

The field experiment was conducted on *clickworker.de*, an online crowdsourcing marketplace. Crowdsourcing marketplaces allow people to work on tasks that are usually easy for humans but difficult to automate. Most tasks on such platforms require a couple of minutes to complete and include assignments like the processing of images or the cleaning of data (see Difallah et al. (2015) and Jain et al. (2017) for an overview of everyday tasks). Online labor markets have become increasingly popular in recent years (Difallah et al., 2015), with 0.5% of the US adult population working in the “sharing economy” in 2016 (Farrell and Greig, 2017). For many workers, these jobs are a substitute for traditional offline work in times of economic downturn (Borchert et al., 2022). For them, online

labor markets are a regular work environment, which makes it a perfect testbed to study volunteering at the workplace.

Each worker was only allowed to participate once and required to speak German, but we did not impose any further restrictions on the pool of workers. Every active worker on the platform fulfilling those requirements was free to join. In total, 3,344 workers joined the assignment and read the explanation of the task. The assignment was made unavailable once each main treatment had reached 300 workers who had finished the first stage.³ Overall, 2,203 workers finished the first stage of the field experiment and decided whether they would like to continue working on the task, thus volunteering within their group.⁴ Altogether, 2,142 workers reached the end of the experiment. For the analysis, we consider only those workers who reached the end of the study. We remain with the treatment composition shown in Table 3.1. This subset of observations will be used for all subsequent analysis in Section 3.4, unless stated differently. We obtain a diverse sample with a wide range in age, gender, educational status, and employment (also see Table 3.8). Of all participants, 26% report an age between 18 and 25 years, but a sizeable share (14%) is also above 45 years of age. Further, the sample comprises self-employed, employed, and unemployed people with various educational backgrounds.

Table 3.1: Number of observations for each treatment.

Group Size	Number of Observations
UNINCENTIVIZED ($N = 1$)	192
INCENTIVIZED ($N = 1$)	196
SMALL GROUP ($N = 3$)	588
MEDIUM GROUP ($N = 30$)	582
BIG GROUP ($N = 300$)	584

Note: The number of workers for each treatment, with N being the group size.

3.2.2 The advertised job

We offered a standard job to all workers active at that time, via an advertisement on the platform. Importantly, there was no mention of an experiment or similar. The workers

³For the baseline treatments, the task was made unavailable after 200 workers had made a volunteering decision.

⁴Some treatments were filled with slightly more than 300 workers since workers could join the task simultaneously.

were invited to rate user comments from another study, a job that can be frequently found in online labor markets. We provided a short description of the task and informed the workers that they could earn 0.90 €, the standard wage set by the platform, and how long it would approximately take them to complete the task.

Once they had clicked on our link on the platform, they were redirected to our oTree server (Chen et al., 2016a). The workers then received detailed instructions about the task. Crucially, they were made aware that they would be working in a team for this assignment. We clarified that they would first be working alone on the described assignment, and would then be offered a voluntary group project afterwards (see the Appendix for the full instructions).

The Task Workers were asked to evaluate user comments from an online forum. Each comment was made with reference to a picture and possibly comments from other users in this online forum. The picture’s theme was always related to migration, refugees, or cultural differences. The workers rated the comments regarding the expressed sentiment and evaluated whether the comments contained hate speech. The coding scheme and a screenshot of the task can be found in Appendix 3.D.2. In both possible experiment stages, 30 distinct comments had to be rated, randomly drawn from a set of 13.356 comments.⁵ All comments had been collected as auxiliary data in a different study and are not part of our research question. Companies and corporations frequently use online labor markets for similar tasks to better understand customer comments or reviews. According to Difallah et al. (2015), those verification and validation tasks belonged to the most common assignments between 2009 and 2014 on Mturk, a US-based competitor of *clickworker.de*. According to *clickworker.de*’s own information, verification tasks and sentiment analyses are the most common tasks in their industry. Thus, we argue that most workers were familiar with this type of task and perceived it as a regular assignment rather than as part of a research project. Furthermore, the ratings of the comments will be used in the study for which they had been collected in the first place. Thus, the work was indeed meaningful and important.

⁵The rating of the comments was used as a dependent variable in Álvarez-Benjumea and Winter (2018), Álvarez-Benjumea and Winter (2020), and Álvarez-Benjumea (2020).

3.2.3 The volunteering decision

After finishing the first stage of the experiment, we explained to the workers that we needed precisely one volunteer in their group to ensure data quality and better evaluate the quality of the ratings within their group. We clearly stated that one volunteer within the group was sufficient for the task. Each worker received the offer to volunteer in their team and to continue to work on the task in the second stage. If at least one group member volunteered, all members received a bonus payment of 0.90 €. If no person in the group volunteered, no group member received any bonus. The bonus payment did not increase if more than one person volunteered. Furthermore, we explained to the workers that, even if they had not volunteered themselves, they might still receive the bonus payment if one of their teammates volunteered. Importantly, workers were not made aware if others already did the task before making their own choice. This allows us to measure the actual willingness to volunteer. To avoid reputation effects, we clarified that their decision would not have an influence on the payoff of the first stage or their user rating in the online labor market. The volunteering decision will be our key dependent variable in the analysis in Section 3.4. A screenshot and a translation of the decision screen can be found in Appendix 3.D.3.

We used the time spent in the first part of the job as a cost measure for volunteering. Time spent on a job is a fair measure of opportunity costs because those who spend more time on the task miss out on more opportunities to work, e.g., on another job on the platform or on more leisure. As we show later, workers differ substantially in the time it takes them to complete the task. Since both parts of the job consist of the same task, the time spent on the first part is also a good predictor for time spent on the second part ($\rho = 0.73$, $p < .001$, Spearman's rank correlation coefficient). Workers were made aware of the time it had taken them to complete the first stage.

Following our theoretical model, we induced commonly known beliefs about the distribution of costs of other workers before making the volunteering decision. To this end, we informed them that other workers usually required between 7.5 to 15 minutes to complete

the rating of the comments.⁶ This gave workers a rough estimate of whether they had high or low costs of volunteering relative to the other workers.

3.2.4 Treatments

Volunteer’s Dilemma Treatments Workers in our main treatments faced team incentives in the form of a Volunteer’s Dilemma. Within the field experiment, we varied three group sizes, the SMALL GROUP with 3 workers, the MEDIUM GROUP with 30 workers, and the BIG GROUP with 300 workers.⁷

Instructions were adapted accordingly for each group size. Since the information about the group size is the key point in our study, we made sure that the information was clear to workers. The group size was mentioned several times, particularly right before workers made a choice.

Baseline Treatments (N=1) In order to provide a benchmark for volunteering rates in our setting, we designed two control treatments: INCENTIVIZED and UNINCENTIVIZED. The basic setup in this study was identical to our main experiment, the difference being that the workers did not face a Volunteer’s Dilemma. That is, the instructions were the same as in the Volunteer’s Dilemma Treatments, including the fact that workers operated in a team. We pointed out, however, that their actions did not influence the payoffs of their team members, and vice versa.

Workers were notified at the end of the first stage that we needed a volunteer to continue to work on the task.⁸ We varied two conditions. In the INCENTIVIZED condition, subjects were paid the same bonus as in the Volunteer’s Dilemma treatments (90 cents) for completing the second part. Yet, if they did not volunteer, they did not receive the bonus even if another team member volunteered. This condition would provide an upper bound

⁶These numbers were collected in a pilot study with 100 workers. We showed the workers the 20th and 80th percentile of the time values, but referred to these values as the time it took “most workers”. This is obviously a deviation from our theoretical model, which assumes common knowledge of the entire distribution. We made this simplification not to overburden our workers and maintain the atmosphere of a typical job.

⁷To implement a truly random and independent draw of the group size, there was an individual random draw for each worker. Then, for each worker, a random team of a specific size was generated, and the worker’s payments were based on the actions of these workers. This method is similar to the method used by Boosey et al. (2017).

⁸Note that in the main treatments, we explicitly told the workers that we only required a single person to volunteer. In this baseline case, we did not specify how many volunteers are required, but rather that we needed volunteers. In fact, the INCENTIVIZED condition can be understood as a Volunteer’s dilemma with a single worker.

of volunteering without any team incentives. In the UNINCENTIVIZED condition, there was no bonus, which allows us to control for intrinsic or non-monetary motivations to finish the task. The volunteering rates of our main treatment, where strategic considerations play a role, should lie between these two conditions.

3.2.5 Belief elicitation

We elicited the participants' beliefs about the volunteering decision of other group members. We asked each participant how likely it is for at least one of their team members to volunteer for the task. Participants reported the probability on a scale from 0 % ("For sure no other person") to 100% ("Surely at least one other person") using a slider. The slider did not have an initial value to avoid any anchoring effect.⁹ The belief elicitation was incentivized using the binarized scoring rule, and the possible bonus is 0.90 € (Hossain and Okui, 2013).¹⁰ The instructions were simple and natural to increase truth-telling (Danz et al., 2022). We randomized the order in which we asked participants for their volunteering decision and their belief. Half of the participants were asked about their beliefs just before their volunteering decision. The other half answered the belief question after their volunteering decision but before the potential second round of coding comments. It allows us to investigate possible order effects. As the order in which we ask participants for their volunteering decision and their belief plays a minor role in the results, we pool all observations along this treatment variation for the analysis in Section 3.4.4.¹¹

After the participants reported their beliefs, we also elicited how certain they were about this probability. To measure cognitive noise, we followed the approach by Enke and Graeber (2019, 2021).

3.2.6 Questionnaire

At the end of the field experiment, workers completed a questionnaire regarding their economic preferences (Falk et al., 2022) and sociodemographic/economic background. Furthermore, we asked the participants different questions about the task they had to perform. We use the responses as an additional cost measure. We kept initial questions

⁹A screenshot is provided in Appendix 3.D.

¹⁰The winning probability was calculated as $p = 1 - (V - \frac{\text{Belief}}{100})^2$ where V is a dummy variable that is equal to one if there was another volunteer in the team.

¹¹See Appendix 3.C.1 for an analysis of the order effects.

to a minimum to avoid disturbing the impression that the experiment was anything other than a typical job.

3.2.7 Replication

We replicate the main treatments of an earlier study with 2,733 participants.¹² The results of volunteering in this earlier study are virtually identical to the results here, giving us even more confidence in the robustness of our findings. The replication allowed us to investigate different channels like beliefs, and we also improved the instructions to avoid any misunderstandings of the rules by the participants.

3.3 Hypotheses

The incentive structure in our field experiment resembles a Volunteer's Dilemma (Diekmann, 1985) with incomplete information and heterogeneous costs. For our model, we thus focus on the setup by Weesie (1994). We introduce the model's key features in the following to derive our hypothesis. The key insights from the model are that volunteering decreases in the group size and in the costs for volunteering. Further details are provided in Appendix 3.A.¹³

To be more formal, there are N players in the game.¹⁴ Each player $i \in \{2, \dots, N\}$ decides simultaneously to volunteer ($a_i = V$) or to defect ($a_i = D$). If at least one player in the game volunteers, all players receive a benefit of b_i . Volunteering is costly, and the costs are denoted by c_i . In line with Weesie (1994), we assume that $\gamma_i := \frac{c_i}{b_i}$ follows some arbitrary probability distribution $\gamma \sim \mathcal{F}$ with a continuous probability density function

¹²The earlier experiment is reported in a previous version of the paper which this chapter builds on and can be accessed here <https://osf.io/4k5y6>. The main results of the original study can also be found in Figure 3.6. In the original study, we also ran treatments where the group size was uncertain (compare Hillenbrand and Winter, 2018), which we do not further explore in this chapter.

¹³While this might seem intuitive for most readers, it is worth mentioning that this prediction reverses the predictions under common knowledge of costs, where those with higher costs volunteer more in equilibrium.

¹⁴The game can be easily extended to allow for stochastic group sizes. The details are also provided in the respective section of the appendix.

f and that $b_i > c_i > 0 \forall i$. The payoff π_i of worker i when X_{-i} others volunteer is

$$\pi_i = \begin{cases} 0 & \text{if } a_i = D \text{ and } X_{-i} = 0 \\ b_i & \text{if } a_i = D \text{ and } X_{-i} > 0 \\ b_i - c_i & \text{if } a_i = V. \end{cases}$$

The payoff structure gives rise to the dilemma situation of the game: If there were another volunteer in the game for sure, workers would never volunteer. However, given that the benefit is greater than the costs, workers would prefer to volunteer if all of their colleagues defected. Formally, player i volunteers if and only if its expected benefit is weakly greater than the benefit from defecting.

$$\begin{aligned} EU(V) &\geq EU(D) \\ \Leftrightarrow b_i - c_i &\geq b_i P(X_{-i} > 0) \\ \Leftrightarrow \gamma_i &\leq 1 - P(X_{-i} > 0) \end{aligned} \tag{3.1}$$

Thus, she volunteers if the cost-benefit ratio $\gamma_i = \frac{c_i}{b_i}$ of player i is weakly below the probability, that there is no other volunteer. Otherwise, she defects. Weesie (1994) shows that there exists a pure strategy equilibrium where players with low cost-benefit ratios volunteer, while those with a γ_i above some threshold do not volunteer.

In our field experiment, we can elicit the main parameters of the model. The additional bonus of 90 cents is the benefit b_i . Volunteering for the task costs time and effort. We argue that those costs to volunteer (c_i) can be approximated by the time it takes participants to complete the first part of the experiment and additional survey measures that we elicit in a questionnaire. Also, we directly ask participants for their belief about $P(X_{-i} > 0)$, i.e., the probability that there is at least one other team member volunteer. As a result, our experimental design allows us to test a wide range of hypotheses, which are derived from the equilibrium analysis by Weesie (1994). We provide further details on the theoretical foundations in Appendix 3.A.

In our baseline treatments, volunteering is an individual choice. It allows us to identify if volunteering is indeed costly and participants react to the dilemma component of the game. If participants perceive the task as costly, they should volunteer less if there is no additional benefit. Thus, we expect volunteering rates in INCENTIVIZED to be

higher than in the UNINCENTIVIZED treatment. The group size treatments mimic a volunteer's dilemma, which introduces freeriding opportunities. Only a single person from the team has to volunteer such that everyone receives an additional bonus. Accordingly, we expect that the volunteering rate is lower in the main group size treatments compared to the INCENTIVIZED treatment, in which participants only influence their own benefit if they volunteer. Yet, as we expect that volunteering is perceived as work and costly, we hypothesize that more participants volunteer in the main treatments compared to the UNINCENTIVIZED treatment. This argument is summarized by Hypothesis 1.¹⁵

Hypothesis 1. *The volunteering rate in the Volunteer's Dilemma treatments is higher than in the UNINCENTIVIZED treatment but lower than in the INCENTIVIZED treatment.*

Participants with lower costs should volunteer more for the additional task in the main treatments. That follows intuitively from Equation 3.1 as the expected utility from volunteering is decreasing in the costs. A formal proof is provided in Appendix 3.A.

Hypothesis 2. *Workers with lower costs volunteer more.*

From a theoretical perspective, the share of volunteers decreases in the group size. Large groups give rise to a higher “diffusion of responsibility”, and players volunteer less on an individual level. We expect that the effect carries over to our workplace environment.

Hypothesis 3. *The volunteering rate decreases in the group size.*

We elicit the participants' beliefs that another person in the team volunteers. We can verify whether participants have correct beliefs using the average volunteering rate in each treatment. Furthermore, the expected value from defecting is increasing in this belief (see Equation 3.1). For higher beliefs, there exists a higher probability that a participant receives the benefit without volunteering herself. Participants with higher beliefs should volunteer less, as they can avoid incurring the costs of volunteering while receiving the same benefit.

Hypothesis 4. *Participants who believe that there is another volunteer with a higher likelihood volunteer less themselves.*

¹⁵The pre-registered hypotheses can be found under <http://dx.doi.org/10.17605/OSF.IO/7RQVH> and <http://dx.doi.org/10.17605/OSF.IO/8TZ57>. The wording of the hypotheses was adapted, and we added hypotheses about beliefs based on theoretical results to fit the structure of the chapter.

Clearly, all these hypotheses build on a purely monetary-driven argument. In a work environment, multiple additional factors such as peer effects or image concerns might play a role. These factors might influence the link between beliefs and actions in different ways.

3.4 Results

In this section, we present the field experiment results, following the hypotheses outlined in Section 3.3. We next consider the stated beliefs of subjects about their pivotality to provide the good and the interplay between beliefs and actions. Furthermore, we provide further robustness tests for our results in Appendix 3.C. The data on the main treatments that we present in this section is based on the replication of our earlier experiments as described in Section 3.2.7.

3.4.1 Workers are sensitive to incentives and strategic situations

To satisfy our first identifying assumption, we must establish whether the volunteering choice was costly. Therefore, we compared two baseline treatments in which we asked participants whether they wanted to volunteer in a second coding round. In these baseline treatments, the payment only depends on the worker's decision and not on other workers.¹⁶ In the INCENTIVIZED condition, participants were paid an additional 0.90 € in case they volunteered for a second round of coding; in the UNINCENTIVIZED condition, they were only paid for the first round.¹⁷ The volunteering rate in the INCENTIVIZED version was 68.4%, but only 27.1% in the UNINCENTIVIZED treatment (See Figure 3.2, $p < 0.01$, χ^2 -test). This allows us to conclude that the task was perceived as a costly effort and establishes the base for the coming results.

Our second identifying assumption was that our participants react to the strategic situation of the Volunteer's Dilemma and show different volunteering rates than in the individual choice situation. We expected that volunteering rates in the main treatments fall between those in the incentivized and the unincentivized individual choice condition. This is also the case. Pooling the data for all main treatments, the volunteering rate is

¹⁶Note that we kept the framing of them being part of the team to keep the decision environment stable.

¹⁷As explained in Section 3.2 for the main treatments, we only consider those participants who had reached the end of the experiment. Out of 400 participants who started the experiment, we ended up with 388 observations: 196 in the INCENTIVIZED condition and 192 in the UNINCENTIVIZED condition.

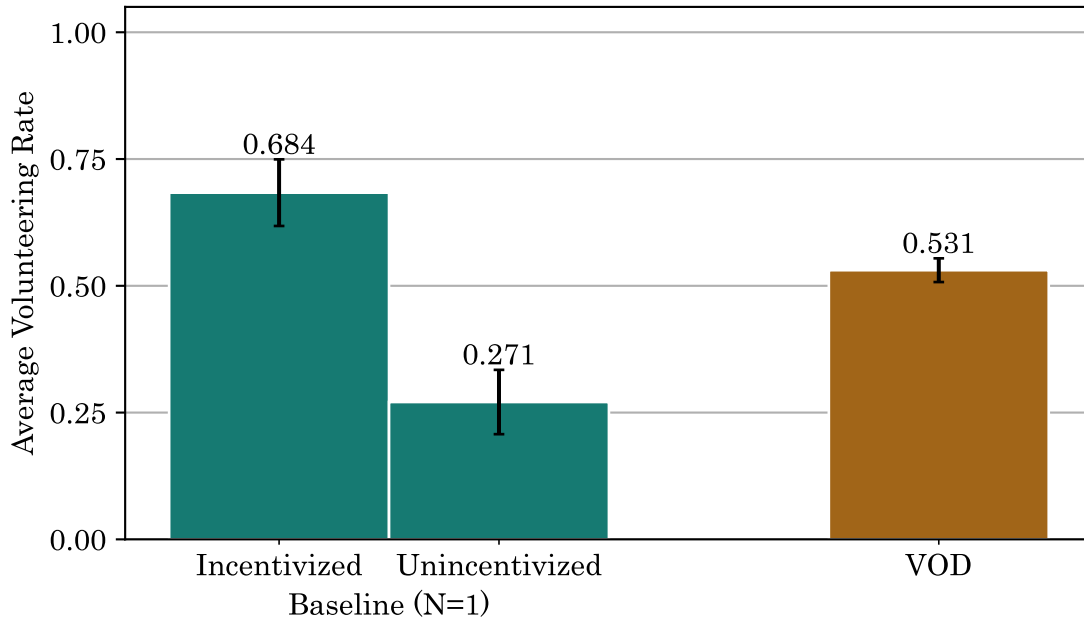


Figure 3.2: Volunteering rates in the incentivized and the unincentivized baseline treatments in comparison to the main Volunteer’s Dilemma (VOD) treatments. The error bars represent the 95% confidence intervals.

significantly lower, at 53.1%, than in the INCENTIVIZED condition ($p < 0.01$, χ^2 -test) and significantly higher than in the UNINCENTIVIZED individual choice condition ($p < 0.01$, χ^2 -test, see Figure 3.2). Hence, we find clear support for Hypothesis 1.

3.4.2 Heterogeneity in costs to volunteer

Hypothesis 2 predicts that participants with higher costs are less likely to volunteer. As explained in Section 3.2, we argue that individual costs of volunteering can be approximated by the time it takes a participant to finish the first stage of the field experiment. Most participants required between 8.23 (20th percentile) and 16.52 minutes (80th percentile) to complete the first stage.¹⁸

Our Hypothesis 2 is not supported by the data since the effect of costs operationalized as time is not significant. Model 1 in Table 3.2 shows the results of a logistic regression

¹⁸Note that those values are similar to the ones attained in the pilot study, which were shown to the participants as an approximation of the cost distribution, e.g., 7.5 (20th percentile) and 15 (80th percentile). The data includes extreme outliers, with some participants having a completion time of more than 30 hours or as little as 1.50 minutes. The participants who took more than 16 hours did not constantly work on the assignment but took long breaks. Figure 3.7(a) presents the estimated distribution of completion times.

model and estimates the probability of volunteering as a function of the z-standardized time. Importantly, the coefficient of TIME is insignificant.¹⁹

We constructed a z-standardized additive index from several control questions in our post-experimental questionnaire as an additional measure of subjective costs. We asked participants on a 7-point Likert scale whether they perceived the task as exhausting ($\mu=3.01$), interesting ($\mu=2.91$, reverse-coded), or emotionally challenging ($\mu=2.30$), and whether it was important to them to contribute to “better data quality” ($\mu=2.44$, reverse-coded).²⁰

The subjective costs have a substantially meaningful and highly statistically significant negative effect on the probability of volunteering (see Model 2 and 3 in Table 3.2 and Table 3.7). We, therefore, conclude the following:

Result 1. *The probability of volunteering decreases in the subjective costs of volunteering.*

Table 3.2: Logistic regression to estimate the volunteering choice as a function of different cost measures

	<i>Dependent variable:</i>		
	Volunteering Choice		
	(1)	(2)	(3)
TIME	−0.007 (0.048)		−0.005 (0.050)
COST INDEX		−0.375*** (0.050)	−0.375*** (0.050)
Constant	0.123*** (0.048)	0.126*** (0.049)	0.126*** (0.049)
Observations	1,754	1,754	1,754
Log Likelihood	−1,212.443	−1,183.000	−1,182.994
Akaike Inf. Crit.	2,428.886	2,370.001	2,371.989

Note: *p<0.1; **p<0.05; ***p<0.01
Standard errors in parentheses

¹⁹This is also the case for the estimated average marginal effect, which can be found in Table 3.7 in the Appendix.

²⁰All questions have been answered on a Likert scale between 1 ("Strongly disagree") and 7 ("Strongly agree").

3.4.3 Insensitivity to team size in the volunteers dilemma

Hypothesis 3 predicts that the volunteering rate decreases in the group size. This is clearly not the case (see Figure 3.3). Volunteering rates are relatively high and statistically indistinguishable with regard to group size (all $p > 0.1$, χ^2 -tests).²¹ We, therefore, conclude that

Result 2. *The volunteering rate does not vary in the group size.*

This result is surprising and interesting concerning our hypothesis and in light of the vast literature on the volunteer's dilemma. It constitutes our main result. In the remainder of the chapter, we will provide a potential explanation based on conditional volunteering and show that the finding is exceptionally robust.

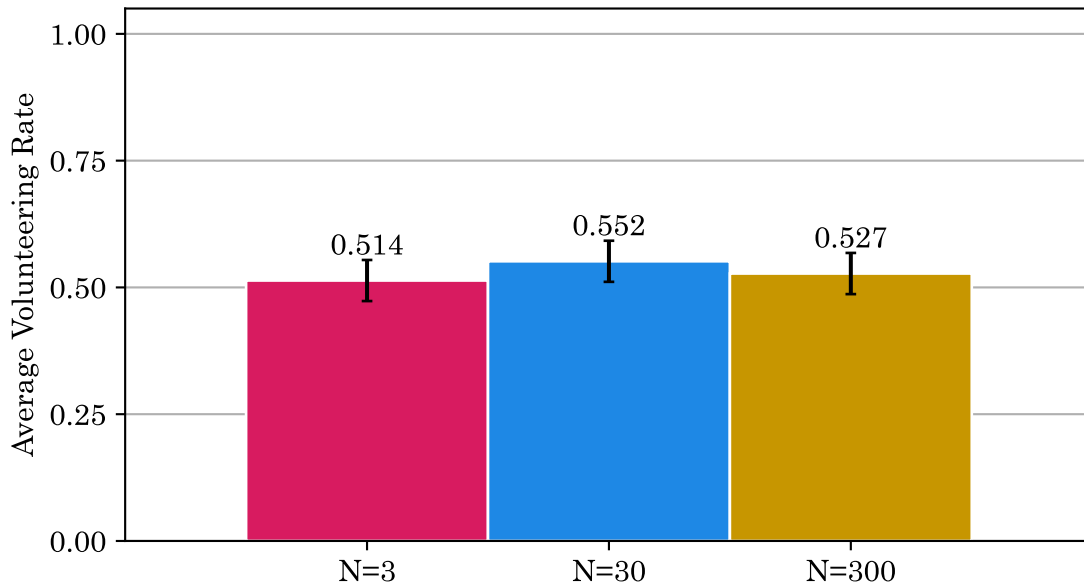


Figure 3.3: The average volunteering rate across the different group sizes. The error bars represent 95% confidence intervals.

3.4.4 Beliefs

In this section, we report the beliefs of the participants about the volunteering decision of their team members. We compare the beliefs with the actual probability that at least one other person volunteers for the given group size.²²

²¹We also test for difference in volunteering across group sizes using a logistic regression in Table 3.5. The average marginal effect estimates in Table 3.6 show that the effect of group size is small and not statistically significant at any conventional level.

²²The probability that one other person volunteers is given by $P(X_{-i} > 0) = 1 - (1 - s)^{N-1}$ where s denotes the share of volunteers in a given treatment.

Table 3.3 reports the results. In each treatment, participants report beliefs that are, on average smaller than the correct belief, i.e., they overestimate their pivotality. The differences are statistically significant ($p < 0.01$ for all treatments).²³ Hence, participants do not form correct beliefs about the volunteering decisions of other team members.

Table 3.3: Belief overview by treatment

Group size	Correct belief	Average belief	Share with correct* belief
SMALL GROUP (N=3)	76.34	60.66	0.11
MEDIUM GROUP (N=30)	100.00	70.24	0.30
BIG GROUP (N=300)	100.00	70.89	0.30

* The correct belief is calculated based on the average volunteering rate in the respective treatment. Here, we define the belief of a participant as correct if it is in the interval of ± 5 pp around the correct belief.

Figure 3.4 shows the discrete distribution of the beliefs by group size. There exists a considerable variation across beliefs. While about a third of participants in $N = 30$ and $N = 300$ reports beliefs close to the correct belief, beliefs are, on average, too low. In $N = 3$, even fewer participants report a belief close to the actual probability that one other team member volunteers. This is also driven by a large share of participants who believe that no one else volunteers.

Result 3. *Beliefs about volunteering are, on average, incorrect.*

Next, we consider differences in beliefs across treatments. Participants in the $N = 3$ treatment report statistically significantly smaller beliefs than in the other two treatments ($p < 0.01$ for both comparisons, two-sided Mann–Whitney U test). There are no differences in beliefs between $N = 30$ and $N = 300$ ($p = 0.64$, two-sided Mann–Whitney U test).

Result 4. *Beliefs are lower in small teams compared to larger team sizes.*

The differences in beliefs make intuitive sense, given the observed volunteering rates in the field experiment. As volunteering rates are similar, the probability that another person in the team is a volunteer increases in the group size.²⁴

²³We regress the difference between the correct belief and the reported belief in a Tobit regression on a constant.

²⁴Furthermore, we observe a higher degree of cognitive uncertainty about the belief in the SMALL GROUP compared to the other two group sizes ($p < 0.01$, two-sided MWU tests). There is no difference between the MEDIUM GROUP and the BIG GROUP ($p > 0.1$, two-sided MWU test). Following the argument by Enke and Graeber (2019), it indicates that subjects shrink their beliefs more towards an ignorant prior. While this ignorant prior is unobserved in our experimental design, it is reasonable to assume that it is

From a participant's perspective, one would expect that lower beliefs lead to higher volunteering and, thus, to see more volunteers in small groups. However, participants volunteer at similar rates across treatments. A possible explanation is that a non-trivial link exists between actions and beliefs in the workplace environment, which is not captured by the classical Volunteer's Dilemma framework. We investigate the relationship between beliefs and actions in the following section.

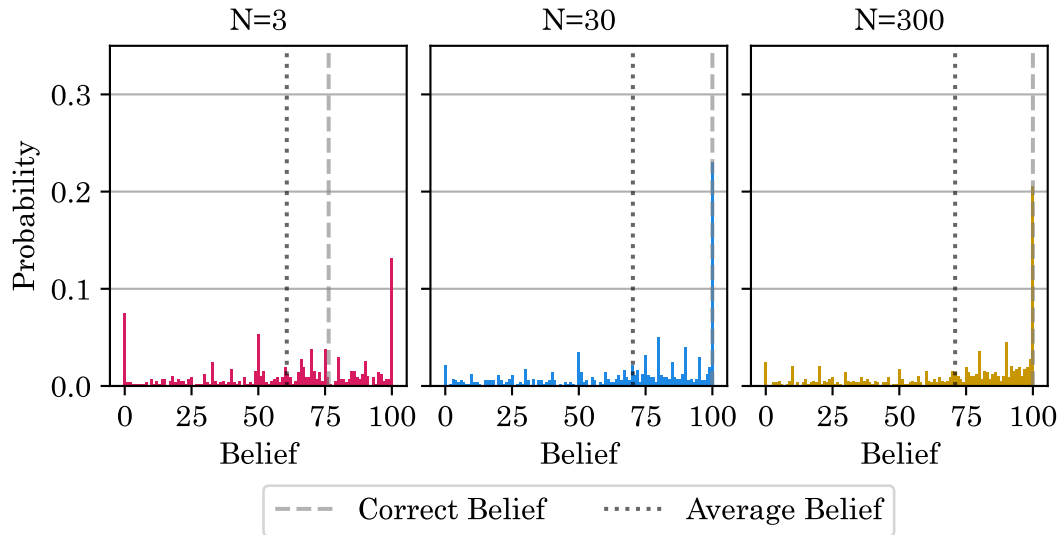


Figure 3.4: The belief that another person in the group volunteers by group size. To calculate the correct belief we use the average volunteering rate in the treatment.

3.4.5 Conditional volunteering

We now focus on how beliefs translate into actions. As discussed in Hypothesis 4, differences in beliefs should translate into differences in the action domain. The higher the probability that someone else in the team volunteers, the larger the likelihood that the player still receives the benefit without volunteering herself. From a purely monetary perspective, subjects who are sure that another participant volunteers (i.e., report a belief of 100%) should never volunteer themselves as volunteering is costly, and the benefit does not increase if there are multiple volunteers. On the other hand, participants with pessimistic beliefs about the volunteering of others should be more likely to volunteer.

We regress the volunteering decision on beliefs in a linear probability model to investigate

50% given the binary outcome. As a result, cognitive uncertainty can also explain the belief differences between the SMALL GROUP and the other group sizes. In line with Enke and Graeber (2019) we find a hump shape relationship between reported probabilities and cognitive uncertainty. In other words, participants who report beliefs close to 50% tend to encounter a higher degree of cognitive noise. The effect is especially strong in smaller groups. See Figure 3.8 for further details.

this relation in the data. Furthermore, Figure 3.5 visualizes those effects in a local linear regression.

Table 3.4: The volunteering choice explained by Beliefs and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>		
	Volunteering Choice		
	(1)	(2)	(3)
BELIEF	0.005*** (0.0004)	0.005*** (0.0004)	0.007*** (0.001)
MEDIUM GROUP		−0.009 (0.027)	0.090 (0.062)
BIG GROUP		−0.037 (0.028)	0.205*** (0.064)
MEDIUM GROUP × BELIEF			−0.002** (0.001)
BIG GROUP × BELIEF			−0.004*** (0.001)
Constant	0.203*** (0.026)	0.214*** (0.029)	0.108*** (0.040)
Observations	1,754	1,754	1,754
<i>Note:</i> *p<0.1; **p<0.05; ***p<0.01 Robust standard errors in parentheses.			

Table 3.4 shows the results, which are in sharp contrast to Hypothesis 4 and the theoretical predictions from the volunteer’s dilemma. There exists a strong positive correlation between beliefs and the volunteering decision. In other words, subjects volunteer *more* if they believe that it is more likely that there is another volunteer in their team. An increase in the belief by one percentage point increases the probability of volunteering by 0.5 percentage points (Model 1 in Table 3.4).

Clearly, this connection between beliefs and actions is not individually rational from a monetary perspective. However, it is well in line with results from the literature on (social or self-) image concerns (e.g., Bénabou and Tirole, 2006) or the vast literature on descriptive norms as well as on peer effects (see, e.g., Cornelissen et al., 2017, for a recent study on peer effects at the workplace). We interpret this effect as a normative driver

for the outcomes of our field experiments. Compared to a standard volunteers dilemma in the laboratory, it is reasonable to assume that subjects in the workplace environment want to comply with the decision of their co-workers, i.e., they do not want to be seen (or see themselves) as unsocial or lazy by acting differently than what they believe is the “correct” thing to do.

However, the degree of the conditional volunteering rate, i.e., the marginal increase of volunteering on beliefs, differs across treatments. In particular, in the small group, the average marginal reaction to an increase in beliefs by 1% is larger compared to the medium and the big group (0.7 vs. 0.5 vs. 0.4 percentage points, see Model 3 in Table 3.4).

Together with the finding that beliefs are, on average, lower in the small group compared to the two larger group sizes (Result 4), this can explain why we find no overall effect on volunteering. That is, while participants, on average, have lower beliefs about the volunteering of their group members, those with higher beliefs react more strongly to the beliefs compared to those with similar beliefs in the larger groups. A possible explanation is that participants fear sticking out more in smaller groups. This effect can also be observed when considering only subjects who believe that there will be another volunteer with certainty (a belief of 100). Among those participants, 82% volunteer in the SMALL GROUP, 73% in the MEDIUM GROUP and 62% in the BIG GROUP.²⁵

Result 5. *Workers are conditional volunteers and volunteer more if they believe others volunteer as well.*

3.4.6 Robustness of our results

Our results on volunteering behavior, as well as beliefs, are surprising. Thus, our first step was to replicate our results in Volunteer’s Dilemma treatments, and the results are presented above. The finding on volunteering rates is nearly identical to the first experiment.²⁶

²⁵The differences between the three groups are statistically significant when tested jointly ($p < 0.01$ χ^2 -test).

²⁶The results from the first experiments are provided in Figure 3.6 and the earlier version of the paper, which this chapter builds on <https://osf.io/4k5y6/>.

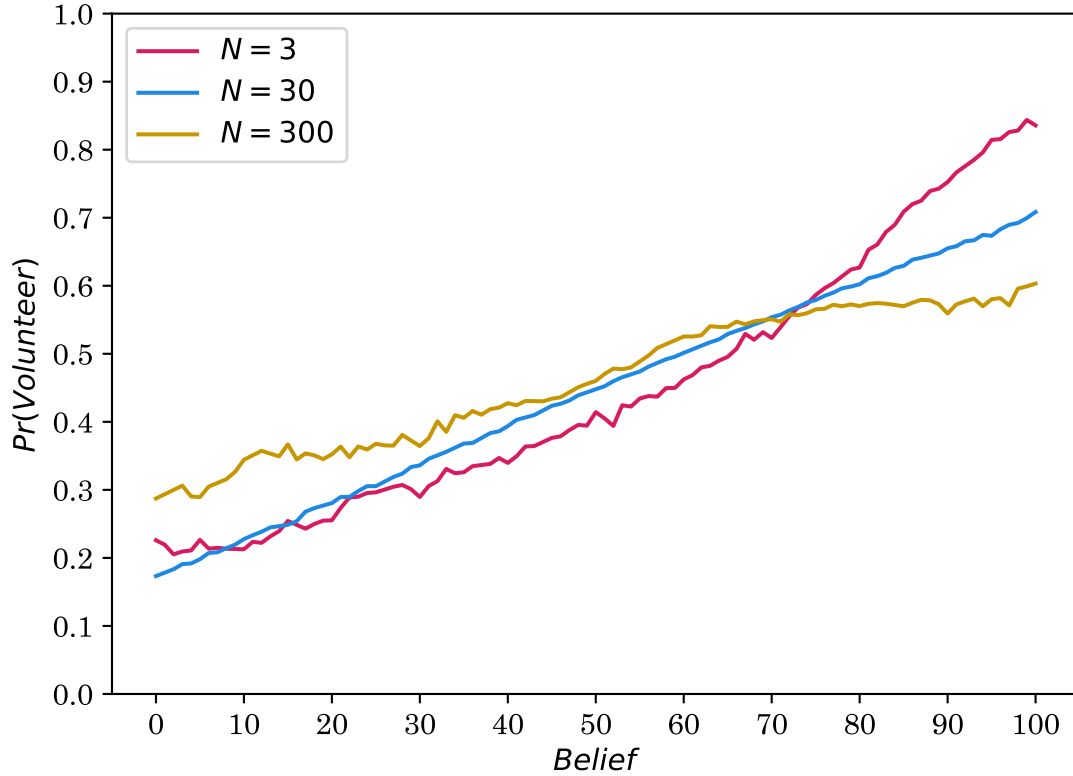


Figure 3.5: Fitted values from a local linear regression describing the probability to volunteer by the belief. We use the Wang Ryzin Kernel and leave-one-out cross-validated bandwidth.

Additionally, to ensure that our results are robust, we thoroughly analyze additional measures we obtained in the questionnaire. We find many factors influencing individual volunteering, which provides interesting insights into the drivers of volunteering. See Appendix 3.C for the complete analysis. Importantly, however, beliefs are by far the strongest predictor of volunteering compared to any other measure we elicit in the experiment (see Section 3.C.6 in the Appendix for a discussion). Also, other factors are balanced across treatments, while our treatment manipulation influences beliefs. This speaks to the robustness of our main result and the explanation of beliefs as a normative driver of behavior.

3.5 Discussion and concluding remarks

In this chapter, we study volunteering at the workplace. While volunteering as an allocation mechanism in work environments is now widespread, and the adaptation is championed by the industry, economic analysis and empirical results suggest that two factors

impede volunteering. First, the incentives create a coordination problem; second, workers prefer others to do the work, often leading to a situation similar to the Volunteer's Dilemma. In particular, our game-theoretical model and past research predict a higher diffusion of responsibility, i.e., lower volunteering rates, in larger groups.

We report the results of an online field experiment with more than 2,000 workers in an online labor market. In our experiment, our workers first individually worked on a standard classification task and were then asked to volunteer for a similar task to secure a bonus for all workers of their team. We exogenously vary the team size and thus study the causal effect of the team size on volunteering behavior. Additionally, we elicit workers' beliefs about the probability of volunteering of their co-workers.

We find that workers react to the strategic situation of the Volunteer's Dilemma and shifts in the incentive structure. Workers volunteer much more if they are alone instead of embedded in a group, and thus a Volunteer's Dilemma. They also volunteer much less if they are not paid for the bonus task. Furthermore, in line with our theoretical predictions, workers with lower subjective costs are more likely to volunteer. Nevertheless, in stark contrast to previous results that study volunteering outside of the remote work context, we find no effect of the group size on volunteering behavior. On average, about 53% of workers choose to volunteer across all our main treatments. We identify workers' beliefs about the volunteering decision of their team members as the main driver of this result. Workers who believe that it is more likely that at least one other worker volunteers are *more* likely to volunteer themselves. While, on average, members of small teams consider it less likely that someone else volunteers, the workers in those teams react to them more strongly compared to participants with similar beliefs in larger teams.

This *conditional volunteering* behavior is puzzling from a purely material perspective. However, it is in line with findings on descriptive norms. If workers perceive volunteering as a stronger descriptive norm, indicated by higher beliefs, they might be more inclined to follow these norms. A similar explanation can stem from image concerns. Workers do not want to be perceived as selfish by others (in this case, by the employer) or themselves for not volunteering if they believe that most other workers are volunteering. Note that while workers in the small group have lower beliefs than in the two other groups, the conditional volunteering effect is also higher here, making those workers with a high belief much more

likely to volunteer than the larger groups. Image concerns can again explain this more substantial volunteering effect: Workers do not want to ‘stick out’ in smaller groups.

Our results suggest that volunteering at the workplace is distinctively different from other volunteering situations. Workers like to conform to the prevalent behavior in their team. It appears that this effect overshadows the classical strategic considerations one would expect and which we discuss in Section 3.3.

Those findings have several implications for firms’ organizational structure or, more generally, social groups that rely on volunteering. We deliberately set up the work environment so that workers are not informed whether others worked before them, which allows us to measure their individual willingness to volunteer. However, depending on the exact organizational structure, our results might be good or bad news for volunteering at the workplace. Interpreted strictly within the context of the volunteer’s dilemma, the high level of volunteering is wasteful and inefficient. Especially in work settings where it is not possible for a manager to prevent this over-provision, the allocation mechanism has adverse welfare effects. Indeed, these inefficiencies are discussed in open source software development (McConnell, 1999; Kenwood, 2001). It is especially relevant in workplaces with limited communication between workers or with little oversight by a manager, like in remote work contexts that have become more popular in recent years. Here, managers should focus on smaller team sizes to prevent too much volunteering and reduce inefficiencies.

For organizations that cannot prevent over-provision and only care about finding at least one volunteer, i.e., who only care about producing the good, our results are good news. Contrary to what theory or intuition would suggest, giving workers the freedom to self-organize does not hurt commitment. Using larger team sizes might be desirable from the perspective of a manager. It is true even without oversight and without any disciplinary measures. Our findings might thus relate to high volunteering rates in other contexts, such as writing open-source code or Wikipedia articles, even though the numbers of potential volunteers are sometimes vast.

Appendix

Appendix 3.A A formal model of volunteering at the workplace

Here we lay out a general model of volunteering under cost uncertainty and apply it to our context of volunteering at the workplace. The model captures three essential characteristics of volunteering decisions in the workplace. First, volunteering is chosen simultaneously and without communication, which sets a lower bound for our purposes. Second, there is heterogeneity in, and incomplete information about, the costs of volunteering; and third, the group size might differ. Following Hillenbrand and Winter (2018), our model also allows for population uncertainty, e.g., the case when the exact group size is unknown, but this is not the scope of this contribution.²⁷

Arguably, the costs of volunteering are usually not homogeneous for all workers, and only in rare cases are the costs known precisely. We, therefore, include ideas from the volunteering models with heterogeneous effort costs (Diekmann, 1986), as well as incomplete information about the distribution of costs (see Weesie (1994) for a formal model and Healy and Pate (2018) for recent experiment evidence).

Player i can volunteer ($a_i = V$) or defect ($a_i = D$) with individual-specific benefit b_i and costs c_i . A single volunteer in the group is sufficient to produce a benefit for all its members, which no one receives in case no volunteer can be found. Furthermore, the individual benefit from the public good is greater than the costs of volunteering. The

²⁷We considered population uncertainty as an experimental variation in an earlier version of the paper and kept the theoretical considerations here for completeness. See Figure 3.6 for the experimental results on population uncertainty. Further details are provided in an earlier version of the paper that can be accessed here <https://osf.io/4k5y6/>.

payoff π_i of worker i is given by

$$\pi_i = \begin{cases} 0 & \text{if } a_i = D \text{ and } X_{-i} = 0 \\ b & \text{if } a_i = D \text{ and } X_{-i} > 0 \\ b - c & \text{if } a_i = V. \end{cases}$$

Here, X_{-i} denotes the number of volunteers in the game that are not player i . In line with Weesie (1994), we assume that $\gamma_i := \frac{c_i}{b_i}$ follows some arbitrary probability distribution $\gamma \sim \mathcal{F}$ with a continuous probability density function f . The following assumption is made on the support of the distribution

$$\text{supp}(\mathcal{F}) = [\underline{\omega}, \bar{\omega}] \subseteq [0, 1] \quad (3.2)$$

with $\underline{\omega} < \bar{\omega}$ and $[0, 1]$ being the unit interval.²⁸ We will commonly refer to γ_i as the type of player i .

Following the approach of Hillenbrand and Winter (2018), we let the number of workers n be drawn from a discrete probability distribution h . The probability mass function is denoted by $h(\cdot)$. Furthermore, let $n \in \tilde{N} = \{2, \dots, \bar{N}\}$, with $\bar{N} \in \mathbb{N}$ being the largest possible number of workers in the game. Since we are interested in the effect of the mean group size on volunteering, we define $E[n] = N$. Importantly, a fixed group size $\hat{n}$, i.e., a situation without population uncertainty, is then just a special case in this setup with a degenerate distribution $h(\hat{n}) = 1$. This captures our main treatment variation in the group size.

For the theoretical discussion instead of discussing the full set of potential equilibria, we focus on a pure-strategy equilibrium as in Weesie (1994). Importantly, as he points out, there is no equilibrium in mixed strategies, i.e., where a player of a given type plays V or D with a positive probability. For a more general discussion on possible equilibria in the Volunteer's Dilemma, see Diekmann (1985, 1986) and Weesie (1994). The pure-strategy equilibrium that we discuss has some nice properties and makes intuitive sense for an applied setting such as ours.

We show that in a Volunteer's Dilemma with incomplete information about costs and population uncertainty there exists an equilibrium where players with cost-benefit ratios

²⁸We further assume that the cumulative distribution function (cdf) denoted by $F(\cdot)$ is atomless.

below some threshold $\tilde{\gamma}$ volunteer, while those above defect. We further show that this threshold, and thus the individual probability to volunteer, is smaller in situations with a large probability of being in bigger groups.

Proposition 1 (Pure Strategy Nash Equilibrium). *Let $\mathcal{F}$ be an arbitrary probability distribution over γ with a continuous pdf and an atomless cdf. Let h be a discrete probability distribution over n . Then*

1.1 *There exists a type-specific pure strategy Nash Equilibrium with some threshold $\tilde{\gamma}_h$, which depends on $\mathcal{F}$ and h .*

1.2 *The equilibrium strategy of worker i can be described as*

$$a_i^* = \begin{cases} V & \text{if } \gamma_i \leq \tilde{\gamma}_h \\ D & \text{if } \gamma_i > \tilde{\gamma}_h \end{cases}$$

1.3 *Let j and k be two discrete probability distributions describing the stochastic population size of the game. Assume that j first-order stochastically dominates k . Define $\tilde{\gamma}_j$ as the equilibrium threshold for distribution j and $\tilde{\gamma}_k$ as the threshold for the distribution k . Then, we have $\tilde{\gamma}_k > \tilde{\gamma}_j$.*

1.4 *Let g and z be two discrete probability distributions describing the stochastic population size of the game. Assume that z is a mean-preserving spread of g . Define $\tilde{\gamma}_z$ as the equilibrium threshold for distribution z and $\tilde{\gamma}_g$ as the threshold for the distribution g . Then, we have $\tilde{\gamma}_z > \tilde{\gamma}_g$.*

The proof for Proposition 1 can be found in Appendix 3.A.1. The three take-aways from the proposition are that, first, players volunteer less given a larger probability of being in bigger groups. In other words, an increase in the (mean) group size decreases volunteering given that the increase is due to a shift in the distribution according to first-order stochastic dominance. Note that this also takes into account an increase in the group size when the size is certain. Second, those players with lower costs should be more likely to volunteer, as $P(\gamma_i \leq \tilde{\gamma}_h) = F(\tilde{\gamma}_h)$ describes the probability of an arbitrary player volunteering given a_i^* . Lastly, higher uncertainty about the group size leads to higher volunteering rates.

3.A.1 Proofs

Proof of Proposition 1: Existence of the Equilibrium: Assume players play V (Volunteer) and D (Defect) in pure strategies. Thus, for player i it must be the case that $EU_i(V) \geq EU_i(D)$ or vice versa.

Hence, we have

$$\begin{aligned}
 EU_i(V) &\geq EU_i(D) \\
 b_i - c_i &\geq b_i P(X_{-i} > 0) \\
 b_i - c_i &\geq b_i(1 - P(X_{-i} = 0)) \\
 \frac{c_i}{b_i} &\leq P(X_{-i} = 0) \\
 \gamma_i &\leq P(X_{-i} = 0)
 \end{aligned} \tag{3.3}$$

Note that for player i the strategy π_i is then:

$$a_i = \begin{cases} V & \text{if } \gamma_i \leq P(X_{-i} = 0) \\ D & \text{if } \gamma_i > P(X_{-i} = 0) \end{cases} \tag{3.4}$$

Note that there exists a type $\tilde{\gamma}_h$, which is indifferent between volunteering and defection.

Note that for this indifferent type $\tilde{\gamma}_h$ we have $\tilde{\gamma}_h = P(X_{-i} = 0)$ and that players with $\gamma_i > \tilde{\gamma}_h$ will not volunteer. The probability for some player i to be above the threshold $P(\gamma_i > \tilde{\gamma}_h)$ is equal to $1 - F(\tilde{\gamma}_h)$. Thus, we have

$$P(X_{-i} = 0) = \sum_{n \in \tilde{N}} h(n)(1 - F(\tilde{\gamma}_h))^{n-1} \tag{3.5}$$

Consider now the indifferent type and note that for him

$$\begin{aligned}
 \tilde{\gamma}_h &= P(X_{-i} = 0) \\
 \Leftrightarrow \tilde{\gamma}_h &= \sum_{n \in \tilde{N}} h(n)(1 - F(\tilde{\gamma}_h))^{n-1}
 \end{aligned} \tag{3.6}$$

holds with $\tilde{\gamma}_h$ being the solution to the fixed point condition. This solution will depend on $F(\cdot)$ and h . Given the strategy in (3.4)

$$a_i^* = \begin{cases} V & \text{if } \gamma_i \leq \tilde{\gamma}_h \\ D & \text{if } \gamma_i > \tilde{\gamma}_h \end{cases} \tag{3.7}$$

describes the strategy of player i in equilibrium. Note that $P(\gamma_i = \tilde{\gamma}_h) = 0$ as $F(\cdot)$ is atomless. There will never be a player i with a cost-benefit ratio, which is exactly equal to the threshold. Thus, we can neglect the case of $\gamma_i = \tilde{\gamma}_h$.²⁹ This proves that there exists a type-specific pure strategy Nash Equilibrium, which is described by the equilibrium threshold $\tilde{\gamma}_h$.

Uniqueness of the solution $\tilde{\gamma}_h$:

The right-hand side (RHS) of Equation 3.6, $P(X_{-i} = 0) = \sum_{n \in \tilde{N}} h(n)(1 - F(\tilde{\gamma}_h))^{n-1}$ is decreasing in $\tilde{\gamma}_h$ and strictly decreasing for $\tilde{\gamma}_h \in [\underline{\omega}, \bar{\omega}]$. The left-hand side (LHS) of Equation 3.6 is strictly increasing in $\tilde{\gamma}_h$. For $\tilde{\gamma}_h = \underline{\omega} < 1$, the RHS becomes one, as $F(\underline{\omega}) = 0$. Similarly, for $\tilde{\gamma}_h = \bar{\omega} > 0$, the RHS equals to zero as $F(\bar{\omega}) = 1$. Naturally $P(X_{-i} = 0)$ will be bounded by zero and one, as it is a common probability. Thus, there can only be one solution to the fixed point condition 3.6. This proves that $\tilde{\gamma}_h$ is unique.

Proof of Proposition 1.3: Consider two discrete probability distributions, j and k , which describe the group size in the game, and assume that j first-order stochastically dominates k . Note that this implies that the mean group size N is greater given distribution j compared to k . By the definition of first-order stochastic dominance, we have for every strictly increasing function $u(n)$ that $\sum_{n \in \tilde{N}} j(n)u(n) > \sum_{n \in \tilde{N}} k(n)u(n)$. Conversely, we get $\sum_{n \in \tilde{N}} k(n)c(n) > \sum_{n \in \tilde{N}} j(n)c(n)$ for $c(n) = -u(n)$.

Note that $(1 - F(\tilde{\gamma}_h))^{n-1}$ is strictly decreasing in n given $\tilde{\gamma}_h \in (\underline{\omega}, \bar{\omega})$. The latter follows directly from the assumption that $\tilde{N}$ is finite, $\{1\} \cap \tilde{N} = \emptyset$ and $\underline{\omega} < \bar{\omega}$. Thus, setting $c(n) := (1 - F(\tilde{\gamma}_h))^{n-1}$, we therefore get

$$\sum_{n \in \tilde{N}} k(n)(1 - F(\tilde{\gamma}_h))^{n-1} > \sum_{n \in \tilde{N}} j(n)(1 - F(\tilde{\gamma}_h))^{n-1} \quad (3.8)$$

Importantly, Equation 3.8 holds for any equilibrium threshold $\tilde{\gamma}_h$ with any arbitrary discrete probability distribution h . This is because $F(\cdot)$ stays unchanged and we only alter the distribution h , which then has an influence on the solution $\tilde{\gamma}_h$. Given Equality 3.6 and Inequality 3.8, we then have

²⁹The existence of a player with a threshold is not required for the threshold to exist.

$$\tilde{\gamma}_k = \sum_{n \in \tilde{N}} k(n)(1 - F(\tilde{\gamma}_k))^{n-1} > \sum_{n \in \tilde{N}} j(n)(1 - F(\tilde{\gamma}_k))^{n-1} \quad (3.9)$$

and

$$\sum_{n \in \tilde{N}} k(n)(1 - F(\tilde{\gamma}_j))^{n-1} > \sum_{n \in \tilde{N}} j(n)(1 - F(\tilde{\gamma}_j))^{n-1} = \tilde{\gamma}_j, \quad (3.10)$$

where $\tilde{\gamma}_k$ and $\tilde{\gamma}_j$ denote the equilibrium thresholds for the respective probability distributions of n . Lastly, assume by contradiction that $\tilde{\gamma}_k \leq \tilde{\gamma}_j$. Then we have

$$\begin{aligned} \tilde{\gamma}_k &\leq \tilde{\gamma}_j \\ \Leftrightarrow \tilde{\gamma}_k &\leq \sum_{n \in \tilde{N}} j(n)(1 - F(\tilde{\gamma}_j))^{n-1} \\ \Leftrightarrow \sum_{n \in \tilde{N}} j(n)(1 - F(\tilde{\gamma}_k))^{n-1} &< \sum_{n \in \tilde{N}} j(n)(1 - F(\tilde{\gamma}_j))^{n-1} \end{aligned} \quad (3.11)$$

using Equation 3.6 in the second and Equation 3.9 in the last step.

However, $(1 - F(\tilde{\gamma}_h))^{n-1}$ is (weakly) decreasing in $\tilde{\gamma}_h$ for $\forall n \in \tilde{N}$. Hence, $\tilde{\gamma}_k \leq \tilde{\gamma}_j$ implies that

$$\begin{aligned} (1 - F(\tilde{\gamma}_k))^{n-1} &\geq (1 - F(\tilde{\gamma}_j))^{n-1} \forall n \in \tilde{N} \\ \Leftrightarrow j(n)(1 - F(\tilde{\gamma}_k))^{n-1} &\geq j(n)(1 - F(\tilde{\gamma}_j))^{n-1} \forall n \in \tilde{N} \\ \Leftrightarrow \sum_{n \in \tilde{N}} j(n)(1 - F(\tilde{\gamma}_k))^{n-1} &\geq \sum_{n \in \tilde{N}} j(n)(1 - F(\tilde{\gamma}_j))^{n-1}, \end{aligned} \quad (3.12)$$

as $j(n) \geq 0 \forall n$ and $(1 - F(\tilde{\gamma}_{j,k}))^{n-1} > 0 \forall n$ by the observation that $\tilde{\gamma}_{j,k} \in (\underline{\omega}, \bar{\omega})$. This yields the desired contradiction if we compare Equation 3.11 and Equation 3.12 and proves that $\tilde{\gamma}_k > \tilde{\gamma}_j$. □

Proof of Proposition 1.4:

Let h be some arbitrary discrete probability distribution describing the population size. Denote $\tilde{\gamma}_h$ as the equilibrium threshold, as in Proposition 1. Note that given $\tilde{\gamma}_h \in (\underline{\omega}, \bar{\omega})$ we have $(1 - F(\tilde{\gamma}_h)) \in (0, 1)$ for any equilibrium threshold. Hence, for any given distribution h the function $(1 - F(\tilde{\gamma}_h))^{n-1}$ will be strictly decreasing and strictly convex in $n \geq 2$ for a fixed equilibrium threshold $\tilde{\gamma}_h$.

Assume that z and g are two arbitrary discrete probability distributions of n and that z is a mean-preserving spread of g ($z_{mps}g$). Note that this implies that g second-

order stochastically dominates z . Thus, for every strictly increasing and strictly concave function $u(n)$, it must hold that $\sum_{n \in \tilde{N}} z(n)u(n) < \sum_{n \in \tilde{N}} g(n)u(n)$. Conversely, we get $\sum_{n \in \tilde{N}} z(n)c(n) > \sum_{n \in \tilde{N}} g(n)c(n)$ for any strictly convex function $c(n) = -u(n)$. Setting $c(n) := (1 - F(\tilde{\gamma}_h))^{n-1}$, we therefore get

$$\sum_{n \in \tilde{N}} z(n)(1 - F(\tilde{\gamma}_h))^{n-1} > \sum_{n \in \tilde{N}} g(n)(1 - F(\tilde{\gamma}_h))^{n-1} \quad (3.13)$$

By following an analogous approach as in the proof of Proposition 1.3, it then follows that $\tilde{\gamma}_z > \tilde{\gamma}_g$.

□

Appendix 3.B Additional figures and tables

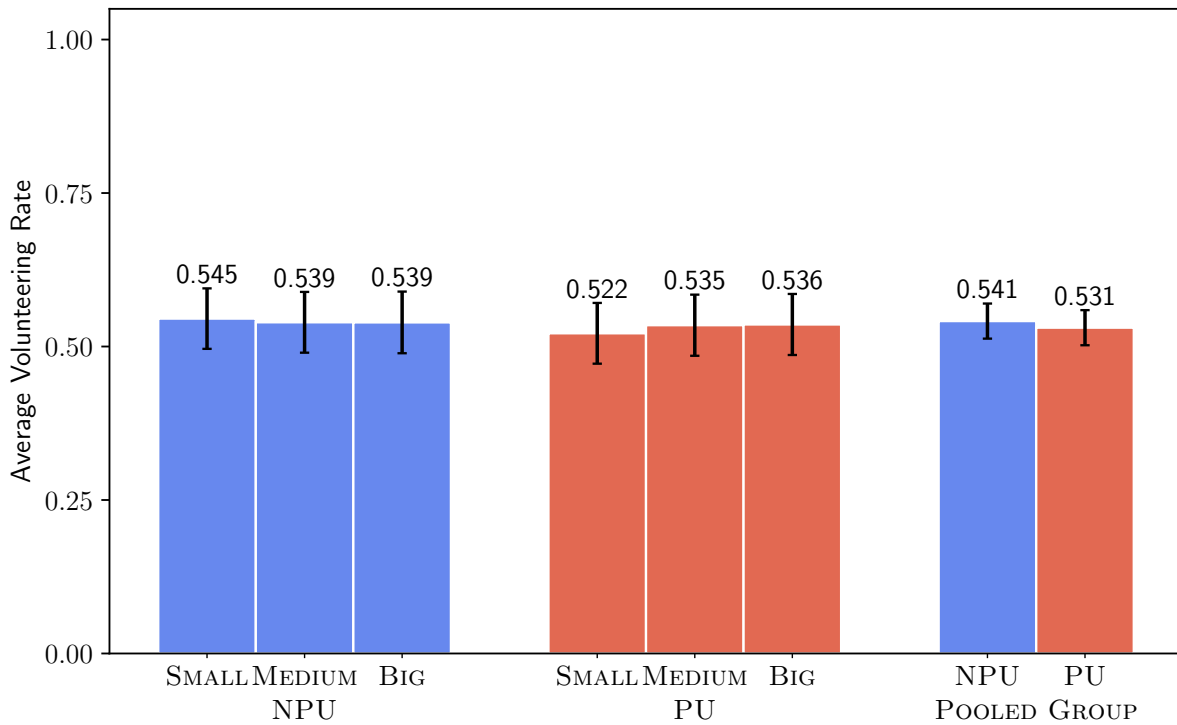


Figure 3.6: The average volunteering rate across the different mean group size for Population Uncertainty (PU) and no population uncertainty (NPU). The error bars represent 95%-confidence intervals. The data was gathered for an earlier version of this paper. For main results of this paper, we only use the data from the replication study that does not consider population uncertainty. The earlier version of the paper can be accessed here <https://osf.io/4k5y6/>.

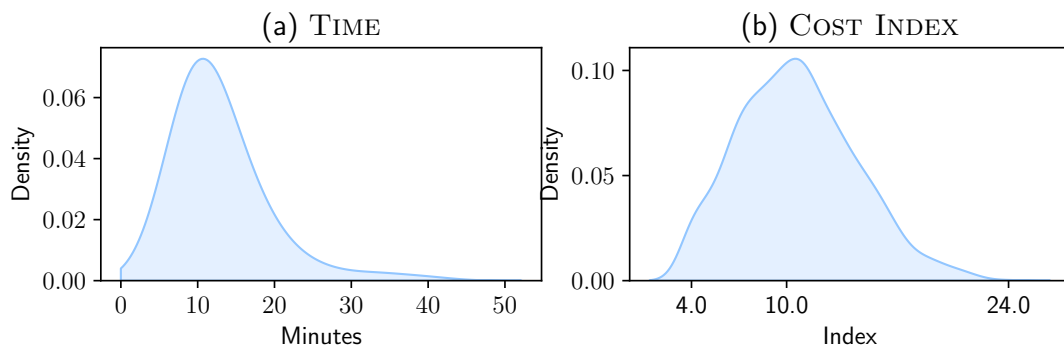


Figure 3.7: Kernel density estimation of the distributions of the costs of volunteering as measured by the time it took participants to complete the first stage of the field experiment (left) and the subjective effort costs (right). The bandwidth (bw) for TIME is chosen with k-fold cross validation for $k = 7$ at $\text{bw} = 0.562$ for the Gaussian kernel. For the bandwidth of the COST INDEX, Silverman's Rule of Thumb was used. We restrict the data in the left plot to values between the 2nd and 98th percentile for this visualization.

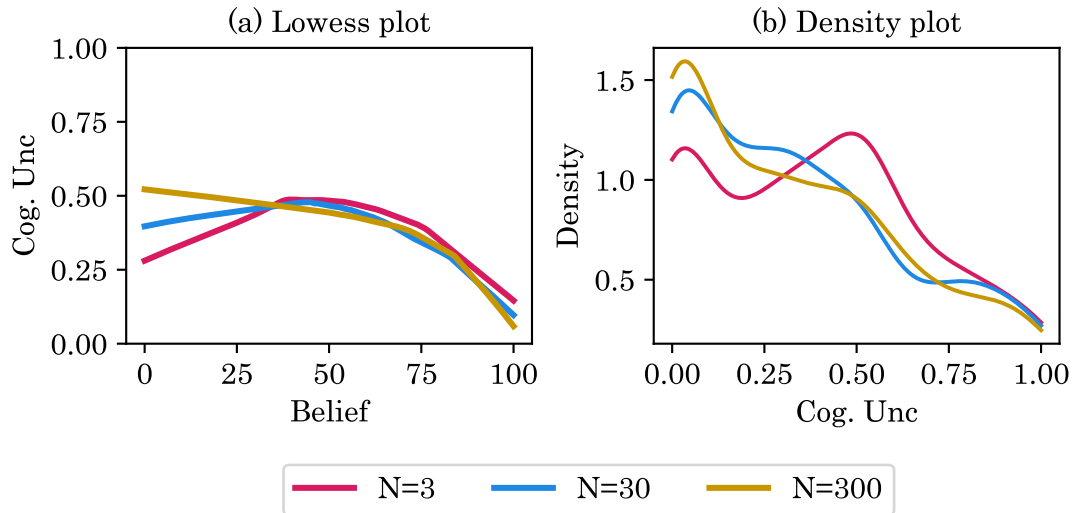


Figure 3.8: The local regression shows the hump shaped relationship between cognitive uncertainty and beliefs. The kernel density plot for cognitive uncertainty uses a Gaussian kernel and Silverman's Rule of Thumb to derive the bandwidth. The values are restricted to be between zero and one.

Table 3.5: Logistic Regression estimating the probability to volunteer

<i>Dependent variable:</i>	
Volunteering Choice	
MEDIUM GROUP	0.152 (0.117)
BIG GROUP	0.055 (0.117)
Constant	0.054 (0.083)
Observations	1,754
Akaike Inf. Crit.	2,429.174
<i>Note:</i> *p<0.1; **p<0.05; ***p<0.01 Standard errors in parentheses	

Table 3.6: Estimated average marginal effects for the main treatment variable

	AME	SE	z	p	lower	upper
BIG GROUP	0.06	0.12	0.47	0.64	-0.17	0.28
MEDIUM GROUP	0.15	0.12	1.30	0.19	-0.08	0.38

Table 3.7: Estimated average marginal effects for model specification (3) in Table 3.2

	AME	SE	z	p	lower	upper
COST INDEX	-0.09	0.01	-8.02	0.00	-0.11	-0.07
TIME	-0.00	0.01	-0.11	0.91	-0.02	0.02

Table 3.8: The mean of different socioeconomic and demographic variables for all treatments

Region & City Size		Employment Status	
Statistic	Mean	Variable	Mean
City size: 200k-500k	0.12	Employed	0.57
City size: 50k-200k	0.19	Searching	0.04
City size: 500k-1500k	0.14	Retired	0.02
City size: Less 50k	0.41	Not employed	0.02
City size: More 1500k	0.14	Self employed	0.11
		Other	0.03
		Student	0.21
Educational Level		Age & Gender	
Statistic	Mean	Statistic	Mean
High school	0.32	Female	0.48
Bachelor	0.19	Age: 18-25	0.26
Apprenticeship	0.27	Age: 26-35	0.40
Master	0.19	Age: 36-45	0.19
School not finished	0.02	Age: 46-55	0.09
		Age: 55+	0.05

Appendix 3.C Confirming the robustness of the results

Conditional cooperation is a strong driver of volunteering and can explain the lack of average group size effects. This result is robust to various factors. In this section, we rule out several other possible explanations. First, we show that randomization into treatments worked. This can be seen in Figure 3.9, where we plot the average values for different control variables by treatment. Second, we show that the way how we elicited the beliefs does not influence outcomes. Then, we argue that the group size was salient and that participants were aware of it. We rule out that workplace-related reputational concerns, economic preferences, or gender differences drive results. Lastly, we highlight

that beliefs are the most important explanatory variable relative to other measures of interest.

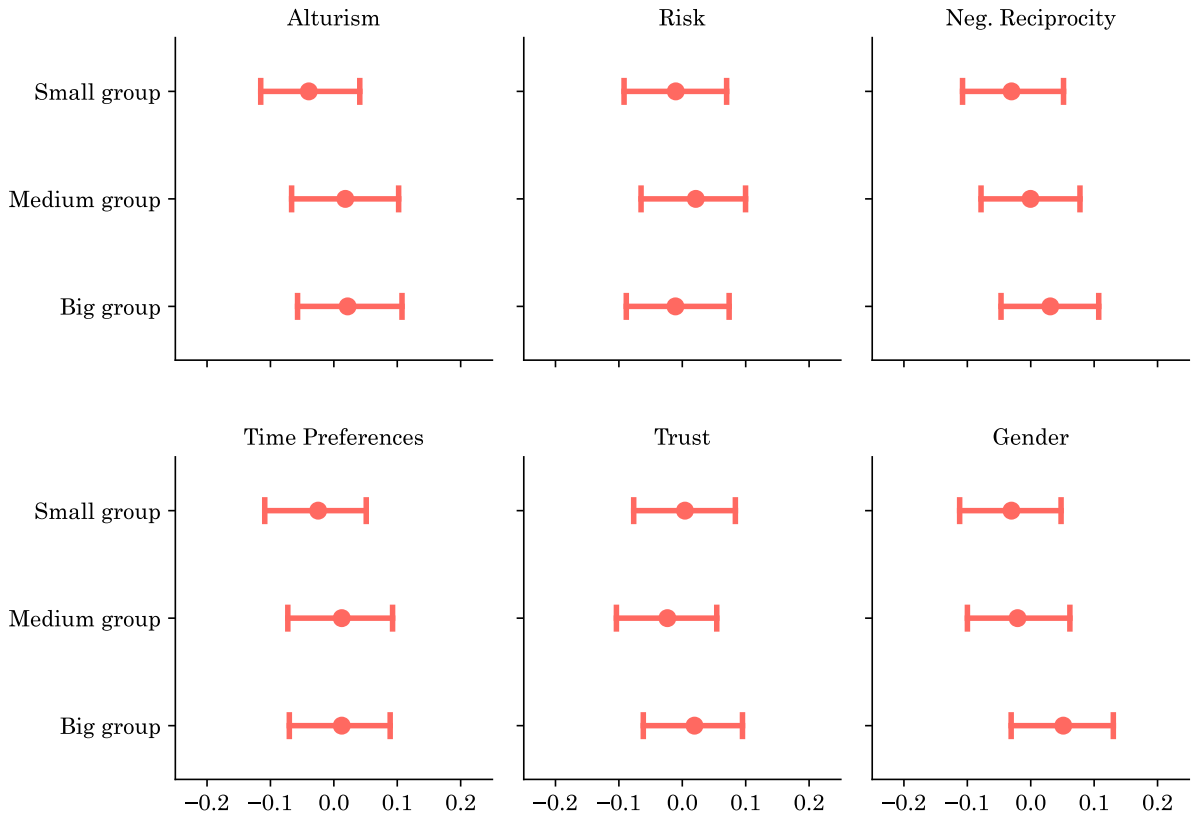


Figure 3.9: Average values and 95% confidence intervals of several observables across treatments.

3.C.1 Belief order effects

In the experiment, we randomized the order of the belief elicitation and the volunteering decision. In ACTION-BELIEF, we asked first for the volunteering decision and then for the belief about the decision of the other team members. In BELIEF-ACTION, the order was reversed. We find no evidence that this order influenced the outcomes systematically. The average reported belief in BELIEF-ACTION is 1.17 percentage points higher than in ACTION-BELIEF. Yet, the difference is either only weakly significant or not statistically significant depending on the exact test specification (two-sided MWU test $p = 0.09$, two-sided t-test $p = 0.43$). Importantly, there are no significant differences in volunteering rates between ACTION-BELIEF and BELIEF-ACTION (χ^2 -test $p = 0.34$). In Table 3.9, we replicate the results on conditional cooperation (see Table 3.4) but include the dummy variable AB that is equal to one for the ACTION-BELIEF variation. The respective coef-

ficient and the interaction terms are insignificant in all model specifications, which shows that conditional cooperation is also robust to order effects.

3.C.2 Prominence of group size

One possible explanation for why subjects did not react to the group size in our main treatments might be that they, in general, did not pay attention to the description of the game and the group size in particular. For example, it might be the case that they simply overlooked the number of players in the team. To minimize this concern, the instructions were simple and easy to understand. Second, the size of the team was mentioned two times on the decision screen and prominently directly before making the choice (see Instructions in Appendix 3.D.3).

It is still possible that workers who read the instructions did not fully contemplate their decision, the payoff consequences of their actions, and the meaning of the team size for their potential payoffs. We note that workers spent more time on the decision screen in the main treatments than in the baseline treatments (Mean VOD = 57.83 seconds, Mean BASELINE (N=1) = 42.46 seconds, $p\text{-value} < 0.001$, two-sided Mann-Whitney test).³⁰ This suggests that workers indeed took the strategic situation into account.

Additionally, we asked participants in the survey at the end of the experiment if they could recall their group size. Table 3.10 shows the share of participants that reported the correct group size. The large majority of participants in all treatments correctly remembered the group size with a correct recall among 93% of all participants. Hence, participants noticed the group size when they took their volunteering decision and that the lack of average treatment effects is not driven by a lack of salience of the group size.

3.C.3 Workplace-related reputational concerns

One might further argue that many participants volunteered because of reputational concerns. Even though we explicitly pointed out that the volunteering choice has no effect on their reputation, they could have been concerned that they might receive a negative rating on the platform. These concerns should be most pronounced for those workers who

³⁰For the analysis of the attention time, we restrict the sample to observations with values between the 98th and the 2nd percentile. Including those outliers does not alter the results substantially. For the main treatments, we restrict the data to participants in the Action-Belief treatment variation as it has the same decision screen as in the baseline.

Table 3.9: The Volunteering choice explained by the order of the belief elicitation (AB), beliefs and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>		
	Volunteering Choice		
	(1)	(2)	(3)
AB	−0.017 (0.023)	0.0001 (0.058)	−0.007 (0.081)
BELIEF	0.005*** (0.0004)	0.005*** (0.0005)	0.007*** (0.001)
MEDIUM GROUP	−0.009 (0.027)	−0.010 (0.039)	0.133 (0.087)
BIG GROUP	−0.037 (0.028)	−0.050 (0.039)	0.156* (0.087)
BELIEF × AB		−0.0004 (0.001)	−0.0004 (0.001)
MEDIUM GROUP × AB		0.001 (0.055)	−0.083 (0.125)
BIG GROUP × AB		0.027 (0.056)	0.105 (0.129)
BELIEF × MEDIUM GROUP			−0.002** (0.001)
BELIEF × BIG GROUP			−0.003*** (0.001)
BELIEF × MEDIUM GROUP × AB			0.001 (0.002)
BELIEF × BIG GROUP × AB			−0.001 (0.002)
Constant	0.223*** (0.031)	0.215*** (0.039)	0.111** (0.053)
Observations	1,754	1,754	1,754

Note:

*p<0.1; **p<0.05; ***p<0.01
Robust standard errors in parentheses.

Table 3.10: Belief overview by treatment

Group size	Share of participants with correct group recall
N=3	0.91
N=30	0.95
N=300	0.94
Pooled	0.93

have frequently worked on the site in the past, and by implication, also expect so in the future. Therefore, we asked participants how often they had worked on the platform before to construct our REPUTATION score and observe no statistically significant differences between treatments ($p = 0.48$, F-test). We find a small statistically significant effect of REPUTATION on volunteering in a linear probability model with robust standard errors (see Table 3.12 in the Appendix). Yet, we do not find systematic interaction effects of REPUTATION with our treatment variables, which could explain our absence of treatment effects. It means that the null effects are robust to reputational concerns.

3.C.4 Economic preferences

One might also expect differences in risk preferences, altruism, or “economic preferences” to explain our results more generally. As explained above, we collected several measures of economic preferences based on the survey measure by Falk et al. (2022), including trust, time preferences, negative reciprocity, and altruism, plus an additional measure of efficiency concerns closely linked to the Volunteer’s Dilemma (see Section 3.D.5 for the exact wording of the questions). Also here, all variables are balanced across treatments (see Figure 3.9). The differences between the indicators are never statistically significant (F-test p-values between 0.31 and 0.82). We find positive main effects for time preferences, negative reciprocity, and efficiency concerns. (see Tables 3.12-3.19). Importantly, there is again almost no systematic significant interaction effect between treatments on any of the preference measures, confirming the robustness of our finding. A notable exception is efficiency consideration. In line with expectations, participants volunteer more likely if they report that they are more willing to make an effort if many profit from it. Furthermore, this effect is increasing in the group size (see Table 3.14). Yet, the effect size is small relative to the influence of beliefs on volunteering. Thus, while efficiency con-

siderations influence the behavior of participants, they play a subordinate role compared to conditional cooperation (see Table 3.11 for a comparison).

3.C.5 Gender differences in volunteering

Finally, we find the same pattern as above for gender differences in volunteering. Men and women are fairly balanced across treatments ($p=0.31$, F-test). Women in general volunteer more often, but also here the insignificant interaction effect with our treatments confirms the robustness of our result (see Table 3.13 in the Appendix).

3.C.6 Beliefs relative to other explanatory variables

In Table 3.11, we display linear probability models, in which we regress the volunteering choice on the set of control variables that we elicited at the end of the experiment. All variables are z-transformed to allow for comparisons. While different factors are correlated with the volunteering decision, the belief participants have about the volunteering decision of the other team members (BELIEF) is the strongest predictor (Model specification 1). Furthermore, BELIEF and EFFICIENCY³¹ are the only two variables we elicit in the experiment that shows a systematic interaction with the treatment variables (Model specification 2 and 3).³² Also, when accounting for the significant interactions in the model specification (4), the beliefs are the strongest predictor. Hence, we find clear evidence that beliefs are a strong driver of volunteering compared to other possible channels.

³¹For a description of the control question see here 3.D.5

³²For the other variables the interactions are either insignificant or only a single group size shows a (weakly) significant interactions.

Table 3.11: The volunteering choice explained by different control variables in a linear probability model with robust standard errors

	<i>Dependent variable:</i>			
	Volunteering Choice			
	(1)	(2)	(3)	(4)
MEDIUM GROUP	−0.017	−0.028	−0.017	−0.029
BIG GROUP	−0.034	−0.037	−0.033	−0.037
BELIEF	0.139***	0.200***	0.139***	0.204***
EFFICIENCY	0.110***	0.111***	0.060***	0.055***
TIME PREF.	0.027**	0.028**	0.028**	0.029**
TRUST	−0.016	−0.016	−0.016	−0.016
RISK	−0.041***	−0.041***	−0.042***	−0.042***
N. RECIPROCITY	0.012	0.012	0.012	0.012
ALTRUISM	0.004	0.004	0.003	0.003
REPUTATION	0.027**	0.027**	0.027**	0.027**
COST INDEX	−0.051***	−0.048***	−0.050***	−0.047***
FREQ. TASK	−0.005	−0.005	−0.003	−0.003
MEDIUM GROUP × BELIEF		−0.061**		−0.067***
BIG GROUP × BELIEF		−0.121***		−0.128***
MEDIUM GROUP × EFFICIENCY			0.077***	0.082***
BIG GROUP × EFFICIENCY			0.076***	0.086***
Constant	0.553***	0.564***	0.553***	0.564***
Observations	1,694	1,694	1,694	1,694
Adjusted R ²	0.167	0.175	0.171	0.181

Note:

*p<0.1; **p<0.05; ***p<0.01

Robust standard errors are used to derive the p-values.

3.C.7 Robustness checks with regard to control variables

Table 3.12: The volunteering choice explained by REPUTATION and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>			
	Volunteering Choice			
	(1)	(2)	(3)	(4)
REPUTATION	0.026** (0.012)	0.025** (0.012)	0.050** (0.022)	0.050* (0.026)
MEDIUM GROUP		0.024 (0.030)	0.044 (0.042)	0.044 (0.041)
BIG GROUP		0.013 (0.030)	0.016 (0.042)	0.016 (0.042)
AB		-0.019 (0.024)	-0.004 (0.042)	-0.004 (0.042)
MEDIUM GROUP $\times$ AB			-0.041 (0.059)	-0.041 (0.059)
BIG GROUP $\times$ AB			-0.008 (0.059)	-0.008 (0.059)
MEDIUM GROUP $\times$ REPUTATION			-0.037 (0.028)	-0.026 (0.040)
BIG GROUP $\times$ REPUTATION			-0.050* (0.029)	-0.061 (0.039)
AB $\times$ REPUTATION			0.006 (0.024)	0.005 (0.036)
MEDIUM GROUP $\times$ AB $\times$ REPUTATION				-0.020 (0.055)
BIG GROUP $\times$ AB $\times$ REPUTATION				0.028 (0.059)
Constant	0.539*** (0.012)	0.536*** (0.024)	0.529*** (0.029)	0.529*** (0.029)
Observations	1,703	1,703	1,703	1,703

Note:

*p<0.1; **p<0.05; ***p<0.01
Robust standard errors in parentheses.

Table 3.13: The Volunteering choice explained by gender and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>			
	Volunteering Choice			
	(1)	(2)	(3)	(4)
WOMAN	0.064*** (0.012)	0.065*** (0.012)	0.042* (0.024)	0.055* (0.029)
MEDIUM GROUP		0.037 (0.029)	0.052 (0.041)	0.051 (0.041)
BIG GROUP		0.009 (0.029)	0.016 (0.041)	0.015 (0.041)
AB		-0.028 (0.024)	-0.014 (0.041)	-0.014 (0.041)
MEDIUM GROUP $\times$ AB			-0.029 (0.058)	-0.029 (0.058)
BIG GROUP $\times$ AB			-0.018 (0.058)	-0.019 (0.058)
MEDIUM GROUP $\times$ WOMAN			0.023 (0.029)	0.015 (0.041)
BIG GROUP $\times$ WOMAN			0.082*** (0.029)	0.048 (0.041)
AB $\times$ WOMAN			-0.023 (0.024)	-0.050 (0.041)
MEDIUM GROUP $\times$ AB $\times$ WOMAN				0.016 (0.058)
BIG GROUP $\times$ AB $\times$ WOMAN				0.067 (0.058)
Constant	0.531*** (0.012)	0.530*** (0.024)	0.522*** (0.029)	0.523*** (0.029)
Observations	1,754	1,754	1,754	1,754

Note:

*p<0.1; **p<0.05; ***p<0.01
Robust standard errors in parentheses.

Table 3.14: The Volunteering choice explained by EFFICIENCY and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>			
	Volunteering Choice			
	(1)	(2)	(3)	(4)
EFFICIENCY	0.132*** (0.011)	0.132*** (0.011)	0.078*** (0.022)	0.082*** (0.026)
MEDIUM GROUP		0.034 (0.028)	0.058 (0.040)	0.059 (0.040)
BIG GROUP		0.015 (0.028)	0.018 (0.040)	0.019 (0.040)
AB		-0.020 (0.023)	-0.003 (0.041)	-0.003 (0.041)
MEDIUM GROUP $\times$ AB			-0.051 (0.057)	-0.050 (0.057)
BIG GROUP $\times$ AB			-0.006 (0.057)	-0.005 (0.057)
MEDIUM GROUP $\times$ EFFICIENCY			0.080*** (0.028)	0.088** (0.038)
BIG GROUP $\times$ EFFICIENCY			0.074*** (0.027)	0.055 (0.037)
AB $\times$ EFFICIENCY			0.008 (0.022)	0.0001 (0.040)
MEDIUM GROUP $\times$ AB $\times$ EFFICIENCY				-0.018 (0.056)
BIG GROUP $\times$ AB $\times$ EFFICIENCY				0.040 (0.054)
Constant	0.531*** (0.011)	0.525*** (0.023)	0.516*** (0.029)	0.516*** (0.029)
Observations	1,754	1,754	1,754	1,754

Note:

*p<0.1; **p<0.05; ***p<0.01
Robust standard errors in parentheses.

Table 3.15: The Volunteering choice explained by ALTRUISM and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>			
	Volunteering Choice			
	(1)	(2)	(3)	(4)
ALTRUISM	0.017 (0.012)	0.016 (0.012)	0.002 (0.024)	−0.004 (0.028)
MEDIUM GROUP		0.037 (0.029)	0.053 (0.041)	0.052 (0.041)
BIG GROUP		0.013 (0.029)	0.017 (0.041)	0.017 (0.041)
AB		−0.023 (0.024)	−0.012 (0.041)	−0.011 (0.041)
MEDIUM GROUP × AB			−0.029 (0.058)	−0.028 (0.058)
BIG GROUP × AB			−0.006 (0.059)	−0.007 (0.059)
MEDIUM GROUP × ALTRUISM			0.030 (0.029)	0.074* (0.038)
BIG GROUP × ALTRUISM			0.048 (0.029)	0.026 (0.040)
AB × ALTRUISM			−0.027 (0.024)	−0.011 (0.044)
MEDIUM GROUP × AB × ALTRUISM				−0.088 (0.058)
BIG GROUP × AB × ALTRUISM				0.039 (0.059)
Constant	0.531*** (0.012)	0.525*** (0.024)	0.519*** (0.029)	0.519*** (0.029)
Observations	1,754	1,754	1,754	1,754

*Note:**p<0.1; **p<0.05; ***p<0.01
Robust standard errors in parentheses.

Table 3.16: The volunteering choice explained by N. RECIPROCITY and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>			
	Volunteering Choice			
	(1)	(2)	(3)	(4)
N. RECIPROCITY	0.037*** (0.012)	0.037*** (0.012)	0.032 (0.024)	0.013 (0.028)
MEDIUM GROUP		0.037 (0.029)	0.055 (0.041)	0.055 (0.041)
BIG GROUP		0.012 (0.029)	0.014 (0.041)	0.013 (0.041)
AB		-0.023 (0.024)	-0.009 (0.041)	-0.008 (0.041)
MEDIUM GROUP $\times$ AB			-0.035 (0.058)	-0.036 (0.058)
BIG GROUP $\times$ AB			-0.005 (0.058)	-0.005 (0.058)
MEDIUM GROUP $\times$ N. RECIPROCITY			0.010 (0.029)	0.027 (0.041)
BIG GROUP $\times$ N. RECIPROCITY			0.024 (0.029)	0.062 (0.040)
AB $\times$ N. RECIPROCITY			-0.012 (0.024)	0.026 (0.041)
MEDIUM GROUP $\times$ AB $\times$ N. RECIPROCITY				-0.036 (0.058)
BIG GROUP $\times$ AB $\times$ N. RECIPROCITY				-0.081 (0.059)
Constant	0.531*** (0.012)	0.526*** (0.024)	0.519*** (0.029)	0.519*** (0.029)
Observations	1,754	1,754	1,754	1,754

Note:

*p<0.1; **p<0.05; ***p<0.01
Robust standard errors in parentheses.

Table 3.17: The volunteering choice explained by RISK and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>			
	Volunteering Choice			
	(1)	(2)	(3)	(4)
RISK	−0.007 (0.012)	−0.007 (0.012)	0.001 (0.024)	0.005 (0.029)
MEDIUM GROUP		0.038 (0.029)	0.053 (0.041)	0.053 (0.041)
BIG GROUP		0.014 (0.029)	0.016 (0.041)	0.015 (0.041)
AB		−0.023 (0.024)	−0.010 (0.041)	−0.010 (0.041)
MEDIUM GROUP × AB			−0.030 (0.058)	−0.031 (0.058)
BIG GROUP × AB			−0.007 (0.058)	−0.007 (0.059)
MEDIUM GROUP × RISK			0.004 (0.029)	−0.025 (0.041)
BIG GROUP × RISK			0.026 (0.029)	0.040 (0.041)
AB × RISK			−0.036 (0.024)	−0.045 (0.041)
MEDIUM GROUP × AB × RISK				0.059 (0.059)
BIG GROUP × AB × RISK				−0.031 (0.059)
Constant	0.531*** (0.012)	0.525*** (0.024)	0.519*** (0.029)	0.519*** (0.029)
Observations	1,754	1,754	1,754	1,754

Note:

*p<0.1; **p<0.05; ***p<0.01

Robust standard errors in parentheses.

Table 3.18: The volunteering choice explained by TRUST and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>			
	Volunteering Choice			
	(1)	(2)	(3)	(4)
TRUST	0.019 (0.012)	0.018 (0.012)	0.043* (0.024)	0.027 (0.030)
MEDIUM GROUP		0.038 (0.029)	0.057 (0.041)	0.056 (0.041)
BIG GROUP		0.014 (0.029)	0.015 (0.041)	0.014 (0.041)
AB		−0.021 (0.024)	−0.008 (0.041)	−0.009 (0.041)
MEDIUM GROUP × AB			−0.034 (0.058)	−0.034 (0.058)
BIG GROUP × AB			−0.004 (0.059)	−0.003 (0.059)
MEDIUM GROUP × TRUST			−0.004 (0.029)	0.024 (0.040)
BIG GROUP × TRUST			0.018 (0.029)	0.035 (0.042)
AB × TRUST			−0.059** (0.024)	−0.028 (0.042)
MEDIUM GROUP × AB × TRUST				−0.058 (0.058)
BIG GROUP × AB × TRUST				−0.033 (0.059)
Constant	0.531*** (0.012)	0.524*** (0.024)	0.516*** (0.029)	0.517*** (0.029)
Observations	1,754	1,754	1,754	1,754

Note:

*p<0.1; **p<0.05; ***p<0.01
Robust standard errors in parentheses.

Table 3.19: The volunteering choice explained by TIME and all treatment dummies in a linear probability model with robust standard errors

	<i>Dependent variable:</i>			
	Volunteering Choice			
	(1)	(2)	(3)	(4)
TIME	0.075*** (0.012)	0.075*** (0.012)	0.076*** (0.022)	0.063** (0.027)
MEDIUM GROUP		0.035 (0.029)	0.057 (0.041)	0.056 (0.041)
BIG GROUP		0.011 (0.029)	0.014 (0.041)	0.013 (0.041)
AB		-0.018 (0.024)	-0.0002 (0.041)	0.001 (0.041)
MEDIUM GROUP $\times$ AB			-0.043 (0.058)	-0.044 (0.058)
BIG GROUP $\times$ AB			-0.007 (0.058)	-0.007 (0.058)
MEDIUM GROUP $\times$ TIME			-0.021 (0.029)	-0.009 (0.040)
BIG GROUP $\times$ TIME			0.008 (0.028)	0.036 (0.038)
AB $\times$ TIME			0.007 (0.023)	0.034 (0.039)
MEDIUM GROUP $\times$ AB $\times$ TIME				-0.026 (0.057)
BIG GROUP $\times$ AB $\times$ TIME				-0.058 (0.056)
Constant	0.531*** (0.012)	0.524*** (0.024)	0.516*** (0.029)	0.516*** (0.029)
Observations	1,754	1,754	1,754	1,754

Note:

*p<0.1; **p<0.05; ***p<0.01
Robust standard errors in parentheses.

Appendix 3.D Experiment instructions

Below are the translated instructions the participants received during the experiment. The experiment was conducted in German. The screenshots show the German version.

3.D.1 The introduction

Welcome and thank you for your support of our project. In this task, you will rate comments from users of an internet forum. With your help, we will better able to understand and assess the behavior of our users. A more detailed explanation of the task can be found in the section "Your task".

Please note:

The quality of the data resulting from rating the comments is very important to us. Furthermore, there are a lot of comments from the Internet forum, all of which should be evaluated.

To ensure data quality and efficient handling of the tasks, you will be working in a team. Your team consists of [PU: $\mu - s$ to $\mu + s$ members; NPU: exactly μ members]. You will initially work individually on the task described below.

Your task:

Below we will show you a series of pictures, each with a comment. These comments come from different users of an internet forum. For each of these comments, we have the following four questions for you:

1. "Is the comment friendly or hostile towards the group which is displayed in the photo?"

We would like to know from you how hostile you find the comment with regard to the topic shown in the picture. You should rate the comments on a scale of 1 to 9. 1 means very friendly, 9 means very hostile.

[Possible answers: Value in a Likert scale between 1 ("very friendly") to 9 ("very hostile") or "Not possible to rate"]

In some cases, the comments are difficult to evaluate. Please click on the checkbox "Not to rate" in these cases.

2. "Is the comment addressed to another user?"

In part, the comments you see are directed towards other users on the internet forum. We are interested in whether the comment is directed towards another user and, if so, whether she agrees or disapproves of the other user. Therefore, we ask you please to answer the following question for each comment: "Is the comment addressed to another user?"

[Possible Answers: "No", "Yes, agreeing", "Yes, rejecting", "Not possible to rate"]

In some cases, the comments are difficult to evaluate. Please click on the checkbox "Not possible to rate" in these cases.

3. "Should the comment be allowed in an online forum?"

Do you personally think that the comment should be allowed in an internet forum?

[Possible Answers: "No", "Yes", "Not possible to rate"]

4. "Which features apply to the comment?"

You will also find a list of features below the scale. Please click on any features that you think apply to the comment. If none of the features apply, just do not click on any.

[Possible Answers: "Contains negative prejudices", "Uses racist insults", "Contains offensive, degrading or derogatory words", "Calls for violence, threats or discrimination", "Uses sexist insults" or/and "Sexual orientation or gender is degraded or stigmatized"]

You will only see each comment once. When you have finished a page, please click on "Next". You will rate a total of 30 comments.

Please press "Next" to start the task. Thanks for your support.

3.D.2 The task

Kommentar 1 von insgesamt 30 Kommentaren

Das Bild:



Der Kommentar:

Wir sollten das gemeinsame Wohlergehen der Menschen über ein unreflektiertes temporäres "unwohlsein" stellen.

Ist der Kommentar freundlich oder feindselig gegenüber der im Foto dargestellten Gruppe?

sehr freundlich ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 sehr feindselig

☐ Nicht zu bewerten

Richtet sich der Kommentar an einen anderen Nutzer?

☐ Nein ☐ Ja, er ist zustimmend ☐ Ja, er ist ablehnend

☐ Nicht zu bewerten

Sollte dieser Kommentar in einem Internetforum erlaubt sein?

☐ Nein ☐ Ja ☐ Ich weiß nicht

Welche Merkmale treffen auf den Kommentar zu?

- ☐ Beinhaltet negative Vorurteile
- ☐ Nutzt rassistische Beleidigungen
- ☐ Beinhaltet beleidigende, erniedrigende oder abwertende Worte
- ☐ Ruft zu Gewalt, Drohungen oder Diskriminierung auf
- ☐ Nutzt sexistische Beleidigungen
- ☐ Die sexuelle Orientierung oder das Geschlecht/Gender wird herabgesetzt oder stigmatisiert

Weiter

Figure 3.10: Screenshot of the task.

3.D.3 The volunteering decision

Additional task for securing the data quality in the team

Thank you for rating the 30 comments. You completed now the main task and will receive 0.90 € for it. In total, it took you TIME minutes to complete the task. Participants in this task normally need between 7.5 Minutes and 15.0 minutes to rate 30 comments.

We now need exactly one volunteer from your team of exactly N person.

Instructions for the additional task in the team:

To assess the comment ratings in your team better, and to improve data quality, 30 more comments need to be rated by your team.

- All team members receive a bonus of 0.90 € each, if one person completes this additional task.
- All persons in your team are offered this additional task.
- It is enough for one person in your team of exactly N person to volunteer. All team members will then receive the bonus payment. It is possible for more than one person to perform the task. The bonus then does not increase.
- You will also receive the bonus if you do not volunteer, but another person is found.
- If nobody volunteers in your team, nobody will receive a bonus of 0.90 €.
- Your decision to volunteer in your team or not will not affect your reputation on the Clickworker platform or the money you earned so far.
- All participants in your group will make this decision for themselves and, if necessary, also work on the task. It is therefore possible that all participants, a smaller number, or no participants actually work on the task. We only calculate the payment of the bonus once all participants in the group have made their decision.

Do you want to volunteer in your team, which consists of exactly N person, and rate the additional 30 comments? [Possible Answers: "Yes" or "No"]

Zusatzaufgabe zur Sicherung der Datenqualität im Team

Vielen Dank für die Bewertung der 1 Kommentare. Sie haben die Hauptaufgabe damit abgeschlossen und erhalten dafür 0,90 €. Insgesamt haben Sie **0,06 Minuten** für die Aufgabe benötigt. Teilnehmer dieser Aufgabe benötigen normalerweise zwischen **0,25 Minuten** und **0,5 Minuten** um 1 Kommentare zu bewerten.

Wir benötigen jetzt **genau einen Freiwilligen** aus Ihrem Team von **genau 300 Personen**.

Anleitung zur Zusatzaufgabe:

Um eine bessere Einschätzung der Kommentarbewertungen in ihrem Team zu ermöglichen und die Datenqualität zu verbessern, müssen 1 weitere Kommentare bewertet werden.

- **Alle** Teammitglieder erhalten jeweils einen Bonus von 0,90 €, wenn eine Person diese zusätzliche Aufgabe absolviert.
- Alle Personen in Ihrem Team bekommen diese Zusatzaufgabe angeboten.
- Es genügt, dass sich **eine Person** in Ihrem Team von **genau 300 Personen** freiwillig meldet. Alle Teammitglieder erhalten dann die Bonuszahlung. Es ist möglich, dass mehr als eine Person die Aufgabe durchführt. Der Bonus steigt dann nicht.
- Sie bekommen den Bonus auch, wenn Sie sich nicht freiwillig melden, aber sich mindestens eine andere Person freiwillig meldet.
- Wenn sich in Ihrem Team niemand freiwillig meldet, dann bekommt auch niemand einen Bonus von 0,90 €.
- Ihre Entscheidung wird keinen Einfluss auf Ihre Bewertung auf der Clickworker Plattform oder auf Ihr bisher verdientes Geld haben.
- Alle Teilnehmer Ihrer Gruppe werden für sich diese Entscheidung treffen, und gegebenenfalls die Aufgabe auch bearbeiten. Es ist also möglich, dass alle Teilnehmer, ein kleinerer Teil, oder kein Teilnehmer die Aufgabe tatsächlich bearbeitet. Die Auszahlung des Bonus wird von uns erst dann berechnet, wenn alle Teilnehmer der Gruppe ihre Entscheidung getroffen haben.

Ihre Entscheidung

Wollen Sie sich in Ihrem Team, das aus **genau 300 Personen** besteht, **freiwillig melden** und die **1 weiteren Kommentare** bewerten?

☐

Weiter

Figure 3.11: Screenshot of the volunteering decision described in Section 3.2.3.

Ihre Entscheidung

Wollen Sie sich in Ihrem Team, das aus **genau 300 Personen** besteht, freiwillig melden und die 1 weiteren Kommentare bewerten? Ihre Entscheidung wird keinen Einfluss auf Ihre Bewertung auf der Clickworker Plattform oder auf Ihr bisher verdientes Geld haben.

Wollen Sie sich in Ihrem Team, das aus genau 300 Personen besteht, freiwillig melden und die 1 weiteren Kommentare bewerten?

Weiter

Zur Erinnerung:

Vielen Dank für die Bewertung der 1 Kommentare. Sie haben die Hauptaufgabe damit abgeschlossen und erhalten dafür 0,90 €. Insgesamt haben Sie **0,07 Minuten** für die Aufgabe benötigt. Teilnehmer dieser Aufgabe benötigen normalerweise zwischen **0,25 Minuten** und **0,5 Minuten** um 1 Kommentare zu bewerten.

Wir benötigen jetzt **genau einen Freiwilligen** aus Ihrem Team von **genau 300 Personen**.

Anleitung zur Zusatzaufgabe:

Um eine bessere Einschätzung der Kommentarbewertungen in ihrem Team zu ermöglichen und die Datenqualität zu verbessern, müssen 1 weitere Kommentare bewertet werden.

- **Alle** Teammitglieder erhalten jeweils einen Bonus von 0,90 €, wenn eine Person diese zusätzliche Aufgabe absolviert.
- Alle Personen in Ihrem Team bekommen diese Zusatzaufgabe angeboten.
- Es genügt, dass sich **eine Person** in Ihrem Team von **genau 300 Personen** freiwillig meldet. Alle Teammitglieder erhalten dann die Bonuszahlung. Es ist möglich, dass mehr als eine Person die Aufgabe durchführt. Der Bonus steigt dann nicht.
- Sie bekommen den Bonus auch, wenn Sie sich nicht freiwillig melden, aber sich mindestens eine andere Person freiwillig meldet.
- Wenn sich in Ihrem Team niemand freiwillig meldet, dann bekommt auch niemand einen Bonus von 0,90 €.
- Ihre Entscheidung wird keinen Einfluss auf Ihre Bewertung auf der Clickworker Plattform oder auf Ihr bisher verdientes Geld haben.
- Alle Teilnehmer Ihrer Gruppe werden für sich diese Entscheidung treffen, und gegebenenfalls die Aufgabe auch bearbeiten. Es ist also möglich, dass alle Teilnehmer, ein kleinerer Teil, oder kein Teilnehmer die Aufgabe tatsächlich bearbeitet. Die Auszahlung des Bonus wird von uns erst dann berechnet, wenn alle Teilnehmer der Gruppe ihre Entscheidung getroffen haben.

Figure 3.12: Screenshot of the volunteering decision if we asked for the belief before the participant took the decision.

3.D.4 Belief elicitation

Your assessment of your team

We are now first interested in your assessment of your team members' decisions. How willing do you think your team members are to volunteer? Please give us your assessment in percent (0-100). The higher the number, the more likely you think it is that at least one other person on your team will volunteer. **You can also receive a further bonus of €0.90 for your estimate. For this purpose, enter the value that actually corresponds to your estimation. The chance to get the bonus**

will depend on how good your estimation is.

How likely is it that **at least one of your team members** from your team, which consists of **3 people**, will volunteer for this task?

[Possible Answers: Slider without initial value between 0 and 100]

Ihre Einschätzung zu Ihrem Team

Wir sind nun zunächst an Ihrer Einschätzung über die Entscheidungen Ihrer Teammitglieder interessiert.

Was glauben Sie, wie hoch ist die Bereitschaft ihrer Teammitglieder sich freiwillig zu melden? Bitte geben Sie hierfür eine Einschätzung in Prozent (0-100) an. Je höher die Zahl, desto wahrscheinlicher ist es Ihrer Einschätzung nach, dass sich mindestens eine andere Person aus Ihrem Team freiwillig meldet. **Auch für Ihre Einschätzung können Sie einen weiteren Bonus von 0,90 € erhalten. Geben Sie dazu den Wert an, der tatsächlich Ihrer Einschätzung entspricht. Die Chance den Bonus zu erhalten richtet sich daran wie gut Ihre Einschätzung ist.**

Wie wahrscheinlich ist es, dass sich von Ihren Teammitgliedern aus Ihrem Team, das aus **30 Personen** besteht **mindestens eine Person** für diese Aufgabe meldet?

Sicher keine
andere Person
(0 %)



Sicher
mindestens eine
andere Person
(100 %)

Bitte klicken Sie auf den Schieberegler, um die Wahrscheinlichkeit anzugeben. Sie können Ihre Entscheidung im Anschluss noch verändern.

Weiter

Zur Erinnerung:

Vielen Dank für die Bewertung der 1 Kommentare. Sie haben die Hauptaufgabe damit abgeschlossen und erhalten dafür 0,90 €. Insgesamt haben Sie **0,07 Minuten** für die Aufgabe benötigt. Teilnehmer dieser Aufgabe benötigen normalerweise zwischen **0,25 Minuten** und **0,5 Minuten** um 1 Kommentare zu bewerten.

Wir benötigen jetzt **genau einen Freiwilligen** aus Ihrem Team von **genau 30 Personen**.

Anleitung zur Zusatzaufgabe:

Um eine bessere Einschätzung der Kommentarbewertungen in ihrem Team zu ermöglichen und die Datenqualität zu verbessern, müssen 1 weitere Kommentare bewertet werden.

- **Alle** Teammitglieder erhalten jeweils einen Bonus von 0,90 €, wenn eine Person diese zusätzliche Aufgabe absolviert.
- Alle Personen in Ihrem Team bekommen diese Zusatzaufgabe angeboten.
- Es genügt, dass sich **eine Person** in Ihrem Team von **genau 30 Personen** freiwillig meldet. Alle Teammitglieder erhalten dann die Bonuszahlung. Es ist möglich, dass mehr als eine Person die Aufgabe durchführt. Der Bonus steigt dann nicht.
- Sie bekommen den Bonus auch, wenn Sie sich nicht freiwillig melden, aber sich mindestens eine andere Person freiwillig meldet.
- Wenn sich in Ihrem Team niemand freiwillig meldet, dann bekommt auch niemand einen Bonus von 0,90 €.
- Ihre Entscheidung wird keinen Einfluss auf Ihre Bewertung auf der Clickworker Plattform oder auf Ihr bisher verdientes Geld haben.
- Alle Teilnehmer Ihrer Gruppe werden für sich diese Entscheidung treffen, und gegebenenfalls die Aufgabe auch bearbeiten. Es ist also möglich, dass alle Teilnehmer, ein kleinerer Teil, oder kein Teilnehmer die Aufgabe tatsächlich bearbeitet. Die Auszahlung des Bonus wird von uns erst dann berechnet, wenn alle Teilnehmer der Gruppe ihre Entscheidung getroffen haben.

Figure 3.13: Screenshot of the belief elicitation.

Zusatzaufgabe zur Sicherung der Datenqualität im Team

Vielen Dank für die Bewertung der 1 Kommentare. Sie haben die Hauptaufgabe damit abgeschlossen und erhalten dafür 0,90 €. Insgesamt haben Sie **0,07 Minuten** für die Aufgabe benötigt. Teilnehmer dieser Aufgabe benötigen normalerweise zwischen **0,25 Minuten** und **0,5 Minuten** um 1 Kommentare zu bewerten.

Wir benötigen jetzt **genau einen Freiwilligen** aus Ihrem Team von **genau 300 Personen**.

Anleitung zur Zusatzaufgabe:

Um eine bessere Einschätzung der Kommentarbewertungen in ihrem Team zu ermöglichen und die Datenqualität zu verbessern, müssen 1 weitere Kommentare bewertet werden.

- **Alle** Teammitglieder erhalten jeweils einen Bonus von 0,90 €, wenn eine Person diese zusätzliche Aufgabe absolviert.
- Alle Personen in Ihrem Team bekommen diese Zusatzaufgabe angeboten.
- Es genügt, dass sich **eine Person** in Ihrem Team von **genau 300 Personen** freiwillig meldet. Alle Teammitglieder erhalten dann die Bonuszahlung. Es ist möglich, dass mehr als eine Person die Aufgabe durchführt. Der Bonus steigt dann nicht.
- Sie bekommen den Bonus auch, wenn Sie sich nicht freiwillig melden, aber sich mindestens eine andere Person freiwillig meldet.
- Wenn sich in Ihrem Team niemand freiwillig meldet, dann bekommt auch niemand einen Bonus von 0,90 €.
- Ihre Entscheidung wird keinen Einfluss auf Ihre Bewertung auf der Clickworker Plattform oder auf Ihr bisher verdientes Geld haben.
- Alle Teilnehmer Ihrer Gruppe werden für sich diese Entscheidung treffen, und gegebenenfalls die Aufgabe auch bearbeiten. Es ist also möglich, dass alle Teilnehmer, ein kleinerer Teil, oder kein Teilnehmer die Aufgabe tatsächlich bearbeitet. Die Auszahlung des Bonus wird von uns erst dann berechnet, wenn alle Teilnehmer der Gruppe ihre Entscheidung getroffen haben.

Ihre Einschätzung zu Ihrem Team

Wir sind nun zunächst an Ihrer Einschätzung über die Entscheidungen Ihrer Teammitglieder interessiert.

Was glauben Sie, wie hoch ist die Bereitschaft ihrer Teammitglieder sich freiwillig zu melden? Bitte geben Sie hierfür eine Einschätzung in Prozent (0-100) an. Je höher die Zahl, desto wahrscheinlicher ist es Ihrer Einschätzung nach, dass sich mindestens eine andere Person aus Ihrem Team freiwillig meldet. **Auch für Ihre Einschätzung können Sie einen weiteren Bonus von 0,90 € erhalten. Geben Sie dazu den Wert an, der tatsächlich Ihrer Einschätzung entspricht. Die Chance den Bonus zu erhalten richtet sich daran wie gut Ihre Einschätzung ist.**

Wie wahrscheinlich ist es, dass sich von Ihren Teammitgliedern aus Ihrem Team, das aus **300 Personen** besteht **mindestens eine Person** für diese Aufgabe meldet?

Sicher keine
andere Person
(0 %)



Sicher
mindestens eine
andere Person
(100 %)

Bitte klicken Sie auf den Schieberegler, um die Wahrscheinlichkeit anzugeben. Sie können Ihre Entscheidung im Anschluss noch verändern.

Weiter

Figure 3.14: Screenshot of the belief elicitation if we asked for the belief before the participant took the decision.

3.D.5 Questionnaire items

- **RISK:** “Are you a person who is generally willing to take risks, or do you try to avoid taking risks?”; Scale: 0 = “completely unwilling to take risks”; 10 = “very willing to take risks”.
- **TIME PREF.:** “In comparison to others, are you a person who is generally willing to give up something today in order to benefit from that in the future, or are you not willing to do so?”; Scale: 0 = “completely unwilling to give up something today”; 10 = “very willing to give up something today”.
- **TRUST:** “As long as I am not convinced otherwise, I assume that people have only the best intentions.”; Scale: 0 = “does not describe me at all”; 10 = “describes me perfectly”.
- **NEGATIVE RECIPROCITY:** “Are you a person who is generally willing to punish unfair behavior even if this is costly?”; Scale: 0 = “not willing at all to incur costs to punish unfair behavior”; 10 = “very willing to incur costs to punish unfair behavior”.
- **ALTRUISM:** “Imagine the following situation: you won 1,000 € in a lottery. Considering your current situation, how much would you donate to charity?”; Values between 0 and 1000 are allowed.
- **EFFICIENCY:** “I am more willing to make an effort if many profit from it.”; Scale: 0 = “does not describe me at all”; 10 = “describes me perfectly”.

Declaration of Contribution

Hereby I, Tobias Felix Werner, declare that the chapter “Willingness to volunteer among remote workers is insensitive to the team size” is coauthored by Adrian Hillenbrand and Fabian Winter. All authors contributed equally to the chapter.

Signature of coauthor (Adrian Hillenbrand): *Adrian Hillenbrand*

Prof. Dr. Adrian Hillenbrand

A handwritten signature in black ink, appearing to read 'Fabian Winter', enclosed within a light gray rectangular box.

Signature of coauthor (Fabian Winter): _____

Dr. Fabian Winter

Chapter 4

What Drives Demand for Loot Boxes?

Co-authored with Simon Cordes and Markus Dertwinkel-Kalt

4.1 Introduction

The market for (online) video games is booming in recent years, with in-game purchases accounting for a substantial share of developers’ revenues. In 2020 alone, so-called “loot boxes” generated \$15 billion of worldwide revenue, and projections suggest that 230 million people will spend money on loot boxes by 2025.¹ Loot boxes are digital lotteries in video games that — similar to gambling — offer *random* rewards to be used in-game. While loot boxes share a lot of similarities with gambling (Drummond and Sauer, 2018), surprisingly little regulation is in place that would restrict their design and the way they are priced.² At the same time, policymakers and the general public alike are concerned that loot boxes induce consumers — in particular, those susceptible to gambling — to overspend on video games.³

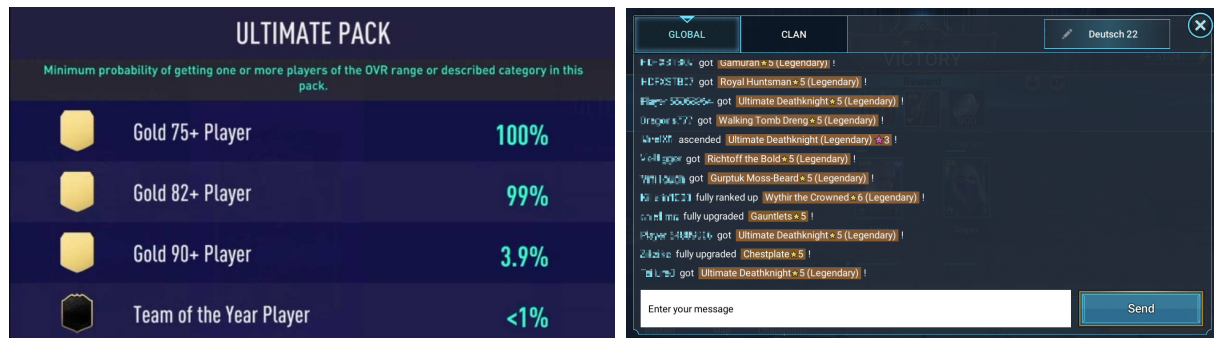


Figure 4.1: Typical design features, censored odds (left) and selected feedback (right), of loot boxes.

This concern is amplified by the fact that loot boxes are designed, and marketed, in ways that obfuscate the chances of winning different rewards. First, the odds are often censored. As a specific example, consider the football simulation *FIFA Ultimate Team*, where gamers build a team of players that vary in strength. Gamers can buy packs that offer lotteries over players. The odds, however, are provided, if at all, only for a coarse set of intervals, bunching together players of very different strength (see the left panel of Figure 1). At the extreme, the worst player in an interval is around 1000 times less

¹See, for instance, [https://www.juniperresearch.com/press/video-game-loot-boxes-to-generate-over-\\$20-billion](https://www.juniperresearch.com/press/video-game-loot-boxes-to-generate-over-$20-billion) (accessed on September 16th, 2022).

²Recently, 20 consumer organizations from 18 European countries suggest that loot boxes should be classified as gambling and therefore regulated (The Norwegian Consumer Counsel, 2022). Additionally, the Federal Trade Commission (FTC) is investigating loot boxes following concerns from U.S. legislators that they may be similar to gambling (Federal Trade Commission, 2020).

³See, for instance, <https://www.theguardian.com/society/2021/apr/02/video-game-loot-boxes-problem-gambling-betting-children> (accessed on September 16th, 2022).

valuable than the best player. The Norwegian Consumer Counsel (2022) argues that gamers, therefore, overestimate the value of these lotteries. Second, gamers often receive highly selected feedback on the rewards other gamers have obtained. In the mobile game *Raid: Shadow Legends*,⁴ for example, gamers receive a notification whenever another player wins a rare reward (see the right panel of Figure 1). As only rare rewards are reported, this provides them with a biased sample of the reward distribution. Going further, game developers not only pay content providers (on *Youtube* or *Twitch*) to open loot boxes on their shows, but they allegedly also offer them better odds.⁵ According to The Norwegian Consumer Counsel (2022, p.44), observing such a biased sample of the reward distribution “reinforces the player’s belief that they might be similarly lucky.” While one could easily imagine that both features contribute to overspending on loot boxes, there is a lack of systematic evidence supporting this claim.

In this chapter, we experimentally investigate what drives the willingness-to-pay for loot boxes. We focus on the effects of censoring the odds and providing gamers with a selected sample of the reward distribution. We do so because these design features hardly provide any utility to gamers. Their sole purpose seems to be making consumers overspend on loot boxes. Indeed, we find evidence that both features increase the willingness-to-pay for lotteries. They do so through different channels, however. Censoring the odds of a lottery increases a subject’s willingness-to-pay via inflating her belief of winning a high reward. Simply providing subjects with a selected sample of the reward distribution, on the other hand, has no statistically significant effect on beliefs. Still, showing subjects such a selected sample significantly increases the average willingness-to-pay for lotteries. Going beyond existing work (e.g., Barron et al., 2019), this suggests that selected feedback affects economic behavior not solely via a beliefs channel. When combined with censored odds, however, a selected sample does significantly increase beliefs and willingness-to-pay. Overall, our results suggest that the design of loot boxes — combining censored odds with selected feedback — contributes to overspending in the market for video games, and thus support a case for regulating loot boxes.

⁴In 2022, three years after its release, the game surpassed \$1bn in lifetime revenue. For details, see <https://gamingonphone.com/news/raid-shadow-legends-surpasses-1-billion-in-lifetime-revenue/> (accessed on January 2nd, 2023).

⁵See, for instance, <https://gamerant.com/ftc-loot-boxes-better-odds-sponsored-streamers/> (accessed on December 22nd, 2022).

We introduce our experimental design in Section 4.2. Subjects repeatedly state their willingness-to-pay (WTP) for different monetary lotteries with three potential prizes, one of which is zero. In a *Control* condition, we transparently describe the odds of the lotteries and do not provide any additional information to the subjects. We assume that this control condition identifies a subject’s true WTP, and define overspending relative to this benchmark. We implement three treatments that capture the features of loot boxes discussed above. In *Censored*, subjects only learn the total probability of winning a non-zero prize, but not the exact probability of winning the highest prize. In *Sample*, we provide subjects with the full prize distribution as well as a selected sample thereof; that is, they observe the five highest outcomes in a sample of 400 draws. Finally, *Joint* combines both: subjects observe the censored prize distribution *and* a selected sample thereof. This last treatment resembles the current design of loot boxes most closely.

At the same time, our experimental design eliminates all features of loot boxes that may provide utility beyond winning a reward, such as a nice design or visual effects. Instead, we isolate the features of loot boxes that almost certainly do not affect a gamer’s material utility, and can thus be interpreted as inducing mistakes. Under the weak assumption that mistakes become weakly stronger as material utility goes up, we then identify a lower bound on overspending (relative to the control condition).

Section 4.3 presents our main results. Compared to the *Control* condition, the average WTP increases by roughly 45% in *Censored* and *Sample*, respectively. The effect in *Censored* is driven by the large and statistically significant increase in the perceived likelihood of winning the highest prize. Once we control for stated beliefs (i.e., the mediator), the average WTP in *Censored* does not differ significantly from that in *Control* anymore. Selected feedback, in contrast, has no statistically significant effect on stated beliefs. In fact, the average beliefs are significantly lower in *Sample* than in *Censored*. Because the average WTPs in *Sample* and *Censored* are almost identical, however, this suggests that selected feedback at least partly operates through a channel different from beliefs. When combined with censored odds, selected feedback does significantly increase both the average WTP *and* beliefs. Compared to the control condition, the average WTP doubles in *Joint*.

In Section 4.4, we provide additional evidence on the underlying mechanism as well as for the relevance of our results. First, we restrict the sample to decisions for which stated

beliefs are “realistic” in that they are consistent with the provided information. Here, we find a precisely estimated zero difference in average beliefs between *Censored* and *Joint*. In either case, subjects tend to assign equal probabilities to the non-zero prizes. This is consistent with evidence on people naively applying a “50-50 heuristic” when being uncertain about the problem they face (e.g., Sonnemann et al., 2013; Enke and Graeber, 2019).

Second, we ran a robustness experiment to address the concern that the lotteries are offered by the experimenter, not a firm, trying to maximize profits. The treatment *Info* replicates *Joint* but adds unbiased information on the reward distribution and explicitly tells the subject that the odds are not 50-50. While the additional information makes subjects less optimistic about winning the highest prize, it does not affect the average WTP. Our findings highlight the robustness of our main treatments and show that even in the presence of further unbiased information, the design features of loot boxes promote overspending.

Third, we study correlates of a survey measure on loot-box overspending. Consistent with the prior literature (e.g., Zendle and Cairns, 2018), survey measures of gambling behavior correlate with overspending on loot boxes. Controlling for these measures of gambling behavior, we still find a positive association between the average WTP for the lotteries in our experiment and survey measures of overspending on loot boxes. Hence, our experimental measure of gambling behavior picks up part of the unexplained variation in loot-box overspending.

We conclude in Section 4.5 by discussing tentative policy implications and challenges in their implementation. Current plans for regulation in Germany include labels for games with loot boxes.⁶ Because it seems to be common design features of loot boxes — not the fact that they offer random rewards — that induce overspending, this regulation is unlikely to be sufficient. Consistent with this observation, firms would probably not invest much effort into obfuscating the odds of loot boxes if it had no effect. Our experimental results thus suggest a case for stronger regulation that restricts the design of loot boxes.⁷ Still, the effects of such regulation are far from obvious. For example, gamers might be overwhelmed by learning the full distribution over the players’ levels in *FIFA Ultimate Team*, and might

⁶See, for instance, <https://usk.de/jugendschutzgesetz-aktualisiert-usk-bereitet-sich-auf-aenderungen-vor/> (accessed on September 19th, 2022).

⁷Alternatively, regulators could ban loot boxes altogether. As the case of Belgium shows, however, such a ban can only work if regulators also introduce proper enforcement mechanisms (Xiao, 2022).

simply ignore it. Moreover, preventing game developers from providing selected feedback does not imply that gamers receive representative feedback on the reward distribution. They might get similarly selected feedback on the reward distribution from talking to others or from watching live streams of professional gamers.

Related Literature First, we contribute to the literature on gaming, specifically loot boxes. A series of papers has established a positive correlation between survey measures of gambling and overspending on loot boxes (Drummond and Sauer, 2018; Zendle and Cairns, 2018; Drummond et al., 2020). We strengthen this link by providing causal evidence on how key design features of loot boxes affect the WTP for lotteries, identifying a lower bound on how these features affect demand for actual loot boxes. Chen et al. (2021) develop a model of optimal loot box pricing, assuming that gamers maximize expected utility (EU). We find, however, that common features of loot boxes result in WTPs that are non-linear in stated beliefs, clearly rejecting the EU hypothesis.

We also connect to the behavioral literature on choice under risk and ambiguity. Specifically, the fact that game developers introduce “ambiguity” by censoring the odds of loot boxes is consistent with recent evidence on ambiguity-seeking behavior (see Chandrasekher et al., 2022, and the literature therein). Such behavior is also consistent with the finding that the censoring of odds increases the average WTP for lotteries in our experiment.

Finally, we add to the literature on biased inferences from (non-)disclosed data. Empirical evidence from the lab and field suggests that individuals often draw wrong inferences from selectively disclosed data in strategic settings (Bolton et al., 2007; Koehler and Mercer, 2009; Brown et al., 2012; Benndorf et al., 2015; Deversi et al., 2021; Jin et al., 2021, 2022) and non-strategic ones (Esponda and Vespa, 2018; Barron et al., 2019; Enke, 2020; López-Pérez et al., 2022). We find that subjects naively bias censored probabilities towards a uniform distribution. By closely resembling common features of loot boxes, our design allows us to provide more nuanced insights to the question of how to regulate their design.

4.2 Experimental design

4.2.1 Experimental setup

We develop an experimental design that allows us to test for the effects of two key features of loot boxes on the willingness-to-pay (WTP) for lotteries. First, the odds of loot boxes are often censored. Second, gamers typically receive (positively) selected feedback on the reward distribution. We implement three treatments and one control condition to identify the effect of each feature in isolation — i.e., the effect of censored odds and selected feedback — as well as how both features interact with each other. Our design abstracts from other features of loot boxes (such as visual effects) that might provide utility to subjects. As we argue in Section 4.2.2, under a weak assumption, we thereby identify a lower bound on the actual effects on the demand for loot boxes.

All subjects sequentially state their WTP for five lotteries. Each lottery pays a non-zero prize with probability $q\%$ and zero otherwise. The non-zero prize is either 10 Coins or x Coins. Both probabilities and prizes vary across decisions: in each decision, we independently draw (without replacement) a probability $q \in \{10, 20, 30, 40, 50\}$ and a prize $x \in \{100, 120, 140, 160, 180\}$. Probability and prize pairs are drawn at the subject level, so that different subjects may observe different lotteries in a different order. The high prize of x Coins is always realized with probability 1%.⁸ Before stating their WTP for a lottery, we ask subjects to state their belief on how often they would win this high prize in 100 draws (see Figure 4.2 for a screenshot and Section 4.2.2 for an interpretation).

Across four conditions, we vary the amount of information that subjects receive on the lotteries. We next describe these four conditions in more detail. Screenshots of the instructions and decision screens can be found in Appendix 4.C.

Control In the control condition, subjects learn the full probability distribution of the different lotteries. More specifically, as illustrated in Figure 4.2, they learn the exact probability of winning 10 Coins and x Coins, respectively. Subjects do not get any additional information on the lotteries.

⁸For example, if $q = 10\%$ and $x = 100$, the lottery pays 0 Coins with probability 90%, 10 Coins with probability 9%, or 100 Coins with probability 1%.

Lottery			
Probability	90 %	9 %	1 %
Payoff	0 coins	10 coins	100 coins

Imagine you would play the lottery **100 times**. How often do you think would you win **100 coins**?

0 times 100 times

I believe that I would win 100 coins exactly **1 time** if I would play the lottery 100 times.

How much are you willing to pay for this lottery?

Coins

Figure 4.2: Decision screen in the Control condition after stating the belief.

Censored In the treatment *Censored*, subjects observe censored versions of the lotteries. As illustrated in Figure 4.3, they only learn the probability of receiving a non-zero prize, but not the exact probabilities of receiving 10 Coins or $x = 100$ Coins, respectively. This mimics the censoring strategies of video game designers such as *EA Sports* who do not provide gamers with the full probability distribution (see the left panel of Figure 4.1). Other than in *Control*, to assess the value of a lottery, subjects have to form a belief about the probability of receiving $x = 100$ Coins. Based on existing research (Fischhoff and Bruine De Bruin, 1999; Sonnemann et al., 2013; Enke and Graeber, 2019), we expect subjects to overestimate this probability, likely biasing it towards 50%.

Lottery		
Probability	90 %	10 %
Payoff	0 coins	either 10 coins or 100 coins

Figure 4.3: Information provided in the treatment Censored.

Sample In the treatment *Sample*, subjects again learn the full probability distribution of each lottery, but on top, observe a sample from this distribution. As illustrated in Figure 4.4, we present subjects the 5 highest outcomes in a sample of 400 actual draws from the lottery. Notably, subjects receive transparent information on how the outcomes are chosen. This treatment is motivated by the common practice of announcing rare prizes other players have obtained (see the right panel of Figure 4.1). Importantly, because subjects observe the full probability distribution, the sample does not contain any new information regarding the value of the lottery. Still, existing research (e.g., Barron et al.,

2019) suggests that observing a series of high draws from a distribution may increase a subject's WTP.

Lottery			
Probability	90 %	9 %	1 %
Payoff	0 coins	10 coins	100 coins

Lottery draws	
#	Value
Draw 1	100 Coins
Draw 2	100 Coins
Draw 3	100 Coins
Draw 4	100 Coins
Draw 5	100 Coins

Figure 4.4: Information provided in the treatment Sample.

Joint The treatment *Joint* combines both of the above: subjects observe a censored version of the lotteries together with the 5 highest outcomes in a sample of 400 draws. Unlike in *Sample*, the sample does contain information about the underlying probability distribution in this case. With censored odds, subjects arguably overestimate the low probability of winning x Coins initially. Then, if all subjects were Bayesian, the average belief upon observing the sample should decrease, moving closer to the truth. If, in contrast, subjects naively infer from a series of good draws that the lottery has to be even better than they initially thought, the average belief upon observing the sample should go up. Hence, compared to *Censored*, also the average WTP should increase.

4.2.2 Conceptual framework

We sketch a simple model to clarify how we think about our experimental design, and how it links to loot boxes typically offered. Consider a lottery Z with a distribution G^* over prizes (such as player cards in *FIFA Ultimate Team*). The lottery is presented in a “frame” f (such as our different treatments), which captures the description of the odds or the feedback provided to gamers. We assume that gamers form a subjective belief G_f over the prize distribution that depends on the frame. Going beyond our experimental design, we allow for features of loot boxes like a nice design or fancy visual effects that directly

provide utility to gamers. We summarize such additional features in the “context” c , and refer to the “neutral context” of our experiment as c_0 .

Willingness-to-Pay We assume that gamers aim to maximize their subjective expected utility. Denote as $u(z, c)$ the utility derived from prize z in context c . A gamer’s WTP for the loot box Z under frame f is then given by

$$\mathbb{E}_{G_f}[u(Z, c)] = \mathbb{E}_{G^*}[u(Z, c)] + \underbrace{\mathbb{E}_{G_f}[u(Z, c)] - \mathbb{E}_{G^*}[u(Z, c)]}_{=:\phi(c, f)},$$

where $\phi(c, f)$ captures a bias that operates through the gamer’s subjective belief.

We make two assumptions to link the above to our experimental design.

Assumption 1. *For any context c and any frame f , $\phi(c, f) \geq \phi(c_0, f)$.*

Our first assumption says that adding contextual features that provide utility induces a weakly larger bias. For example, if gamers get distracted by fancy visual effects, they might become more prone to make statistical errors. A nice design of loot boxes might also result in a more favorable view of the game developer and the odds it offers. Under Assumption 1, our design, which abstracts from features of loot boxes that directly generate utility, allows us to estimate a lower bound on the bias in loot-box demand.

Assumption 2. *The control condition eliminates any bias in the WTP.*

Under Assumption 2, our control condition identifies the average consumption value of a lottery. Moreover, a simple linear regression of the stated WTPs on treatment indicators identifies the average overspending on these lotteries due to censored odds and selected feedback. Under Assumption 1, the same regression provides a lower bound on the average overspending on loot boxes induced by these common design features.

Beliefs To test whether the any bias in WTPs operates via beliefs, we ask subjects how often they believe to win the high prize of x Coins in 100 draws. Denote as b_i the belief of subject i for a lottery Z (under frame f). We think of this belief as follows:

$$b_i = \mathbb{P}_{G_f}[Z = x] + \epsilon_i,$$

where the “noise” term includes implementation errors or general optimism.

Under the assumption that this noise is independent of the frame, a simple linear regression of stated beliefs on treatment indicators identifies the bias in beliefs induced by censored odds and selected feedback on the reward distribution. To make beliefs comparable across lotteries, we normalize stated beliefs by the probability of winning a non-zero prize.⁹ Overly “optimistic” subjects might, however, state beliefs that contradict the objective information they have, meaning that their normalized belief exceeds 100%. We will also provide analyses separately for those subjects who state realistic beliefs.

4.2.3 Implementation and logistics

As is common practice for WTP elicitation, we use the BDM mechanism to incentivize subjects (Becker et al., 1964). We do not incentivize the belief elicitation, however. Recent work by Danz et al. (2022) suggests that standard incentivization techniques (such as the binarized scoring rule) systematically distort reported beliefs. Moreover, we view the belief question as an input to (or mediator of) a subject’s stated WTP, which is our primary outcome of interest. To minimize anchoring effects, we use a slider without an initial value to elicit beliefs (see Figure 4.14). To ensure that subjects engage with the lotteries, they could state beliefs (and afterwards WTPs) only after a 5 second delay.

The design was pre-registered in the AEA RCT registry as trial AEARCTR-0009501.¹⁰ We collected data from 617 subjects located in the United Kingdom (UK) via *Prolific* in July 2022. The experiment consisted of 3 modules. First, we screened out inattentive participants via an attention check at the beginning of the experiment and after the instructions via comprehension questions. Second, all subjects who passed both tests stated their WTPs and beliefs for five lotteries. Third, we collected demographics and potential correlates of interest. Screenshots of all parts of the experiment (including additional survey questions) can be found in Appendix B. Subjects earned a base fee of 1.50 GBP for participation. In addition, 1 out of 6 participants received a bonus payment depending on the WTP stated for one randomly selected lottery. Conditional on receiving a bonus, the average bonus paid was 5.35 GBP. The experiment took, on average, 11 minutes to complete.

⁹Consider the lottery that pays 0 Coins with 90% probability, 10 Coins with 9% probability, and 100 Coins with 1% probability. If a subject believes to win the high prize of 100 Coins 5% of the time, the corresponding normalized belief is $\frac{5\%}{9\%+1\%} = 50\%$.

¹⁰The pre-registration can be found at <https://doi.org/10.1257/rct.9501-3.0>.

4.3 Main results

Our main results are summarized in Table 4.1 and Figure 4.5. Relative to *Control*, the average WTP for the lotteries significantly increases by 43% in *Sample*, by 45% in *Censored*, and by 100% in *Joint*. Common design features of loot boxes, therefore, induce substantial overspending in the context of our experiment.

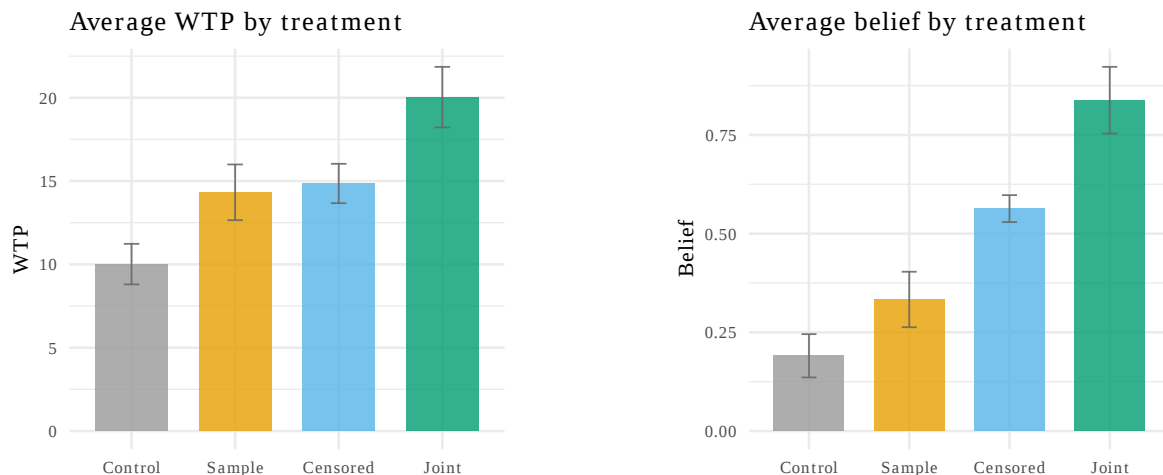


Figure 4.5: Average willingness-to-pay and beliefs by treatment. We include all subjects that finished the experiment. WTP is the willingness to pay for a lottery. Belief is the belief that the high outcome occurs conditionally on the medium or high outcome being drawn. Whiskers are the standard error of the mean.

Consistent with our conceptual framework, the effect of censoring operates through the belief of winning the high prize. To make this point precisely, we transform stated beliefs into subjective probabilities of winning the high prize *conditional* on winning a non-zero prize (see Section 4.2.2). Compared to *Control*, this average conditional belief of winning the high prize significantly increases by 37 p.p. in *Censored*. Along these lines, when controlling for beliefs (see Column (3) of Table 4.1), the effect of censoring on the average WTP is no longer significant. This cleanly demonstrates that censoring operates via a belief channel.¹¹ Providing subjects with a biased sample of the reward distribution has a weaker effect on beliefs, and the effect is only significant at the 10% level. The fact that the increase in average WTP (relative to *Control*) is similar in *Sample* and *Censored* thus suggests that selected feedback affects overspending (at least partly) through a non-belief channel. When controlling for beliefs, the effect of selected feedback on WTPs remains

¹¹Under the null of our conceptual framework, beliefs fully mediate the treatment effects on WTP. Including the potential mediator into our regression allows us to test this hypothesis of full mediation.

weakly significant (see Column (3) of Table 4.1). When combining it with censored odds, however, selected feedback does increase average beliefs by an additional 28 p.p. compared to *Censored*. In line with our conceptual framework, this may explain part, but not all of the increase in the average WTP when moving from *Censored* to *Joint*. Furthermore, when controlling for stated beliefs, the average WTP significantly increases in *Joint* compared to *Control*. It suggests again that selected feedback works at least partly through a non-belief channel. While our experimental design does not allow us to identify this channel, it is an interesting path for future research to investigate it further.

Table 4.1: Regression results — main specification

	WTP			Belief	
	No controls (1)	Controls (2)	Controls + Belief (3)	No controls (4)	Controls (5)
Censored	4.84*** (1.55)	4.83*** (1.53)	2.53 (1.56)	0.373*** (0.055)	0.373*** (0.055)
Sample	4.31** (2.00)	4.17** (1.99)	3.32* (1.83)	0.143* (0.074)	0.138* (0.074)
Joint	10.0*** (1.99)	10.0*** (2.02)	6.04*** (2.05)	0.647*** (0.079)	0.650*** (0.080)
Belief			6.17*** (1.10)		
Observations	3,085	3,085	3,085	3,085	3,085

Notes: Results from ordinary least squares (OLS) regressions on treatment dummies. The outcome variable in columns (1) and (2) is the willingness to pay and in (3) and (4) the transformed belief, i.e. the belief that the high outcome occurs conditional on the medium or high outcome being drawn. Columns (1) and (3) do not include control variables. Columns (2) and (4) control for age, gender and monthly available budget. Standard errors clustered at the subject level in parentheses. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

4.4 Additional results

4.4.1 Treatment effects under realistic beliefs

Next, we restrict our sample to decisions in which subjects stated “realistic” beliefs that could be interpreted as conditional probabilities consistent with the information they observed. Consider, for example, a lottery that pays 10 Coins with 39% probability and 100 Coins with 1% probability. If subjects observe this distribution (in *Control* and

Sample), the only realistic belief is exactly 1%. In *Control*, subjects stated the realistic belief in 85% of the decisions, while in *Sample* they stated it 62% of the time. In treatments *Censored* and *Joint* subjects only learn the probability with which the lottery pays a non-zero prize; here, 40%. Hence, any belief between 0% and 40% is realistic in this case. This leaves us with 97% of the decisions in *Censored* and 86% of the decisions in *Joint*.

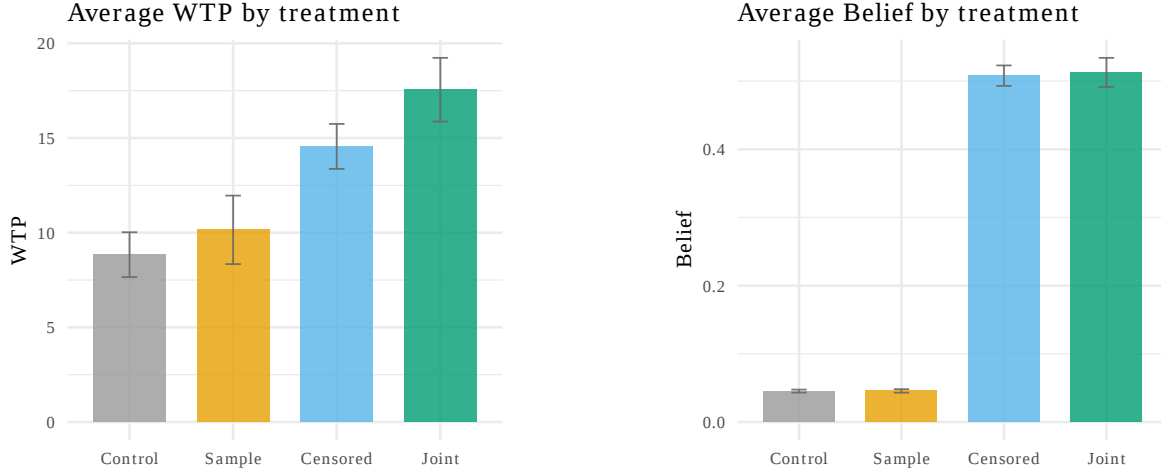


Figure 4.6: Average willingness-to-pay and beliefs by treatment in the main experiment. We include all subjects that finished the experiment. We exclude all decisions in which a subject stated a Belief of larger than 1. WTP is the willingness to pay for a lottery. Belief is the belief that the high outcome occurs conditionally on the medium or high outcome being drawn. Whiskers are the standard error of the mean.

Figure 4.6 and Table 4.3 show the results. By construction of the sample, stated beliefs are identical in *Control* and *Sample*. More interestingly, conditional on stating a realistic belief, there is also no significant difference in stated beliefs across *Censored* and *Joint*. At the same time, even realistic beliefs are massively inflated (compared to the truth) in these treatments. On average, subjects assign almost equal probabilities to the events of receiving 10 Coins and x Coins. This is consistent with existing evidence on people — when being uncertain — biasing probabilities towards a uniform distribution (e.g., Fischhoff and Bruine De Bruin, 1999; Sonnemann et al., 2013; Enke and Graeber, 2019). Across the four conditions, the average WTPs follow the same qualitative patterns as before. However, neither the difference in WTP between *Control* and *Sample* nor the difference in WTP between *Censored* and *Joint* is statistically significant in this subsample.

4.4.2 Robustness experiment

One caveat of our design is that the lotteries are offered by the experimenter, not a firm trying to maximize profits. This might affect the inferences that subjects draw from observing censored probabilities or a selected sample, and it might result in higher WTPs compared to a market setting. We address this concern in a second experiment.¹²

In the treatment *Info*, we make it clear to subjects that the outcomes with censored probabilities are *not* equally likely, and we further provide them unbiased information on the probability distribution (on top of a selected sample). A total number of 414 subjects completed the experiment on *Prolific* in November 2022. The instructions, screenshots of the decision screens, and details on the implementation can be found in Appendix 4.B.

Figure 4.7 summarizes our findings. The additional information significantly decreases the average (conditional) belief of winning the high prize by 37 p.p. compared to *Joint* ($p < 0.01$). While the WTP in *Info* is also slightly below the one in *Joint*, the effect is not significant at the 10% level ($p = 0.69$). Hence, the shift in beliefs is insufficient to influence subjects' WTP to the same degree. Importantly, both the belief and the WTP are significantly higher in *Info* compared to *Control* (see Table 4.4 in Appendix 4.A). Overall, it highlights the robustness of our results and points to the significance of the design features of loot boxes on WTPs.

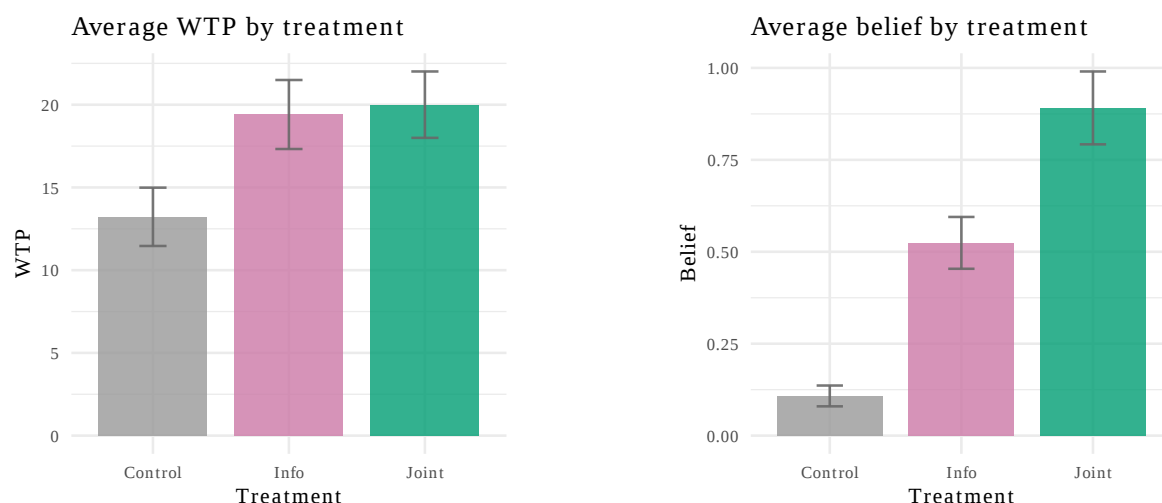


Figure 4.7: Average willingness-to-pay and beliefs by treatment in the Robustness Experiment. We include all subjects that finished the Robustness Experiment. WTP is the willingness to pay for a lottery. Belief is the belief that the high outcome occurs conditionally on the medium or high outcome being drawn. Whiskers are the standard error of the mean.

¹²The pre-registration is available at <https://doi.org/10.1257/rct.10506-1.0>.

4.4.3 Correlates of overspending on loot boxes

Using survey measures of loot-box demand, we study correlates of loot-box overspending. At the end of the experiment, we asked subjects a series of questions on their usage of video games and knowledge of loot boxes. Our subjects spend, on average, 1.2 hours playing video games daily. Moreover, 69% of our subjects know what loot boxes are, 17% state that they spend a positive amount of money on loot boxes every month, and 11% state that they have ever spent more on loot boxes than they initially planned to. Consistent with the prior literature (Zendle and Cairns, 2018), we find a positive correlation ($\rho = 0.26$, $p < 0.01$) between loot-box usage and survey measures of gambling behavior.¹³ Overall, our sample thus seems well-suited to study the demand for loot boxes.

Table 4.2: Regression results — predictors for real world overspending

	=1 if subject overspent on loot boxes			
	(1)	(2)	(3)	(4)
WTP	0.004*** (0.001)			0.003** (0.001)
Belief		0.044 (0.033)		-0.002 (0.034)
Gambling Score			0.869*** (0.271)	0.825*** (0.260)
Observations	425	425	425	425

Notes: Results from ordinary least squares (OLS) regressions of a dummy that is one if the subject state to have spent more than planned on loot boxes in the real world. WTP is the willingness to pay. Belief is the belief that the high outcome occurs conditional on the medium or high outcome being drawn. Gambling score is the score from a self reported gambling questionnaire, scaled from 0 to 1. All variables are subject averages. We include subjects who have stated that to know what loot boxes are. Robust standard errors in parentheses. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

We are mainly interested in whether subjects who have overspent on loot boxes in the past differ in systematic ways from those who did not and whether our experimental measures pick up part of this variation. We asked subjects: “Have you ever spent more than you planned to on loot boxes?” We then regress a dummy variable that takes a value of one if the answer was yes and a value of zero otherwise on a subject’s average WTP and stated beliefs (see Table 4.2). We find a positive association between a subject’s WTP and overspending on loot boxes. An increase in the average WTP from the 5th

¹³We define subjects as loot-box users if they have either (a) ever spent more than they planned to on loot boxes or (b) have a positive monthly spending on loot boxes.

percentile to the median is associated with an increase in the probability to overspend of 3.28 percentage points. While admittedly small, the effect remains significant even when controlling for survey measures of gambling behavior. Hence, the average WTP for monetary lotteries picks up part of the variation in loot-box overspending that these survey measures cannot explain. On the other hand, stated beliefs are not correlated with the tendency to overspend on loot boxes.

4.5 Conclusion

We document experimentally that censored odds and selected feedback increase the demand for lotteries, which gives lower bounds on the effects on actual loot-box demand. Both features hardly provide utility to gamers, so our results support a case for regulating loot-box design. At the same time, it is not obvious what regulation would be effective.

There are currently plans to label games with loot boxes in Germany.¹⁴ However, our results show that the design of loot boxes, rather than the random rewards they provide, encourages players to overspend. Hence, this regulation may not be effective in reducing overspending. While it should be easy to enforce a transparent display of odds, it is not clear that gamers will use this information when making their purchase decisions. Our robustness experiment, for instance, suggests that additional information may not affect WTPs. Moreover, even when learning the full probability distribution over *many* prizes, gamers might not act on it because it is simply too much information to be considered. Instead, regulators must find ways of communicating the odds of loot boxes in an easily understandable way. Preventing gamers from being confronted with selected feedback on the reward distribution comes with the additional challenge that it is not only game developers who provide it to gamers. Even if game developers are not allowed to announce prizes others have won selectively, gamers may get similar (biased) feedback from talking to their peers or watching their videos on *Youtube* or *Twitch*.

¹⁴See, for instance, <https://usk.de/jugendschutzgesetz-aktualisiert-usk-bereitet-sich-auf-aenderungen-vor/> (accessed on September 19th, 2022).

Appendix

Appendix 4.A Additional tables

Table 4.3: Regression results — main specification — realistic beliefs

	WTP		Belief	
	No controls (1)	Controls (2)	No controls (3)	Controls (4)
Censored	5.35*** (1.43)	5.36*** (1.40)	0.435*** (0.016)	0.435*** (0.016)
Sample	2.46 (1.72)	2.33 (1.72)	0.042*** (0.013)	0.041*** (0.013)
Joint	8.34*** (1.84)	8.42*** (1.86)	0.440*** (0.019)	0.441*** (0.019)
Observations	2,872	2,872	2,872	2,872

Notes: Results from ordinary least squares (OLS) regressions on treatment dummies. The outcome variable in columns (1) and (2) is the willingness to pay and in (3) and (4) the transformed belief, i.e. the belief that the high outcome occurs conditional on the medium or high outcome being drawn. Columns (1) and (3) do not include control variables. Columns (2) and (4) control for age, gender and monthly available budget. Standard errors clustered at the subject level in parentheses. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

Table 4.4: Regression results — robustness experiment

	WTP		Belief	
	No controls (1)	Controls (2)	No controls (3)	Controls (4)
Info	6.19** (2.49)	5.82** (2.51)	0.416*** (0.060)	0.419*** (0.060)
Joint	6.78*** (2.60)	6.39** (2.68)	0.783*** (0.078)	0.786*** (0.081)
Observations	2,070	2,070	2,070	2,070

Notes: Results from ordinary least squares (OLS) regressions on treatment dummies. The outcome variable in columns (1) and (2) is the willingness to pay and in (3) and (4) the normalized belief. The independent variables are indicators that equal 1 if the participant is in the respective condition and 0 else. Columns (1) and (3) do not include control variables. Columns (2) and (4) control for age, gender and monthly available budget. Standard errors clustered at the subject level in parentheses. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

Table 4.5: Summary statistics

	Treatments			
	Control (1)	Sample (2)	Censored (3)	Joint (4)
Female	0.56	0.57	0.55	0.49
Age (years)	40.11	38.81	39.05	40.34
College	0.74	0.68	0.75	0.69
Monthly budget	696.67	734.75	642.66	686.66
Observations	151	159	157	150

Notes: Summary statistics. We include all subjects who completed the study.

Table 4.6: Loot box statistics

	Treatments			
	Control (1)	Sample (2)	Censored (3)	Joint (4)
Video Games played (hours)	1.46	1.59	1.72	1.76
Loot Box Spending (month)	25.15	13.18	7.37	15.83
Ever Overspent on Loot Boxes	0.11	0.20	0.15	0.16
Observations	105	100	116	104

Notes: Summary statistics on loot boxes. We include all subjects who state to know what loot boxes are (69% of the total sample).

Appendix 4.B Experimental setup: Robustness experiment

In this section, we present the design of *Robustness Experiment*, which is conducted in a between-subjects design. The experiment consists of three conditions: the *Control* and *Joint* conditions from the *Main Experiment*, as well as a new *Info* condition. The *Info* condition is an extension of the *Joint* condition and includes supplementary information concerning the prevalence of medium and high outcomes in a sample of 50 draws, as well as the information that the probabilities of these outcomes occurring are not equal. All other characteristics of the experiment are equivalent to those of the main Experiment, as described in Section 4.2.1.

We collected data from 617 subjects located in the UK via *Prolific* in November 2022. Screenshots of all parts of the experiment can be found in Appendix 4.C. Subjects earned a base fee of 1.50 GBP for participation. In addition, 1 out of 6 participants received a bonus payment depending on the WTP stated for one randomly selected lottery. Conditional on receiving a bonus, the average bonus paid was 5.08 GBP. The experiment took, on average, 11 minutes to complete.

Appendix 4.C Experimental details

Below we provide screenshots for all pages of the experiment. It includes the initial attention check, instructions, comprehension checks, and all survey questions. The screenshots are presented in the order in which participants progress through the experiment.

Your thoughts on daylight savings

Please write **at least 15 words** describing **your opinion** about **daylight savings in the United States**. Whether you are in favor or against daylight savings does not affect your eligibility to participate in this study. However, we ask that you write at least 15 words on your thoughts about this topic.

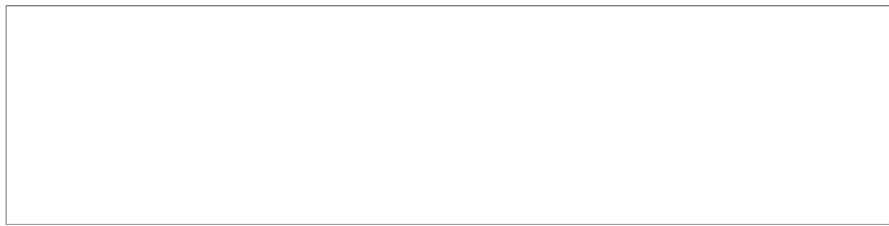
[Next](#)

Figure 4.8: Attention check at the beginning of the experiment.

Welcome

Thank you for participating in this study.

This study will take approximately 10 minutes to complete.

You will earn a fixed reward of £1.5 for completing the study. To complete the study, you will need to read all instructions carefully and answer the corresponding comprehension questions correctly.

You have a chance to win an **additional bonus** that depends on your decisions in this study.

Please enter your Prolific ID below to make sure that you will **receive your participant payment**:

[Next](#)

Figure 4.9: Welcome page.

Instructions

Please read these instructions carefully. There will be comprehension checks.

The Lotteries

In this study, we will ask you about your perception of, and preferences toward, different lotteries. Each lottery offers the chance to win different monetary prizes. The probabilities of these prizes may differ. You can find an example below:

Lottery			
Probability	40 %	30 %	30 %
Payoff	20 coins	40 coins	60 coins

Here, the lottery pays either 20 coins with 40% probability or 40 coins with 30% probability or 60 coins with 30% probability. Importantly, a lottery always only pays out one of the amounts; that is, it pays **either** 20 coins **or** 40 coins **or** 60 coins.

The Decisions and Bonus Payment

In total, you will see **5 lotteries**. For each lottery, we will ask you two questions. First, we will ask you **how often you think you would win the highest prize** of the lottery (i.e., 60 coins in the example above). After that, we will ask you for your **willingness to pay** for each lottery; that is, for the maximum amount of money you would pay to participate in this lottery.

Please indicate your **true** willingness to pay for each lottery. A payment mechanism will determine your potential bonus payment based on your stated willingness to pay. **The payment mechanism is designed in a way that gives you an incentive to indicate your true willingness to pay.** If you are interested in the details, you can find a complete description of the payment mechanism at the bottom of the page.

At the end of the experiment, the computer will randomly select one of the 5 lotteries for payment. The computer will then determine, based on your stated willingness to pay, your potential bonus payment. One out of every 6 subjects will be randomly selected to receive this additional bonus.

During the experiment, we use a fictitious currency called coins. For the bonus payment coins will be converted into British pound at a rate of 13 coins = 1 British pound.

Details on the payment mechanism

Next

Figure 4.10: Instructions on the experiment.

Details on the payment mechanism

Next

The actual price of a lottery is set randomly. For this, a price will be chosen randomly. This randomly chosen price is between 0 and the maximum amount that the respective lottery can pay out. If the price is lower than or equal to your willingness to pay, you play the lottery. Your payoff will then be determined by a random draw from that lottery. Otherwise, the randomly selected price will be your payoff.

Example 1:

The lottery pays exactly 50 coins with a probability of 40% and nothing otherwise. You state that you are willing to pay a maximum of 20 coins to play the lottery. A random number between 0 and 50 is then drawn, which determines the price of 15 coins. Since this price is lower than your willingness to pay, you play the lottery. A random draw from the lottery determines your payoff.

Example 2:

The lottery pays exactly 0 coins with a probability of 20%, exactly 60 coins with 30% probability and exactly 70 coins with 50% probability. You state that you are willing to pay a maximum of 30 coins to play the lottery. A random number between 0 and 70 is then drawn, which determines the price of 35 coins. Since this price is higher than your maximum price, you do not play the lottery. You will receive the price of 35 coins as a payoff.

Figure 4.11: Additional info box with details on the payment mechanism.

Comprehension questions

The questions below test your understanding of the instructions.

Important: If you fail to answer any one of these questions correctly, you will not be allowed to participate in the rest of the study, and you will not be able to earn a bonus payment.

1. Which payoffs are possible for the following lottery?

Lottery			
Probability	20 %	40 %	40 %
Payoff	0 coins	30 coins	80 coins

Please select one of the statements:

☐ It is possible that I get paid both 30 coins and 80 coins, i.e., I may receive a total amount of 110 coins from this lottery.

☐ I receive EITHER 30 coins OR 80 coins OR 0 coins from this lottery.

☐ I will receive at least some coins with certainty.

2. What is the probability to get a payoff of 30 coins for the following lottery?

Lottery			
Probability	50 %	40 %	10 %
Payoff	0 coins	30 coins	200 coins

Please select one of the statements:

☐ The probability to receive 30 coins is 60 %.

☐ The probability to receive 30 coins is 40 %.

☐ The probability to receive 30 coins is 0 %.

[Next](#)

Figure 4.12: Comprehension questions.

You have correctly answered the comprehension questions. We will now start with the study.

The first decision begins on the next page.

[Next](#)

Figure 4.13: Start of the experiment.

Please take a decision

Below you find information about the lottery that you can buy.

Lottery			
Probability	90 %	9 %	1 %
Payoff	0 coins	10 coins	100 coins

Imagine you would play the lottery **100 times**. How often do you think would you win **100 coins**?

0 times 100 times

Please click on the slider to indicate your belief. You can still change your decision afterwards.

How much are you willing to pay for this lottery?

Coins

Next

Details on the payment mechanism

Figure 4.14: Belief & WTP elicitation in the “Control” treatment.

Please take a decision

Below you find information about the lottery that you can buy. Moreover, you see the 5 highest values of in total 5000 random lottery draws.

Lottery			
Probability	90 %	9 %	1 %
Payoff	0 coins	10 coins	100 coins

Lottery draws	
#	Value
Draw 1	100 Coins
Draw 2	100 Coins
Draw 3	100 Coins
Draw 4	100 Coins
Draw 5	100 Coins

Imagine you would play the lottery **100 times**. How often do you think would you win **100 coins**?

0 times 100 times

Please click on the slider to indicate your belief. You can still change your decision afterwards.

How much are you willing to pay for this lottery?

Coins

Next

Details on the payment mechanism

Figure 4.15: Belief & WTP elicitation in the “Sample” treatment.

Please take a decision

Below you find information about the lottery that you can buy.

Lottery		
Probability	90 %	10 %
Payoff	0 coins	either 10 coins or 100 coins

Imagine you would play the lottery **100 times**. How often do you think would you win **100 coins**?

0 times  100 times

Please click on the slider to indicate your belief. You can still change your decision afterwards.

How much are you willing to pay for this lottery?

Coins

Next

Details on the payment mechanism

Figure 4.16: Belief & WTP elicitation in “Censored” treatment.

Please take a decision

Below you find information about the lottery that you can buy. Moreover, you see the 5 highest values of in total 5000 random lottery draws.

Lottery		
Probability	90 %	10 %
Payoff	0 coins	either 10 coins or 100 coins

Lottery draws	
#	Value
Draw 1	100 Coins
Draw 2	100 Coins
Draw 3	100 Coins
Draw 4	100 Coins
Draw 5	100 Coins

Imagine you would play the lottery **100 times**. How often do you think would you win **100 coins**?

0 times 100 times

Please click on the slider to indicate your belief. You can still change your decision afterwards.

How much are you willing to pay for this lottery?

 Coins

Next

Details on the payment mechanism

Figure 4.17: Belief & WTP elicitation in “Joint” treatment.

Please take a decision

Below you find information about the lottery that you can buy. Moreover, you see the 5 highest values of in total 400 random lottery draws.

Lottery		
Probability	90 %	10 %
Payoff	0 coins	either 10 coins or 100 coins

Lottery draws	
#	Value
Draw 1	100 Coins
Draw 2	100 Coins
Draw 3	100 Coins
Draw 4	100 Coins
Draw 5	100 Coins

Note:

- 10 and 100 coins are **not equally likely**.
- Of the last 50 random lottery draws, **4 draws were equal to 10 coins** and **0 draws were equal to 100 coins**.

You will be able to enter your belief and your willingness to pay after 15 seconds.

Figure 4.18: Additional info in “Info” treatment.

Please take a decision

Below you find information about the lottery that you can buy. Moreover, you see the 5 highest values of in total 400 random lottery draws.

Lottery		
Probability	90 %	10 %
Payoff	0 coins	either 10 coins or 100 coins

Lottery draws	
#	Value
Draw 1	100 Coins
Draw 2	100 Coins
Draw 3	100 Coins
Draw 4	100 Coins
Draw 5	100 Coins

Note:

- 10 and 100 coins are **not equally likely**.
- Of the last 50 random lottery draws, **4 draws were equal to 10 coins** and **0 draws were equal to 100 coins**.

Imagine you would play the lottery **100 times**. How often do you think would you win **100 coins**?

0 times 100 times

Please click on the slider to indicate your belief. You can still change your decision afterwards.

How much are you willing to pay for this lottery?

 Coins

Next

Details on the payment mechanism

Figure 4.19: Belief & WTP elicitation in “Info” treatment.

Please answer the following questions

Thank you for participating. Please answer the following questions. The survey will last approximately for another 10 minutes. Afterwards, you will receive your payment.

How old are you?

What is your gender?

What is your highest level of education:

What do you study? / What is your occupation?

How much money do you have available each month (after deducting fixed costs such as rent, insurance, etc., British pound)?

Next

Figure 4.20: General control questions.

Please answer the following questions

Loot boxes in video games

In the following questions we will ask you about your experiences with loot boxes in video games. Loot boxes are digital items that contain a random reward. In many games it is possible to buy such loot boxes for real world money. Examples include Packs in Fifa 22, Hextech Chests in League of Legends or Eggs in Pokémon Go. Below you also see an exemplary picture of a loot box.



How many hours do you spend playing videogames per day on average?

Have you ever heard of loot boxes in videogames before?

- ☐ Yes
☐ No

Next

Figure 4.21: Questions about loot box experience.

Please answer the following questions

How much did you spend on loot boxes during the last year on average? (in British pound):

Have you ever spend more than you planned to on loot boxes?

- ☐ Yes
- ☐ No

Next

Figure 4.22: Questions for loot box users.

Please answer the following questions

Using the scale provided, please indicate how much each of the following statements reflects how you typically are.

	Not at all				Very much
I am good at resisting temptation.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
I have a hard time breaking bad habits.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
I am lazy.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
I say inappropriate things.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
I do certain things that are bad for me, if they are fun.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
I refuse things that are bad for me.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
I wish I had more self-discipline.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5

Figure 4.23: Self-control survey questions (Part 1).

People would say that I have iron self-discipline.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
Pleasure and fun sometimes keep me from getting work done.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
I have trouble concentrating.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
I am able to work effectively toward long-term goals.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
Sometimes I can't stop myself from doing something, even if I know it is wrong.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
I often act without thinking through all the alternatives.	☐ 1	☐ 2	☐ 3	☐ 4	☐ 5

Next

Figure 4.24: Self-control survey questions (Part 2).

Please answer the following questions

In this part, we are interested in your gambling behavior.

Thinking about the last 12 months ...

Have you bet more than you could really afford to lose?	☐ Never	☐ Sometimes	☐ Most of the time	☐ Always
Have you needed to gamble with larger amounts of money to get the same feeling of excitement?	☐ Never	☐ Sometimes	☐ Most of the time	☐ Always
Have you gone back on another day to try to win back the money you lost?	☐ Never	☐ Sometimes	☐ Most of the time	☐ Always
Have you borrowed money or sold anything to gamble?	☐ Never	☐ Sometimes	☐ Most of the time	☐ Always
Have you felt that you might have a problem with gambling?	☐ Never	☐ Sometimes	☐ Most of the time	☐ Always

Figure 4.25: Gambling survey questions (Part 1).

Have people criticised your betting or told you that you had a gambling problem, whether or not you thought it was true?	☐ Never	☐ Sometimes	☐ Most of the time	☐ Always
Have you felt guilty about the way you gamble or what happens when you gamble?	☐ Never	☐ Sometimes	☐ Most of the time	☐ Always
Has gambling caused you any health problems, including stress or anxiety?	☐ Never	☐ Sometimes	☐ Most of the time	☐ Always
Has your gambling caused any financial problems for you or your household?	☐ Never	☐ Sometimes	☐ Most of the time	☐ Always

[Next](#)

Figure 4.26: Gambling survey questions (Part 2).

You have finished all tasks

You have completed this study in its entirety. **You will receive the fixed reward of £1.5 credited to your Prolific account.**

You are eligible for the bonus payment. After the end of the entire study, one out of every 6 subjects will be randomly selected to receive an additional bonus. If you are selected for payment, your additional bonus will be determined as follows:

The computer will randomly select one lottery for payout, with equal probability, and determine a price for the lottery.

The actual price of a lottery is set randomly. For this, a price will be chosen randomly. This randomly chosen price is between 0 and the maximum amount that the respective lottery can pay out. If the price is lower than or equal to your willingness to pay, you play the lottery. Your payoff will then be determined by a random draw from that lottery. Otherwise, the randomly selected price will be your payoff.

If you won the bonus, we will inform you in the next days using your Prolific-ID.

Next

Figure 4.27: Last page in the experiment.

Declaration of Contribution

Hereby I, Tobias Felix Werner, declare that the chapter “What drives demand for loot boxes?” is coauthored by Simon Cordes and Markus Dertwinkel-Kalt. All authors contributed equally to the chapter.

Signature of coauthor (Simon Cordes): S. Cordes
Simon Cordes

Signature of coauthor (Markus Dertwinkel-Kalt): M. Dertwinkel-Kalt
Prof. Dr. Markus Dertwinkel-Kalt

Bibliography

- Abada, Ibrahim and Xavier Lambin**, “Artificial Intelligence: Can Seemingly Collusive Outcomes Be Avoided?,” *Management Science* (*forthcoming*), 2022.
- Abreu, Dilip**, “On the Theory of Infinitely Repeated Games with Discounting,” *Econometrica*, 1988, *56* (2), 383–396.
- Agrawal, Ajay, Joshua Gans, and Avi Goldfarb**, *The Economics of Artificial Intelligence: An Agenda*, University of Chicago Press, 2019.
- Álvarez-Benjumea, Amalia**, “Exposition to xenophobic content and support for right-wing populism: The asymmetric role of gender,” *Social Science Research*, 2020, *92*, 102480.
- **and Fabian Winter**, “Normative Change and Culture of Hate: An Experiment in Online Environments,” *European Sociological Review*, 03 2018, *34* (3), 223–237.
- **and —**, “The breakdown of antiracist norms: A natural experiment on hate speech after terrorist attacks,” *Proceedings of the National Academy of Sciences*, 2020, *117* (37), 22800–22804.
- Arulkumaran, Kai, Marc Peter Deisenroth, Miles Brundage, and Anil Anthony Bharath**, “Deep reinforcement learning: A brief survey,” *IEEE Signal Processing Magazine*, 2017, *34* (6), 26–38.
- Asker, John, Chaim Fershtman, and Ariel Pakes**, “The Impact of AI Design on Pricing,” *Nber Working Paper 28535*, 2022.
- Assad, Stephanie, Emilio Calvano, Giacomo Calzolari, Robert Clark, Vincenzo Denicolò, Daniel Ershov, Justin Johnson, Sergio Pastorello, Andrew Rhodes, Lei Xu, and Matthijs Wildenbeest**, “Autonomous algorithmic collusion:

- Economic research and policy implications,” *Oxford Review of Economic Policy*, 2021, 37 (3), 459–478.
- , **Robert Clark, Daniel Ershov, and Lei Xu**, “Algorithmic Pricing and Competition : Empirical Evidence from the German Retail Gasoline Market,” *CESifo Working Paper 8521*, 2022.
- Babcock, Linda, Maria P Recalde, Lise Vesterlund, and Laurie Weingart**, “Gender Differences in Accepting and Receiving Requests for Tasks with Low Promotability,” *American Economic Review*, 2017, 107 (3), 714–47.
- Barredo, Alejandro Arrieta, Natalia Díaz-Rodríguez, Javier Del Ser, Adrien Bennetot, Siham Tabik, Alberto Barbado, Salvador Garcia, Sergio Gil-Lopez, Daniel Molina, Richard Benjamins, Raja Chatila, and Francisco Herrera**, “Explainable Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI,” *Information Fusion*, 2020, 58, 82–115.
- Barron, Greg and Eldad Yechiam**, “Private e-mail requests and the diffusion of responsibility,” *Computers in Human Behavior*, 2002, 18 (5), 507–520.
- Barron, Kai, Steffen Huck, and Philippe Jehiel**, “Everyday econometricians: Selection neglect and overoptimism when learning from others,” Working Paper 2019.
- Becker, Gordon M., Morris H. DeGroot, and Jacob Marschak**, “Measuring utility by a single-response sequential method,” *Behavioral science*, 1964, 9 (3), 226–232.
- Bénabou, Roland and Jean Tirole**, “Incentives and Prosocial Behavior,” *American Economic Review*, 2006, 96 (5), 1652–1678.
- Benndorf, Volker, Dorothea Kübler, and Hans-Theo Normann**, “Privacy concerns, voluntary disclosure of information, and unraveling: An experiment,” *European Economic Review*, 2015, 75, 43–59.
- Bick, Alexander, Adam Blandin, and Karel Mertens**, “Work from home before and after the Covid-19 outbreak,” Working Paper 2022.

- Bó, Pedro Dal**, “Cooperation under the Shadow of the Future : Experimental Evidence from Infinitely Repeated Games,” *American Economic Review*, 2005, *95* (5), 1591–1604.
- **and Guillaume R. Fréchette**, “The Evolution of Cooperation in Infinitely Repeated Games: Experimental Evidence,” *American Economic Review*, 2011, *101* (1), 411–429.
- **and —**, “On the determinants of cooperation in infinitely repeated games: A survey,” *Journal of Economic Literature*, 2018, *56* (1), 60–114.
- **and —**, “Strategy choice in the infinitely repeated Prisoner’s Dilemma,” *American Economic Review*, 2019, *109* (11), 3929–52.
- Bolton, Patrick, Xavier Freixas, and Joel Shapiro**, “Conflicts of interest, information provision, and competition in the financial services industry,” *Journal of Financial Economics*, 2007, *85* (2), 297–330.
- Boosey, Luke, Philip Brookins, and Dmitry Ryvkin**, “Contests with group size uncertainty: Experimental evidence,” *Games and Economic Behavior*, 2017, *105*, 212–229.
- Borchert, Kathrin, Matthias Hirth, Michael Kummer, Ulrich Laitenberger, Olga Slivko, and Steffen Viete**, “Unemployment and Online Labor - Evidence from Microtasking,” *MIS Quarterly (forthcoming)*, 2022.
- Borenstein, Severin and Andrea Shepard**, “Dynamic Pricing in Retail Gasoline Markets,” *The RAND Journal of Economics*, 1996, *27* (3), 429–451.
- Brown, Alexander L., Colin F. Camerer, and Dan Lovallo**, “To Review or Not to Review? Limited Strategic Thinking at the Movie Box Office,” *American Economic Journal: Microeconomics*, 2012, *4* (2), 1–26.
- Brown, Zach and Alexander MacKay**, “Competition in Pricing Algorithms,” *American Economic Journal: Microeconomics (forthcoming)*, 2022.
- Bruttel, Lisa**, “The Effects of Recommended Retail Prices on Consumer and Retailer Behaviour,” *Economica*, 2018, *85* (339), 649–668.
- Buehler, Stefan and Dennis L. Gärtner**, “Making sense of nonbinding retail-price recommendations,” *American Economic Review*, 2013, *103* (1), 335–59.

- Bundeskartellamt and Autorité de la concurrence**, “Algorithms and Competition,” Available at: https://www.bundeskartellamt.de/SharedDocs/Publikation/EN/Berichte/Algorithms_and_Competition_Working-Paper.pdf 2019.
- Byrne, David P. and Nicolas De Roos**, “Learning to coordinate: A study in retail gasoline,” *American Economic Review*, 2019, *109* (2), 591–619.
- Cadsby, C. Bram, Elizabeth Maynes, and Viswanath Umashanker Trivedi**, “Tax compliance and obedience to authority at home and in the lab: A new experimental approach,” *Experimental Economics*, 2006, *9* (4), 343–359.
- Calvano, Emilio, Giacomo Calzolari, Vincenzo Denicolò, and Sergio Pastorello**, “Algorithmic Pricing What Implications for Competition Policy?,” *Review of Industrial Organization*, 2019, *55* (1), 155–171.
- , —, —, and —, “Artificial intelligence, algorithmic pricing, and collusion,” *American Economic Review*, 2020, *110* (10), 3267–97.
- , —, —, and —, “Algorithmic collusion with imperfect monitoring,” *International Journal of Industrial Organization*, 2021, *79*, 102712.
- , —, —, **Joseph E. Harrington**, and **Sergio Pastorello**, “Protecting consumers from collusive prices due to AI,” *Nature*, 2020, *370* (6520), 1040–1042.
- Campos-Mercade, Pol**, “The volunteer’s dilemma explains the bystander effect,” *Journal of Economic Behavior & Organization*, 2021, *186*, 646–661.
- Chandrasekher, Madhav, Mira Frick, Ryota Iijima, and Yves Le Yaouanq**, “Dual-Self Representations of Ambiguity Preferences,” *Econometrica*, 2022, *90* (3), 1029–1061.
- Chen, Daniel L., Martin Schonger, and Chris Wickens**, “oTree—An open-source platform for laboratory, online, and field experiments,” *Journal of Behavioral and Experimental Finance*, 2016, *9*, 88–97.
- Chen, Le, Alan Mislove, and Christo Wilson**, “An empirical analysis of algorithmic pricing on amazon marketplace,” in “Proceedings of the 25th International Conference

on World Wide Web” International World Wide Web Conferences Steering Committee 2016, pp. 1339–1349.

Chen, Ningyuan, Adam N. Elmachoub, Michael L. Hamilton, and Xiao Lei, “Loot Box Pricing and Design,” *Management Science*, 2021, 67 (8), 4809–4825.

Competition & Markets Authority, “Algorithms: How they can reduce competition and harm consumers,” Available at: <https://www.gov.uk/government/publications/algorithms-how-they-can-reduce-competition-and-harm-consumers> 2021.

Cornelissen, Thomas, Christian Dustmann, and Uta Schönberg, “Peer Effects in the Workplace,” *American Economic Review*, 2017, 107 (2), 425–56.

Crandall, Jacob W., Mayada Oudah, Tennom, Fatimah Ishowo-Oloko, Sherief Abdallah, Jean François Bonnefon, Manuel Cebrian, Azim Shariff, Michael A. Goodrich, and Iyad Rahwan, “Cooperating with machines,” *Nature Communications*, 2018, 9 (1), 1–12.

Croson, Rachel and Melanie Marks, “The effect of recommended contributions in the voluntary provision of public goods,” *Economic Inquiry*, 2001, 39 (2), 238–249.

Danz, David, Lise Vesterlund, and Alistair J. Wilson, “Belief Elicitation and Behavioral Incentive Compatibility,” *American Economic Review*, 2022, 112 (9), 2851–2883.

—, **Neeraja Gupta, Marissa Lepper, Lise Vesterlund, and K Pun Winichakul**, “Going virtual : A step-by-step guide to taking the in-person experimental lab online,” Working Paper 2021.

Darley, John M. and Bibb Latané, “Bystander intervention in emergencies: Diffusion of responsibility.,” *Journal of Personality and Social Psychology*, 1968, 8 (4, Pt.1), 377–383.

Davies, Stephen, Matthew Olczak, and Heather Coles, “Tacit collusion, firm asymmetries and numbers: Evidence from EC merger cases,” *International Journal of Industrial Organization*, 2011, 29 (2), 221–231.

- Deversi, Marvin, Alessandro Ispano, and Peter Schwardmann**, “Spin doctors: An experiment on vague disclosure,” *European Economic Review*, 2021, *139*, 1038–1072.
- Diekmann, Andreas**, “Volunteer’s Dilemma,” *Journal of Conflict Resolution*, 1985, *29* (4), 605–610.
- , “Volunteer’s Dilemma. A Social Trap without a Dominant Strategy and some Empirical Results,” in “Paradoxical Effects of Social Behavior,” Physica-Verlag HD, 1986, pp. 187–197.
- , “Cooperation in an asymmetric Volunteer’s dilemma game theory and experimental evidence,” *International Journal of Game Theory*, 1993, *22* (1), 75–85.
- Difallah, Djellel Eddine, Michele Catasta, Gianluca Demartini, Panagiotis G Ipeirotis, and Philippe Cudré-Mauroux**, “The Dynamics of Micro-Task Crowdsourcing: The Case of Amazon MTurk,” in “Proceedings of the 24th International Conference on World Wide Web” International World Wide Web Conferences Steering Committee 2015, pp. 238–247.
- Drummond, Aaron and James D. Sauer**, “Video game loot boxes are psychologically akin to gambling,” *Nature human behaviour*, 2018, *2* (8), 530–532.
- , —, **Lauren C. Hall, David Zendle, and Malcolm R. Loudon**, “Why loot boxes could be regulated as gambling,” *Nature human behaviour*, 2020, *4* (10), 986–988.
- Duffy, John and Nick Feltovich**, “Correlated equilibria, good and bad: an experimental study,” *International Economic Review*, 2010, *51* (3), 701–721.
- Engel, Christoph**, “How much collusion? A meta-analysis of oligopoly experiments,” *Journal of Competition Law & Economics*, 2007, *3* (4), 491–549.
- , “Tacit Collusion: The Neglected Experimental Evidence,” *Journal of Empirical Legal Studies*, 2015, *12* (3), 537–577.
- Enke, Benjamin**, “What You See Is All There Is,” *The Quarterly Journal of Economics*, 2020, *135* (3), 1363–1398.
- **and Thomas Graeber**, “Cognitive Uncertainty,” *Nber Working Paper 26518*, 2019.

- **and** — , “Cognitive Uncertainty in Intertemporal Choice,” *Nber Working Paper 29577*, 2021.
- Esponda, Ignacio and Emanuel Vespa**, “Endogenous sample selection: A laboratory study,” *Quantitative Economics*, 2018, 9 (1), 183–216.
- European Commission**, “Final report on the E-commerce Sector Inquiry,” Available at: https://ec.europa.eu/competition/antitrust/sector_inquiry_final_report_en.pdf 2017.
- Ezrachi, Ariel and Maurice E. Stucke**, “Virtual Competition,” *Journal of European Competition Law & Practice*, 2016, 7 (9), 585–586.
- **and** — , “Artificial intelligence & collusion: When computers inhibit competition,” *University of Illinois Law Review*, 2017, 2017 (5), 1775–1810.
- Faber, Riemer P. and Maarten C.W. Janssen**, “On the effects of suggested prices in gasoline markets,” *The Scandinavian Journal of Economics*, 2019, 121 (2), 676–705.
- Falk, Armin, Anke Becker, Thomas Dohmen, David Huffman, and Uwe Sunde**, “The Preference Survey Module: A Validated Instrument for Measuring Risk, Time, and Social Preferences,” *Management Science (forthcoming)*, 2022.
- Farrell, Diana and Fiona Greig**, “The Online Platform Economy: Has Growth Peaked?,” Working Paper 2017.
- Federal Trade Commission**, “FTC Video Game Loot Box Workshop,” Available at: https://www.ftc.gov/system/files/documents/reports/staff-perspective-paper-loot-box-workshop/loot_box_workshop_staff_perspective.pdf 2020.
- Fischhoff, Baruch and Wändi Bruine De Bruin**, “Fifty-Fifty=50%?,” *Journal of Behavioral Decision Making*, 1999, 12 (2), 149–163.
- Fonseca, Miguel A. and Hans-Theo Normann**, “Explicit vs. tacit collusion—The impact of communication in oligopoly experiments,” *European Economic Review*, 2012, 56 (8), 1759–1772.
- Foros, Øystein and Frode Steen**, “Vertical control and price cycles in gasoline retailing,” *The Scandinavian Journal of Economics*, 2013, 115 (3), 640–661.

- Franzen, Axel**, “Group Size and One-Shot Collective Action,” *Rationality and Society*, 1995, 7 (2), 183–200.
- Friedman, James W.**, “A non-cooperative equilibrium for supergames,” *Review of Economic Studies*, 1971, 38 (1), 1–12.
- Fudenberg, Drew, David G. Rand, and Anna Dreber**, “Slow to anger and fast to forgive: Cooperation in an uncertain world,” *American Economic Review*, 2012, 102 (2), 720–749.
- Gale, Douglas and Hamid Sabourian**, “Complexity and competition,” *Econometrica*, 2005, 73 (3), 739–769.
- Goeree, Jacob K., Charles A. Holt, and Angela M. Smith**, “An experimental examination of the volunteer’s dilemma,” *Games and Economic Behavior*, 2017, 102, 303–315.
- Greiner, Ben**, “Subject pool recruitment procedures: organizing experiments with ORSEE,” *Journal of the Economic Science Association*, 2015, 1 (1), 114–125.
- Hans Böckler Stiftung**, “Studien zu Home Office und Mobiler Arbeit,” Available at: <https://www.boeckler.de/de/auf-einen-blick-17945-Auf-einen-Blick-Studien-zu-Homeoffice-und-mobiler-Arbeit-28040.html> 2021.
- Hansen, Karsten T., Kanishka Misra, and Mallesh M. Pai**, “Frontiers: Algorithmic collusion: Supra-competitive prices via independent algorithms,” *Marketing Science*, 2021, 40 (1), 1–12.
- Harrington, Joseph E.**, “Developing competition law for collusion by autonomous artificial agents,” *Journal of Competition Law and Economics*, 2018, 14 (3), 331–363.
- , “The Effect of Outsourcing Pricing Algorithms on Market Competition,” *Management Science*, 2022, 68 (9), 6889–6906.
- , **Roberto Hernan Gonzalez, and Praveen Kujal**, “The relative efficacy of price announcements and express communication for collusion: Experimental findings,” *Journal of Economic Behavior & Organization*, 2016, 128, 251–264.

- Harstad, Ronald, Stephen Martin, and Hans-Theo Normann**, *Experimental tests of consciously parallel behavior in oligopoly* Applied Industrial Organization, Cambridge University Press, 1998.
- Healy, Andrew J and Jennifer G Pate**, “Cost asymmetry and incomplete information in a volunteer’s dilemma experiment,” *Social Choice and Welfare*, 2018, 51 (3), 465–491.
- Hettich, Matthias**, “Algorithmic Collusion: Insights from Deep Learning,” Working Paper 2021.
- Hill, Dan**, “How much is your spare room worth?,” *IEEE Spectrum*, 2015, 52 (9), 32–58.
- Hillenbrand, Adrian and Fabian Winter**, “Volunteering under population uncertainty,” *Games and Economic Behavior*, 2018, 109, 65–81.
- Hoda, Rashina, Norsaremah Salleh, and John Grundy**, “The Rise and Evolution of Agile Software Development,” *IEEE Software*, 2018, 35 (5), 58–63.
- Holt, Charles A and Douglas Davis**, “The effects of non-binding price announcements on posted-offer markets,” *Economics letters*, 1990, 34 (4), 307–310.
- Horstmann, Niklas, Jan Krämer, and Daniel Schnurr**, “Number Effects and Tacit Collusion in Experimental Oligopolies,” *Journal of Industrial Economics*, 2018, 66 (3), 650–700.
- Hossain, Tanjim and Ryo Okui**, “The Binarized Scoring Rule,” *Review of Economic Studies*, 2013, 80 (3), 984–1001.
- Huck, Steffen, Hans-Theo Normann, and Jörg Oechssler**, “Two are few and four are many: number effects in experimental oligopolies,” *Journal of Economic Behavior & Organization*, 2004, 53 (4), 435–446.
- Hunold, Matthias, Ulrich Laitenberger, and Guillaume Thébaudin**, “Bye-box: An Analysis of Non-Promotion on the Amazon Marketplace,” Working Paper 2022.
- Imhof, Lorens A., Drew Fudenberg, and Martin A. Nowak**, “Tit-for-tat or Win-stay, Lose-shift?,” *Journal of Theoretical Biology*, 2007, 247 (3), 574–580.

- Jain, Ayush, Akash Das Sarma, Aditya Parameswaran, and Jennifer Widom**, “Understanding Workers, Developing Effective Tasks, and Enhancing Marketplace Dynamics: A Study of a Large Crowdsourcing Marketplace,” *Proceedings of the VLDB Endowment*, 2017, 10 (7), 829–840.
- Jeschonneck, Malte**, “Collusion among Autonomous Pricing Algorithms Utilizing Function Approximation Methods,” Working Paper 2021.
- Jin, Ginger Zhe, Michael Luca, and Daniel Martin**, “Is No News (Perceived As) Bad News? An Experimental Investigation of Information Disclosure,” *American Economic Journal: Microeconomics*, 2021, 13 (2), 141–173.
- , —, and —, “Complex Disclosure,” *Management Science*, 2022, 68 (5), 3236–3261.
- Johnson, Justin, Andrew Rhodes, and Matthijs Wildenbeest**, “Platform Design When Sellers Use Pricing Algorithms,” Working Paper 2022.
- Johnson, Justin Pappas**, “Open Source Software: Private Provision of a Public Good,” *Journal of Economics & Management Strategy*, 2004, 11 (4), 637–662.
- Jones, Matthew T.**, “Strategic complexity and cooperation: An experimental study,” *Journal of Economic Behavior and Organization*, 2014, 106, 352–366.
- Kenwood, Carolyn A.**, “A Business Case Study of Open Source Software,” Available at: <https://apps.dtic.mil/sti/pdfs/ADA459563.pdf>, 2001.
- Kim, Jae Wook and J. Keith Murnighan**, “The effects of connectedness and self interest in the organizational volunteer dilemma,” *International Journal of Conflict Management*, 1997, 8 (1), 32–51.
- Klein, Timo**, “Autonomous algorithmic collusion: Q-learning under sequential pricing,” *The RAND Journal of Economics*, 2021, 52 (3), 538–558.
- Köbis, Nils, Jean François Bonnefon, and Iyad Rahwan**, “Bad machines corrupt good morals,” *Nature Human Behaviour*, 2021, 5 (6), 679–685.
- Koehler, Jonathan J. and Molly Mercer**, “Selection Neglect in Mutual Fund Advertisements,” *Management Science*, 2009, 55 (7), 1107–1121.

- Kopányi-Peuker, Anita**, “Yes, I’ll Do it: A Large-Scale Experiment on the Volunteer’s Dilemma,” *Journal of Behavioral and Experimental Economics*, 2019, 80, 211 – 218.
- Kühn, Kai-Uwe and Steve Tadelis**, “Algorithmic Collusion,” *Prepared for CRESSE 2017*, 2017.
- Latané, Bibb and Steve Nida**, “Ten Years of Research on Group Size and Helping,” *Psychological Bulletin*, 1981, 89 (2), 308–324.
- Lee, Seungjoon, Dave Levin, Vijay Gopalakrishnan, and Bobby Bhattacharjee**, “Backbone construction in selfish wireless networks,” *AMC SIGMETRICS Performance Evaluation Review*, 2007, 35 (1), 121–132.
- Lefez, Willy**, “Price recommendations and the value of data: A mechanism design approach,” Working Paper 2021.
- Leisten, Matthew**, “Algorithmic Competition , with Humans,” Working Paper 2022.
- Lerer, Adam and Alexander Peysakhovich**, “Maintaining cooperation in complex social dilemmas using deep reinforcement learning,” Working Paper 2017.
- Li, Jiawei, Stephen Leider, Damian R. Beil, and Izak Duenyas**, “Running On-line Experiments Using Web-Conferencing Software,” *Journal of the Economic Science Association*, 2021, 7 (2), 167–183.
- López-Pérez, Raúl, Ágnes Pintér, and Rocío Sánchez-Mangas**, “Some conditions (not) affecting selection neglect: Evidence from the lab,” *Journal of Economic Behavior & Organization*, 2022, 195, 140–157.
- March, Christoph**, “Strategic interactions between humans and artificial intelligence: Lessons from experiments with computer players,” *Journal of Economic Psychology*, 2021, p. 102426.
- McConnell, Steve**, “Open-Source Methodology: Ready for Prime Time?,” *IEEE Software*, 1999, 16 (4), 6–8.
- Mehra, Salil K.**, “Antitrust and the robo-seller: Competition in the time of algorithms,” *Minnesota Law Review*, 2016, 100 (4), 1323–1375.

- Miklós-Thal, Jeanine and Catherine Tucker**, “Collusion by algorithm: Does better demand prediction facilitate coordination between sellers?,” *Management Science*, 2019, *65* (4), 1552–1561.
- Miller, Nathan H. and Matthew C. Weinberg**, “Understanding the Price Effects of the MillerCoors Joint Venture,” *Econometrica*, 2017, *85* (6), 1763–1791.
- Mnih, Volodymyr, Koray Kavukcuoglu, David Silver, Andrei A Rusu, Joel Veness, Marc G Bellemare, Alex Graves, Martin Riedmiller, Andreas K Fidje-land, Georg Ostrovski, Stig Petersen, Charles Beattie, Amir Sadik, Ioannis Antonoglou, Helen King, Dhharshan Kumaran, Daan Wierstra, Shane Legg, and Demis Hassabis**, “Human-level control through deep reinforcement learning,” *Nature*, 2015, *518* (7540), 529–533.
- Murnighan, J. Keith, Jae Wook Kim, and A. Richard Metzger**, “The Volunteer Dilemma,” *Administrative Science Quarterly*, 1993, *38* (4), 515–538.
- Musolff, Leon**, “Algorithmic Pricing Facilitates Tacit Collusion: Evidence from E-Commerce,” in “Proceedings of the 23rd ACM Conference on Economics and Computation” 2022, pp. 32–33.
- Normann, Hans-Theo and Martin Sternberg**, “Human-algorithm interaction: Algorithmic pricing in hybrid laboratory markets,” *European Economic Review*, 2023, *152*, 104347.
- Nowak, Martin and Karl Sigmund**, “A strategy of win-stay, lose-shift that outperforms tit-for-tat in the Prisoner’s Dilemma game,” *Nature*, 1993, *364* (6432), 56–58.
- O’Connor, Jason and Nathan E. Wilson**, “Reduced demand uncertainty and the sustainability of collusion: How AI could affect competition,” *Information Economics and Policy*, 2021, *54*, 100882.
- Pavlov, Vladimir and Ron Berman**, “Price Manipulation in Peer-to-Peer Markets and the Sharing Economy,” Working Paper 2019.
- Posch, Martin**, “Win-stay, lose-shift strategies for repeated games - Memory length, aspiration levels and noise,” *Journal of Theoretical Biology*, 1999, *198* (2), 183–195.

- Przepiorka, Wojtek and Joël Berger**, “The sanctioning dilemma: A quasi-experiment on social norm enforcement in the train,” *European Sociological Review*, 2016, *32* (3), 439–451.
- Romero, Julian and Yaroslav Rosokha**, “Constructing strategies in the indefinitely repeated prisoner’s dilemma game,” *European Economic Review*, 2018, *104*, 185–219.
- Roth, Alvin E. and J. Keith Murnighan**, “Equilibrium behavior and repeated play of the prisoner’s dilemma,” *Journal of Mathematical Psychology*, 1978, *17* (2), 189–198.
- Schlütter, Frank**, “Managing Seller Conduct in Online Marketplaces and Platform Most-Favored Nation Clauses,” Working Paper 2022.
- Schotter, Andrew**, “Decision making with naive advice,” *American Economic Review*, 2003, *93* (2), 196–201.
- **and Barry Sopher**, “Social learning and coordination conventions in intergenerational games: An experimental study,” *Journal of Political Economy*, 2003, *111* (3), 498–529.
- Schwalbe, Ulrich**, “Algorithms, Machine Learning, and Collusion,” *Journal of Competition Law & Economics*, 2018, *14* (4), 568–607.
- Silver, David, Aja Huang, Chris J. Maddison, Arthur Guez, Laurent Sifre, George Van Den Driessche, Julian Schrittwieser, Ioannis Antonoglou, Veda Panneershelvam, Marc Lanctot, Sander Dieleman, Dominik Grewe, John Nham, Nal Kalchbrenner, Ilya Sutskever, Timothy Lillicrap, Madeleine Leach, and Koray Kavukcuoglu**, “Mastering the game of Go with deep neural networks and tree search,” *Nature*, 2016, *529* (7585), 484–489.
- Silverman, Dan, Joel Slemrod, and Neslihan Uler**, “Distinguishing the role of authority “in” and authority “to”,” *Journal of Public Economics*, 2014, *113*, 32–42.
- Sonnemann, Ulrich, Colin F. Camerer, Craig R. Fox, and Thomas Langer**, “How psychological framing affects economic market prices in the lab and field,” *Proceedings of the National academy of Sciences*, 2013, *110* (29), 11779–11784.

- Sonntag, Axel and Daniel John Zizzo**, “Institutional authority and collusion,” *Southern Economic Journal*, 2015, 82 (1), 13–37.
- Sutton, Richard S. and Andrew G. Barto**, *Reinforcement learning: An introduction*, 2 ed., MIT Press, 2018.
- Teh, Tat-How**, “Platform Governance,” *American Economic Journal: Microeconomics*, 2022, 14 (3), 213–54.
- The Norwegian Consumer Counsel**, “Insert Coin: How the Gaming Industry exploits consumers using loot boxes,” Available at: <https://fil.forbrukerradet.no/wp-content/uploads/2022/05/2022-05-31-insert-coin-publish.pdf> 2022.
- Waltman, Ludo and Uzay Kaymak**, “Q-learning agents in a Cournot oligopoly model,” *Journal of Economic Dynamics and Control*, 2008, 32 (10), 3275–3293.
- Watkins, Christopher John Cornish Hellaby**, “Learning from Delayed Rewards.” PhD dissertation, King’s College, Cambridge 1989.
- **and Peter Dyan**, “Q-Learning,” *Machine Learning*, 1992, 8, 279–292.
- Weesie, Jeroen**, “Incomplete information and Timing in the Volunteer’s Dilemma: A comparison of Four Models,” *Journal of Conflict Resolution*, 1994, 38 (3), 557–585.
- Werner, Tobias**, “Algorithmic and Human Collusion,” Working Paper 2021.
- Wieting, Marcel and Geza Sapi**, “Algorithms in the marketplace: An empirical analysis of automated pricing in e-commerce,” Working Paper 2021.
- Wright, Julian**, “Punishment strategies in repeated games: Evidence from experimental markets,” *Games and Economic Behavior*, 2013, 82, 91–102.
- Xiao, Leon Y.**, “Breaking Ban: Belgium’s ineffective gambling law regulation of video game loot boxes,” Working Paper 2022.
- Zendle, David and Paul Cairns**, “Video game loot boxes are linked to problem gambling: Results of a large-scale survey,” *PLOS ONE*, 11 2018, 13 (11), 1–12.

Zhang, Xiaoquan Michael and Feng Zhu, “Group size and Incentives to Contribute: A Natural Experiment at Chinese Wikipedia,” *American Economic Review*, 2011, *101* (4), 1601–15.

Zhao, Shuchen, Kristian López Vargas, Daniel Friedman, and Marco Antonio Gutierrez Chavez, “UCSC LEEPS Lab Protocol for Online Economics Experiments,” Working Paper 2020.