



universität
wien

MASTERARBEIT / MASTER'S THESIS

Titel der Masterarbeit / Title of the Master's Thesis

Experimental Evidence: Loss of Control & Immigration
Attitudes

verfasst von / submitted by

Julian Polzin, B.Sc.

angestrebter akademischer Grad / in partial fulfilment of the requirements for the degree of
Master of Science (MSc)

Wien, 2022 / Vienna 2022

Studienkennzahl lt. Studienblatt /
degree programme code as it appears on
the student record sheet:

UA 066 913

Studienrichtung lt. Studienblatt /
degree programme as it appears on
the student record sheet:

Masterstudium Volkswirtschaftslehre

Betreut von / Supervisor:

Univ.-Prof. Dr. Jean-Robert Tyran

Acknowledgements

First of all, I want to deeply thank my supervisor, Univ.-Prof. Dr. Jean-Robert Tyran, who gave me the possibility to design and organize an explorative experiment and enabled its funding. I am extremely grateful for the valuable feedback and support that he provided despite his tight schedule as Vice-Rector for Research and International Affairs at University of Vienna. I sincerely appreciate the experience of conducting an experimental study from scratch which he made possible.

Moreover, I want to express special thanks to Dipl.-Math. oec. Dr. Christian Koch who supported me constantly throughout the whole process from designing the experiment to the finalization of this thesis. The brainstorming about design aspects of the experiment and his useful advice helped me a lot and opened up new perspectives for me. I also truly enjoyed the classes in Behavioral & Experimental Economics taught by Assistant Professor Koch at University of Vienna during my Master's program. Finally, I thank my family, friends and colleagues who were by my side throughout my studies.

Universität Wien, Oktober 2022

Abstract

English

This master thesis introduces an experiment which explores the impact of perceived control loss on immigration attitudes. In a sequence of choices that are framed as immigration decisions, subjects face interference from the computer which can be resisted by performing a costly protest. The results show that subjects are willing to pay a control premium to gain back control over their decisions which is in line with observations from the literature. Moreover, immigration decisions seem to be influenced by a compound of social and risk preferences as well as spillover effects from the feeling of control loss. Hence, the novel experimental design indicates that control preferences play a role in the forming of immigration attitudes.

German

Diese Master-These führt ein Experiment ein, das den Einfluss eines empfundenen Kontrollverlustes auf die Einstellung zu Immigration untersucht. In einer Sequenz von Entscheidungen, die einen Immigrationskontext suggerieren, sind die Subjekte der Einmischung eines Computers ausgesetzt. Es besteht die Möglichkeit Kosten in einen Protest zu investieren, um sich gegen diese Störung zu wehren. Die Resultate zeigen, dass die Subjekte bereit sind einen Kontrollaufschlag zu bezahlen, um die Kontrolle über ihre Entscheidungen wiederzuerlangen, was im Einklang mit der einschlägigen Literatur steht. Darüber hinaus scheinen Immigrationsentscheidungen von einer Verbindung aus Risiko- und sozialen Präferenzen wie auch Übertragungseffekten des Gefühls von Kontrollverlust abzuhängen. Daher deutet das neuartige experimentelle Design darauf hin, dass Kontrollpräferenzen eine Rolle in der Bildung von Einstellungen zu Immigration spielen.

Contents

1	Introduction	1
2	Literature Review	4
3	Experimental Design	7
3.1	Experimental Game	7
3.2	Procedures	13
4	Hypotheses	16
4.1	Other-regarding behavior	16
4.2	Control Preferences	18
4.3	Immigration choices	20
5	Results	23
5.1	Subject pool characteristics	23
5.2	Other-regarding behavior	23
5.3	Control Preferences	27
5.4	Immigration choices	38
6	Discussion	44
7	Conclusion	51
	References	53
	Appendix	56
A1	Introductions	56

List of Figures

3.1	Image Identifying Task	8
3.2	Example Choice Situation	9
3.3	Alert Message for an Overruled Decision	10
3.4	Summary of Outcomes for Part 2	15
5.1	Control Premiums and Protest Amounts	28
5.2	Individual Control Premiums	32
5.3	Individual Protest Amounts	35
5.4	Acceptance Rates depending on Sequence Position	40

List of Tables

3.1	Sequence of Choice Situations	11
5.1	Certainty Equivalents	24
5.2	Acceptance Rates	26
5.3	Overruled Scenarios with respective P and CP	28
5.4	Control Premiums in High- vs. Low-Group	33
5.5	Regression Results on Protest Amount P	36
5.6	Probit Model Results on “Allow”-Probability	38
5.7	Preference Reversals	41
6.1	Treatments with Framing Variations	48

1 Introduction

In today's globalized and widely connected world, immigration remains an important subject of current research due to its economic and social consequences that impact societies as a whole. Climate change or intensifying military conflicts are global events that will produce increasing migration flows across the world. Hence, future immigration strategies must be able to tackle the upcoming challenges and steer immigration such that it is widely accepted by the public. For this purpose, it is crucial to understand how natives view immigration and how they form attitudes towards it. If decision makers understand the reasons behind immigration attitudes, it could generate more effective political discourses and efficient policies.

So, what are the driving forces behind immigration attitudes? Obviously, economic concerns play a role in people's reasoning while deciding to accept or reject certain types of immigration. Although immigration seems to have a positive net effect on the economy (Dustmann and Frattini (2014)) which mainly stems from young immigrants (Storesletten (2003)), people might be afraid of negative impacts on the labor market or perceive immigrants as a fiscal burden to the state (Alesina and Tabellini (2022)). Besides material interests, social and cultural aspects are important as well since immigrants and their lifestyle can be perceived as foreign and give rise to prejudices (Hainmueller and Hopkins (2014)). Xenophobia seems to be a significant factor to impact people's stance in questions about immigration (Dustmann and Preston (2007), Hjerm and Nagayoshi (2011)) and is increasing with the presence of immigrants in the country (d'Ancona and Ángeles (2016)). Hence, economic and cultural reasons are identified as key reasons behind people's perspective on immigration and most likely interact with each other.

Despite these well-known drivers, it is important to explore other aspects that might contribute to shaping immigration attitudes to better understand the more subtle causes behind certain views and get a clearer picture overall. There are only a few studies that try to uncover behavioral effects in an experimental setting with specific connection to immigration. Haaland and Roth (2020) analyze if people change their attitudes when research evidence on immigration is given while Florio (2022) extends the analysis to personal contact with immigrants. Khadjavi and Tjaden (2018) investigate migrant's power of self-determination and public debate in an in-group vs. out-group setting with a public good game.

Interestingly, little attention has been paid to a factor that played a major role in the successful Brexit campaign: The feeling of control loss. Brexiteers were frequently proclaiming to "take

back control” over the country and its borders, especially in the context of immigration. The perceived feeling of uncontrolled immigration flows might have had a significant impact on the voting behavior of the British population despite economic or cultural concerns. As Wright (2016) describes, governments need to send control signals to the public to resolve the dilemma between looser economically beneficial immigration policies and the populist pressure for stricter controls. Hence, it seems that the understanding of immigration views can be enriched by insights from possible control preferences.

The aim of this thesis is to gather experimental evidence on the impact of perceived control loss on immigration attitudes. In an experiment, we will try to find out if subjects prefer to control certain outcomes against their own economic interest. The experimental design contains similar aspects of other studies that explore the well-documented existence of a control premium that subjects are willing to pay to get control over their own or other people’s outcomes (Bartling et al. (2014), Bobadilla-Suarez et al. (2017), Fehr et al. (2013), Ferreira et al. (2020), Lübbecke and Schnedler (2020)). In particular, the method on distinguishing between the material value of a choice and the intrinsic value of control relies on the comparison of certainty equivalents. By eliciting equivalence points in a neutral setting and in a scenario where subjects perceive to be in control, one can compare these valuations and draw conclusions about control preferences. If the latter is higher, subjects seem to attach a value to control.

However, the context of this experiment differs from existing literature: Subjects are exposed to a framing that describes their choices as immigration decisions in an abstract sense. Participants are “citizens” and decide if “non-citizens” can join their “state” or not. The feeling of lost control is induced by overruling the subject’s decision repeatedly while giving it the opportunity to pay for “taking back control” over the choice. Varying economic conditions in a choice sequence reveal if subjects are willing to pay higher prices for control than rationality would predict by comparing the elicited certainty equivalents. This might be an indicator that the desire for control also plays a role in immigration choices.

Indeed, the results show that subjects value control over immigration choices by paying a control premium independently of the direction of their choice (allowing or not allowing the “non-citizen” to join). Interestingly, a minority with particularly strong control preferences is responsible for the overall significant premium. Moreover, spillover effects from control loss seem to influence the decision behavior while social and risk preferences play a role as well.

The experimental evidence suggests interesting real-world implications. Immigrants might not only need to add an economic gain to the target country but a sufficiently significant one that

potential concerns about immigration control do not play a role in people's view on immigration. Furthermore, perceived control loss in the past can still influence today's views on immigration and political actors can use this fact to push their agenda. In general, control preferences should be considered by designing immigration policies and political discussions.

In the following chapter, the existing literature with regards to control preferences will be summarized and connected to the novel context of this experiment. In chapter 3, the experimental design and procedures will be outlined. The hypotheses are formulated in chapter 4 which is followed by the presentation of results in chapter 5. Chapter 6 discusses the design aspects of the experiment and potential improvements for further research. Finally, chapter 7 summarizes the main aspects and contains some concluding remarks.

2 Literature Review

The idea that people assign an intrinsic value to control that goes beyond its instrumental or economic worth is not a particularly new one. Langer (1975) already described the “illusion of control” as the preference of triggering a random event oneself instead of another person, partly due to the belief that one can influence the outcome. The main conclusion is that control and decision rights must seem valuable even if they do not give an obvious or objective benefit. Recent studies investigated this bias further by conducting experiments with variations of games where subjects could gain control over obvious random events.

In the experiment of Bouacida and Foucart (2022), subjects prefer to place the bet on one out of two outcomes in a lottery themselves rather than letting a computer determine it. Sloof and von Siemens (2017) observed that a third of all subjects would rather do an uninformed choice between two tasks by themselves than letting another participant do it while 75 % are willing to pay a small amount for this right. In contrast to that, Charness and Gneezy (2010) find that only 10 % of participants are willing to pay to roll a die themselves to determine the value of an asset than letting the experimenter do it. They conclude that control preferences may not be as important in financial decisions such as portfolio choice.

Besides the described phenomenon that people prefer to keep “useless authority” (Sloof and von Siemens (2017)) in simple gambling tasks, control preferences play a role in more elaborated economic contexts as well. Many studies use delegation games to investigate the principal’s desire to keep control although delegating the task would be economically beneficial. Indeed, delegation rates are lower than rational individuals would exhibit, also when controlled for risk preferences that are relevant due to the uncertain nature of the agent’s actions in case of delegation (Fehr et al. (2013), Hausfeld et al. (2020)). Furthermore, subjects are willing to pay a control premium to keep control which is observed in samples from Switzerland (Bartling et al. (2014)), France and Japan (Ferreira et al. (2020)). Bobadilla-Suarez et al. (2017) extend the analysis to the loss domain and observe that gains and losses are equally affected by control preferences. Additionally, subjects are equally interested in keeping control over unfair or fair outcomes in decisions that involve payments to other people (Hausfeld et al. (2020)).

Besides delegation games, the existence of a control premium is also observed in a variety of other experimental designs: A game to bet on their own or other people’s performance in solving a quiz (Owens et al. (2014)), an auction to bid on decision rights (Neri and Rommeswinkel (2017)) or a puzzle game where one subject can overrule the submitted solution

while the other one can resist this interference (Lübbecke and Schnedler (2020)). While these experiments analyze situations where subjects want to keep control over their own outcomes, Pikulina and Tergiman (2020) designed the “power game” where subjects can pay to receive control over other people’s payoffs. They also observe a willingness to pay for control over others that exceeds the economic value of such power.

Although it is a common finding that subjects attach an intrinsic value to decision rights or control, there is no clear picture about the motives behind such preferences. The observed misconception that people think they can influence random events in a favorable way by triggering them (Bouacida and Foucart (2022), Langer (1975), Sloof and von Siemens (2017)) is claimed to be not the only reason behind the preference for control (Bartling et al. (2014), Owens et al. (2014)). Researchers tried to investigate different explanations behind the desire for control which could stem from avoiding disutility caused by actions from other parties or from utility that is inherent to exercising control. However, the results are mixed and do not deliver direct evidence of drivers behind control preferences.

While some researchers identify disutility from interference as the main reason (Fehr et al. (2013), Neri and Rommeswinkel (2017)), other studies do not confirm that independence from others is explanatory (Ferreira et al. (2020)). In the experiment of Pikulina and Tergiman (2020), subjects are willing to pay for power over others although other researchers could not observe a desire for power (Ferreira et al. (2020), Neri and Rommeswinkel (2017)). Lübbecke and Schnedler (2020) suggest that people want to signal authorship of a solution or outcome which causes resistance against interference from others. Ferreira et al. (2020) find a significant effect from the feeling of self-reliance that causes people to retain control over their outcomes.

While the experiment of this thesis does not include a treatment that could clearly contribute to the research of underlying motives, it contains similar design aspects from other studies that concern the control premium. The measuring method probably comes closest to the delegation experiment from Bartling et al. (2014) that was replicated by Ferreira et al. (2020). In the delegation game, principals indirectly determine two lotteries that should be equivalent in utility if rationality prevails: One lottery where they are in charge and one where they delegate the task. At a later stage of the experiment, the equivalence points of these two lotteries are elicited in a neutral context. If the valuation for both lotteries is not the same, one can draw conclusions that being in control adds some value to the lottery compared to delegating the task.

In principle, our design follows the same measuring approach since at the end of the experiment, equivalence points for three lotteries are elicited in a neutral context. However, the lotteries are

not derived from the subject's choices but are predetermined and presented in the choice sequence in an immigration context. Nevertheless, we also use the equivalence points of the lotteries as a benchmark to see if people pay irrational prices to keep control over their choices which will be described in further detail in the next chapter. Since the lottery in our experiment also triggers payments for other subjects, the elicited equivalence points do not only include risk but also social preferences.

Furthermore, the nature of the game differs from delegation games since principals can decide themselves to delegate a task and risk to be disappointed by the outcome while in our design, subjects are overruled without having any say in it. It is comparable to the experiment of Lübbecke and Schnedler (2020) in which a subject can submit a solution to a logical puzzle while another participant can interfere and replace it with another solution. Finally, the overruled subject has the possibility to resist this interference by paying a price. In our experiment, subjects also have the possibility to protest the overruling from the computer but there is a key difference: While in our experiment, the overruling is completely arbitrary and does not follow any reasoning, the interference in the experiment of Lübbecke and Schnedler (2020) happens due to actual concerns about the submitted solution. Hence, the interference from participants is meant to be constructive while ours seemingly does not serve any purpose from the perspective of the deciding subject other than denying its preferred choice.

Finally, the experiment differs in the suggested context which is meant to feel like an immigration choice and contains a stronger framing than in most presented literature. This aspect is elaborated further in the next chapter which describes the experimental design.

3 Experimental Design

The experimental setup follows a within-subject design and consists of three parts. The first part is a real effort task where subjects need to correctly identify images to earn some money. After correctly identifying enough images, subjects proceed to the main part of the experiment which contains the immigration decisions. The choice sequence lets subjects decide if another participant can join their “state” or not and triggers economic consequences for both parties which consist of lotteries or safe payments. Some decisions will be overruled by the computer while subjects can protest the interference by paying a price. In the third and final part, the certainty equivalents of three lotteries from part 2 are elicited in a neutral context. The following paragraphs will describe each part in more detail before the experimental procedures are outlined.

3.1 Experimental Game

Part 1: Real Effort Task

In the instructions, participants are referred to as “citizens” who can work for their “state” and take decisions that affect it. In part 1, subjects do a real effort task following the design from Koch et al. (2022). In particular, they individually face 77 image identifying tasks in which they have to choose the correct image for a given label (Figure 3.1). The more tasks the subjects can solve within the time limit of 5 minutes, the more points they can earn up to a maximum of 20 points (EUR 5). Additionally, only by hitting the threshold of 50 correctly identified images, the subjects were allowed to make payoff relevant decisions in part 2 and earn additional points. For each label, subjects are allowed to choose incorrectly once and are prompted to try again. The image identifying tasks were intended to be sufficiently easy such that no individual should have problems getting the maximum amount of points.

After finishing all 77 image identification tasks or when the time is up, the earned points are displayed and subjects proceed to part 2 if they could at least successfully identify 50 images. If a subject failed to hit the threshold, they could still make the decisions in part 2, albeit only hypothetically without earning any points. However, this case did not occur in either of the two sessions.

Image identifying - (1/77)!

Time left to complete this page: 4:53

What image depicts a **LAMP** (deutsch: **Lampe**)?



Figure 3.1: Image Identifying Task

The image identification task was implemented to induce a feeling of entitlement for a reward in part 2. The instructions described repeatedly that subjects “work” for “their state” and only because of their successful performance, they are able to profit from it and decide who can enter and who cannot in part 2. Subjects might feel like a native who built up their own state and have to decide over benefitting other players that did not contribute to it. Hence, the intended effect was that subjects take the decisions seriously and might perceive it as an immigration decision in an abstract sense.

Part 2: Immigration decisions

In the main part of the experiment, part 2, subjects face a sequence of 15 decisions where they have to choose between a lottery and a safe payoff. By choosing the lottery, they allow another subject that was assigned to them – which is called “non-citizen” in the instructions – to join their state. If the subject chooses the safe payoff instead, the “non-citizen” cannot join the state. The probabilities and outcomes of the lottery as well as the safe payoff vary across the decisions while the payments for the “non-citizen” are always the same – she receives 10 points for joining the state and 0 points for not being allowed to join. Independently of the decision, the “citizen” always receives 10 points per decision which is communicated to the participants as a reward from their performance in part 1.

The described decision between a risky and a safe option shall symbolize that economic gains from immigration can be uncertain and need to be assessed by the citizens of a country to make a choice (Dustmann and Preston (2006)). On the other hand, the economic conditions and well-being for immigrants are usually better in the country they moved to compared to their home country (Jasso et al. (2000)) which is reflected in the obvious profit for the “non-citizens” by

joining the state. The deciding players have to consider their own gain of immigration and the benefit for immigrants by accepting them.

Before the subjects start the sequence, they are shown an example scenario and are asked a few comprehension questions to make sure that the task is clear. In all decisions, the described payoffs for the deciding player and for the other player of the particular scenario are clearly displayed on the screen. For an example choice situation, see Figure 3.2. As indicated before, “non-citizens” receive 0 points in case they are not allowed to join but get 10 points if the “citizen” accepts them. These payoffs for “non-citizens” do not vary across the choice situations. In this example scenario, the “citizen” has to choose between the safe payment of 0 points by not allowing the other subject to join and the lottery which would allow the subject to join. Here, the lottery has a probability of 30 % to produce the good outcome with 10 points and the bad outcome with a probability of 70 % which means a loss of 25 points. Notably, safe payments and lotteries are varied across the different scenarios. A pie chart which visualizes the probabilities of the lottery is added to avoid misinterpretation about the percentage representation. Furthermore, two images illustrate the consequences of allowing or not allowing the “non-citizen” in the state. Their purpose is to strengthen the perceived context of subjects deciding over their own state.

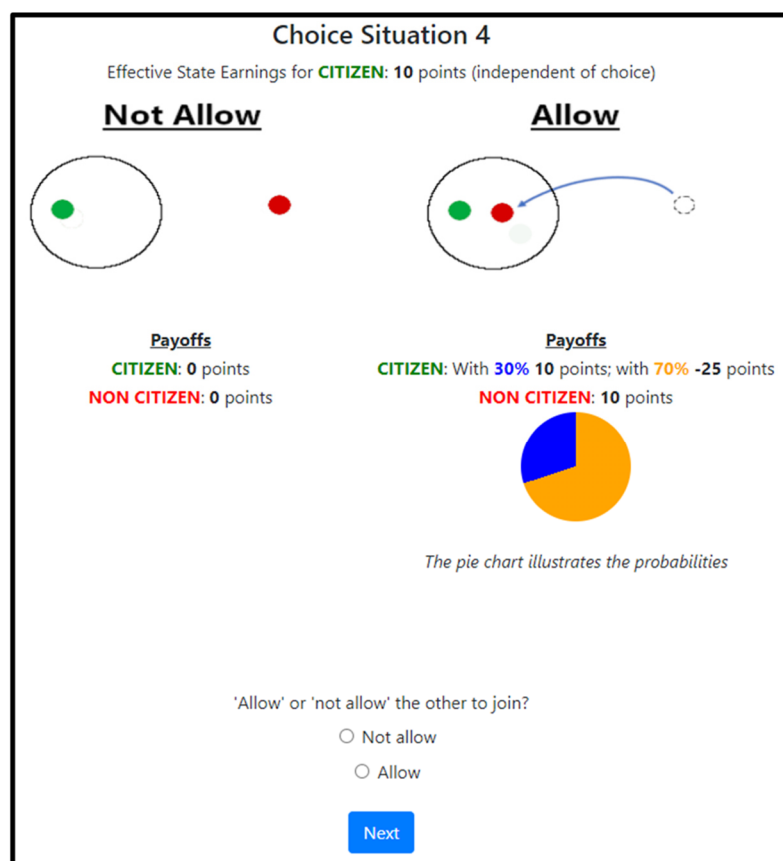


Figure 3.2: Example Choice Situation

After the subject decided to allow or not allow the other player to join, there are two possibilities how the experiment can proceed. Firstly, the decision is confirmed and the subject faces the next choice situation. Alternatively, the decision gets overruled by the computer which means that the opposite of the initial choice gets implemented. For example, the subject chose the option “Allow” but the computer changes the decision to “Not Allow”. The treatment of overruling decisions is meant to create a feeling of lost control for the subjects. Notably, it is not known initially that decisions can be overruled. In addition, to familiarize subjects with the decisions they have to make, the initial three decisions are never overruled. Choice situation 4 represents the first choice that is overruled and comes as a surprise for subjects. After participants have made an (initial) decision in choice situation 4, detailed instructions on how the new overruling mechanic affects the remaining decisions are displayed and the options that arise for the subjects are outlined. For an example of the alert message for an overruled decision, see Figure 3.3.

Choice Situation 4

ATTENTION Your decision was overruled!

In the last choice situation (see below), you decided **NOT to allow** the **NON CITIZEN** to join (i.e. you went for the safe payoff). However, your decision was **overruled**. This means, that now the **NON CITIZEN** is **allowed** to join your state (i.e. you will get the lottery), which is the **opposite of what you decided initially**.

As indicated before, you have the option to **protest** the overruling. What is the **maximum amount** of points you would be willing to invest to switch back to your preferred option “Not Allow”?

[You have to choose points between 0 and 25 and can use non-integer values such as 4.5 or 17.3]

(If you type in a price of 0, so no protest, the allow/lottery-situation (and not your original choice) will be implemented for sure. A random price mechanism will be used to determine whether your protest is successful. This gives you an incentive to state your personal maximum as you only have to pay the potentially lower price or nothing at all in case your protest is unsuccessful.)

Maximum amount you would invest to go back to Not Allow?

Next

Figure 3.3: Alert Message for an Overruled Decision

Before the computer implements the opposite of the subject’s choice, subjects always have the option to protest against this potential overruling by entering the maximum amount of points they would be willing to give up to switch back to their initial choice. If the chosen amount is higher than a randomly drawn price, the protest is successful and the preferred option is implemented. If the amount is lower, the protest fails and the overruled decision stays enacted.

The subject is incentivized to be truthful about their maximum willingness to pay since they only pay the randomly drawn price if their protest is successful, even if their chosen amount was higher. In case of a failed protest, the player does not lose any points. Subjects can enter an amount between 0 and 25 points.

Overall, subjects face fifteen choice scenarios out of which seven certainly get overruled while the other choices are directly implemented according to the decision of the subjects. For a full list of the fifteen scenarios and their payoffs, see Table 3.1. In the second column, the lotteries are denoted in the following way: $(X, p; Y, (1 - p))$ with p being the probability for the outcome X and $(1 - p)$ as the probability for outcome Y . Outcome X is always strictly better than outcome Y .

The column *EV* shows the expected value of the lottery for the deciding player and in the fourth column, the safe payoff for the “Not Allow”-option is displayed. By comparing the expected value with the outside option, one can see that certain lotteries are particularly attractive while others are not as desirable. Some scenarios are identical, namely scenarios A, B and C which is why they occur twice in the column *Scenario*. Other scenarios have identical lotteries, but the safe options are varied. These scenarios have the same letter but marked with a prime, for example E, E' and E'', since there are three variations of the safe option for the same lottery in scenario E. The column *Position* says at which position in the sequence the scenario will be shown to the subject and the last column shows if the decision is subject to the overruling treatment or not.

Scenario	Lottery (Allow)	EV	Safe Option (Not Allow)	Position	Overruled
A	(10, 0.7; -2, 0.3)	6.4	7	Random (1,3)	NO
B	(6, 0.4; -4, 0.6)	0	1	Random (1,3)	NO
C	(10, 0.6; -10, 0.4)	2	3	Random (1,3)	NO
D	(10, 0.3; -25, 0.7)	-14.5	0	4	YES
E	(15, 0.8; -5, 0.2)	11	12	5	YES
F	(20, 0.5; -15, 0.5)	2.5	3	6	YES
B	(6, 0.4; -4, 0.6)	0	1	7	NO
A	(10, 0.7; -2, 0.3)	6.4	7	Random (8,15)	NO
C	(10, 0.6; -10, 0.4)	2	3	Random (8,15)	NO
D'	(10, 0.3; -25, 0.7)	-14.5	-12	Random (8,15)	YES
E'	(15, 0.8; -5, 0.2)	11	8	Random (8,15)	YES
E''	(15, 0.8; -5, 0.2)	11	4	Random (8,15)	YES
F'	(20, 0.5; -15, 0.5)	2.5	-3	Random (8,15)	YES
G	(10, 0.0; 0, 1.0)	0	0	Random (8,15)	NO
H	(11, 0.95; -2, 0.05)	10.35	8	Random (8,15)	NO

Table 3.1: Sequence of Choice Situations

The positions 4 to 7 always contain the same scenarios while all other positions are placed following a randomization rule. The rule is denoted with *Random* (t, T) which means that the scenario can be between and including position t and T . The randomization rules and fixed positions are specifically designed to enable a row of three consecutive overruled decisions at position 4 to 6, followed by scenario B which already appeared somewhere at the first three positions. This fixed order has the purpose to increase the potential effect of the overruling treatment by allocating more overruled decisions to its first occurrence. The density of overruled choices may also lead to a different behavior for the following scenario B compared to the identical scenario at the start of the experiment.

To summarize the whole sequence of 15 decisions, there are three identical scenarios A, B and C that are once presented before the first overruling treatment and once after it. This allows us to test whether the overruling and loss of control has any influence even on scenarios that are not overruled. Notably, scenario B is a special case since it is directly attached to the first series of overruled choices at position 7. The scenarios D, E, F and their variations – where only the outside options are varied while the lottery remains the same – are the only decisions which get overruled. While the first occurrences of the overruled scenarios are fixed, the other ones are placed randomly in the second half of the sequence. Scenarios G and H are neither subject to the overruling treatment nor directly comparable to other scenarios.

The rich design of the choice sequence was chosen to enable an explorative approach by observing the decision behavior. Due to the novel design of the experiment, new behavioral aspects may arise that were not expected before. Besides that, several reasons led to this specific sequence structure that are mainly concerned with an efficient way of measuring the treatment effect of the loss of control. The first three scenarios were placed at the beginning such that the participants could get familiar with the task without bothering about the overruling treatment. After that, three consecutive overruled decisions should serve as a surprise moment and an intense confrontation with the unpleasant interference that was intended to enhance a potential treatment effect. Scenario B was placed directly afterwards to see spillover effects and to compare it to scenarios A and C that have random positions. The second half of the experiment was randomized to analyze if positional effects play a role and to test if a treatment effect exists independently of constructed sequences like in the first half.

After finishing all 15 decisions, including the determination of the maximum amount the subjects are willing to invest into the protest, part 3 of the experiment starts.

Part 3: Elicitation of Certainty Equivalents

In part 3, subjects are asked to choose between a lottery and a safe payoff – which are also payoff relevant – to determine their certainty equivalent (CE) of the respective lottery. While choosing either the safe payment or the lottery, the safe payment gets varied depending on the previous choice while the lottery remains identical. This eventually leads to a value where the subject is indifferent between the lottery and the safe option. The value of the safe payment is then assumed to be the CE for this subject which represents the reservation price to enter the respective lottery. The procedure is repeated for three lotteries in a randomized order, such that three CEs per subject are elicited. The safe payments for each lottery are varied in 6 consecutive steps which results in 18 decisions that the subject has to make in part 3.

The three lotteries are identical to the ones used in all the overruled scenarios D, E, F and their variations in part 2. Additionally, the independent payment of 10 points for the deciding player and the 10 points for the other player by choosing the lottery are replicated. Hence, the decisions between the lottery and safe options are economically identical between part 2 and part 3, except for the different context. While the context in part 3 is neutral with no overruling treatment, the decisions in part 2 were framed as an immigration decision to allow the “non-citizen” in the state or not.

The purpose of part 3 is to elicit the subject’s perceived economic value of the lotteries, including the payment to the other player, in the decisions of part 2 that are subject to the overruling treatment. Subjects could enter a high protest amount in part 2 to switch back to their preferred option which lowers the resulting payoff of this option (lottery or safe payment) because they have to pay a price in case of a successful protest. Since the CE of the respective lottery is elicited in part 3, one can directly see if a subject decides against its material interest. For example, if the subject successfully protests the overruling and switches back to the preferred lottery, it might be the case that the resulting payoff is lower than the subject’s CE of that lottery. Therefore, the CEs serve as a benchmark that can be used to see if the observed protest amounts in the overruling treatment exceed the threshold that a rational subject would choose.

3.2 Procedures

The experiment took place in two sessions in January and March 2022 at the Vienna Center for Experimental Economics (VCEE) at the University of Vienna. Since the experiment requires

the assignment of pairs, the number of participants in each session has to be divisible by two. If an uneven number of subjects were present, a randomly chosen individual received the participation fee and was asked to leave the experiment before the start. This resulted in 40 participants in total divided into two sessions.

The participation fee of EUR 5 is the minimum that subjects could earn in the experiment. On top of that, it is possible to earn significantly more depending on the performance and luck of the individuals as well as the decisions of the subjects they were paired with. In the experiment itself, the currency of potential payments was called “Points” and was introduced in the instructions with an exchange rate EUR / Points of 1:4. On average, participants earned EUR 16.05 which was paid out in cash at the end of the experiment.

Each session lasted approximately 45 minutes. The instructions of the experiment were always displayed on screen of each subject before they became relevant in a task or decision. Additionally, explanatory introductions were read aloud at the start, during and at the end of the experiment. All three parts were played in the same order but the exact details about each part were only revealed at the start of them. As described before, the second part of the experiment contains a surprise element which is not anticipated by the subjects but is well explained in the instructions once it occurs and before a decision by the subjects takes place. After the three described parts of the experiment, a sociodemographic survey with several questions about immigration is answered at the end. The instructions can be found in the Appendix.

To increase meaningful observations, all subjects initially are “citizens” and make decisions as such throughout the whole experiment. Only at the end of the experiment, the roles for the paired individuals are set with one subject being the “citizen” while the other is the “non-citizen”. From the “citizen”, three out of fifteen decisions from part 2 are randomly chosen to be payoff relevant while three out of 18 decisions ($3 \text{ lotteries} \times 6 \text{ steps} = 18$) from the “non-citizen” in part 3 are payoff relevant. If lotteries were chosen in either of the parts, they are played out and the outcomes are displayed on a summary page. Figure 3.4 shows the summary for the part 2 decisions from the perspective of the “citizen” player. The final payment therefore consists of earnings from (i) the real-effort task (ii) three decisions of part 2 and (ii) three decisions of part 3.

Payment for decisions in Part 2:

For the choices in Part 2, the computer **randomly selected you** as the one deciding

The following 3 scenarios were randomly selected: **1) Scenario M , 2) Scenario A , 3) Scenario E** (see table below for details)

For these three scenarios, the computer randomly draw the following lottery outcomes:

- 1) **good outcome!**
- 2) **bad outcome!**
- 3) **bad outcome!**

The following decision were made:

- 1) **NOT allow**. This decision was **overruled** though (your protest was unsuccessful)
- 2) **NOT allow**
- 3) **NOT allow** This decision was **overruled** though (your protest was unsuccessful)

This implies that you earn:

- 1) **15 points**
- 2) **7 points**
- 3) **-5 points**

Figure 3.4: Summary of Outcomes for Part 2

After outlining the experimental design, certain assumptions about the expected behavior will be presented in the next chapter.

4 Hypotheses

Since there are hardly any insights about the impact of a perceived control loss in immigration decisions, the main hypotheses try to connect well established economic concepts with the novel design of the presented experiment. Firstly, the subject's social preferences with regards to the experimental context will be the focus. The second part is concerned with the existence of a control premium which subjects might be willing to pay to implement their preferred choice. Lastly, the general decision-making behavior, especially under the overruling treatment, will be looked at more closely.

4.1 Other-regarding behavior

The experimental design might give a hint about the subject's other-regarding preferences in two ways: Firstly, the elicited certainty equivalents from part 3 potentially show if the subjects take the payment to the other player into account or not. Secondly, the acceptance rate of "non-citizens" might reflect the subjects' social preferences with regards to the different choice situations, especially for those where no CEs were identified.

Certainty Equivalents

In part 3, subjects are asked to choose between a lottery and a safe payment which have the same implications as the choices in part 2: The other player receives a payment if the lottery is chosen while she gets nothing if the safe payment is taken. Therefore, the elicited CE in part 3 not only represents material gains from the lottery but could also reflect risk and social preferences.

In standard economic theory, an agent is classified as risk-seeking if their CE exceeds the expected value (EV) of a lottery, as risk-neutral if the two measures are equal and as risk-averse if the CE is below the EV. While using different tools to measure risk preferences, most research finds that the majority of people are risk-averse as summarized by Harrison and Rutström (2008) or Eckel and Grossman (2008) who focus on gender differences in risk attitudes. Therefore, one would suspect the average CEs of the standalone lotteries, without considering the payment to the other player, to be below their respective EVs. This is also in line with experimental data that relies on the elicitation of CEs, like Hartog et al. (2002) who show that most subjects are risk-averse to risk-neutral while only a little fraction is risk loving.

For the following argumentation, *EV* is defined as the expected value of the lottery for the deciding player and therefore only represents selfish payoffs while the elicited certainty equivalent *CE* from part 3 reflects risk preferences and can also include utility from benefitting others. Based on the assumption that most subjects are risk-averse, an average *CE* below the *EV* would not indicate other-regarding preferences. Hence, subjects that care about the payment for the “non-citizen” would exhibit a *CE* that would at least approach or exceed the selfish *EV* of the respective lottery. The experiment does not allow to distinguish between risk-seeking behavior and social preferences if the *CE* is indeed higher than the *EV*. Hence, a certainty equivalent for a lottery l that satisfies $CE(l) \geq EV(l)$ could reflect a compound of social and risk preferences.

The proposed threshold of *CE* exceeding *EV* is further backed by other empirical studies that have a similar experimental design compared to the choice situations here. In simple dictator games where the deciding player is exposed to more risk by allocating more to the other player, the subjects share significantly less compared to the standard game (Brock et al. (2013), Cappelen et al. (2013), Owens (2016)). This means that other-regarding behavior could be suppressed in a risky environment compared to decisions where no risk is present. In this experiment, subjects are also exposed to more risk while enabling a payment for the other players. Hence, the *EV* of the standalone lottery can be seen as a conservative benchmark because the utility from benefitting the other player might be decreased which diminishes the *CE*.

Acceptance rate

Since the only possibility to enable a payment for the other player is to accept them into the state, other-regarding behavior could be reflected in a high acceptance rate. Only five out of fifteen choice situations would possibly be accepted by an agent who strictly maximizes selfish EVs while ignoring the other person’s payoff. Assuming that extremely high risk aversion is rare and that other-regarding preferences are present in the population, the lower bound of an overall acceptance rate would therefore be: $\frac{\text{"Allow" choices}}{\text{Total choices}} = \frac{5}{15} = 0.33$.

As described above, subjects tend to care less about the well-being of other people when they need to take more risks while doing so. Owens (2016) observes that the sharing behavior under risk is best described by risk-averse preferences which means that subjects avoid very risky lotteries even if they would like to benefit the other player. In this experiment, only two decisions contain a significantly risky lottery that has a highly negative EV (scenario D and

D'). The vast majority of choice situations are designed to have a negligible difference between their EVs and safe payoffs which simplifies the acceptance of “non-citizens” if social preferences play a role.

Hence, the predicted high values for CE should translate into high acceptance rates which could be caused by other-regarding behavior or risk-seeking preferences. Combining the two predictions, the first hypothesis is stated as follows:

H1: *The aggregate acceptance rate is significantly higher than 33 %.*

In case that the hypothesis is confirmed, it could suggest that subjects show other-regarding behavior towards “non-citizens” in an abstract immigration context, even if they have to accept risky outcomes for themselves.

4.2 Control Preferences

The main goal of this thesis is to identify the potential control premium that subjects are willing to pay to switch back to their preferred option after the overruling treatment. The deciding player might be influenced by all sorts of sentiments when his decision is overruled and the perceived loss of control could cause the subject to protest heavily by entering a substantial protest amount in part 2. At first, a hypothesis about the protest amount and its magnitude will be formulized that looks closer at potential reasons behind its evolution. After that, a quantitative definition of the control premium will be derived to identify a potential intrinsic value that subjects attach to gaining back control over their choice.

Protest Amount

By choosing the lottery l or the safe payment s in the immigration choices, the individual clearly displays a preference or is indifferent between the two options. For example, if the lottery is chosen, in terms of a utility function $U(x)$, it must hold $U(l) \geq U(s)$. In the 7 choice situations that get overruled, the subject has the possibility to enter a protest amount P to protest against this overruling. The higher P , the more likely it is that the overruling is undone and the initial choice is implemented while it also decreases the payoff of the preferred option. Hence, a rational subject would choose at maximum the amount $P' = |U(l) - U(s)|$ which is the absolute difference in utility between the two options. Intuitively, a subject wants to move to its desired option but should only pay as much as it would win from the switch. If it would pay more than the potential gain – which means it chooses $P > P'$ – it would worsen its situation

since it would be better to stay with the enforced choice rather than accepting a lower utility from the desired option. If the subject is indifferent, it does not want to pay a price to switch to an equally wanted option and would choose $P = 0$.

Following the described reasoning, a rational subject would choose an amount from the interval $P \in [0, P']$ such that it would still prefer the reduced payment from the desired option instead of the choice that was imposed on it. If the rational price P' increases, the subject should naturally choose a higher value for P since the potential utility gain is higher if the protest is successful. Experimental results validate the assumption that subjects are willing to pay more to get control over the outcome if the stake size increases (Bartling et al. (2014)), if the preferences are strongly directed towards one outcome (Hausfeld et al. (2020)) or if the subjects are more selfish (Gawn and Innes (2019)). Therefore, one would observe a higher protest amount P , the more attractive the preferred option is compared to the undesired outcome.

Despite the material difference between the lottery and the safe payment, other factors could influence the magnitude of P . The scope of the experiment is a perceived loss of control that potentially impacts the evolution of the protest amounts. It is well documented that individuals assign an intrinsic value to control (Bartling et al. (2014), Bobadilla-Suarez et al. (2017), Fehr et al. (2013), Ferreira et al. (2020), Lübbecke and Schnedler (2020), Owens et al. (2014), Pikulina and Tergiman (2020) and others) while mixed evidence is provided regarding the motives behind this valuation. Possible reasons include “disutility from being overruled” (Fehr et al. (2013)) and a preference for “non-interference from others” (Neri and Rommeswinkel (2017)) that explain the utility derived from control. The experimental setup gives the possibility to explore the described phenomena since the decisions of subjects get overruled throughout the whole sequence. The recurring interference with the subject’s choices might lead to a strong reaction that is reflected in the size of P . One could speculate that P increases with the number of past overruled decisions since the subject could show accelerating aversion towards the overruling and therefore reinforce its resistance towards the interference.

Summarizing the described motives that could impact the protest amount P , the resulting hypothesis is the following:

H2a: *The protest amount P increases in (i) the utility difference between the lottery and the safe payment as well as (ii) the number of past overruled decisions.*

Since the rational upper bound for P has been defined as P' , an overshooting of this threshold might indicate a value for control beyond the material gain of the preferred outcome. The mark-up on the protest amount will be defined as the control premium in the following paragraphs.

Control Premium

To provide evidence on the existence of a control premium, it must be shown that the subject is not acting in its best material interest by choosing an excessive protest amount to gain back control. As introduced before, a rational subject with a linear utility function would not pick a higher protest amount than the rational price $P' = |U(l) - U(s)|$, which is the absolute utility difference between the lottery l and the safe payment s . To calculate P' , utility values for both the lottery and the safe payment have to be defined. For linear preferences, it holds $U(s) = s$ while $U(l)$ can be defined by using the CE from part 3. By stating $U(l) = CE(l)$ as a benchmark for the subject's utility from the lottery, one would follow a similar approach like Bartling et al. (2014) who have used the difference in CEs between lotteries to identify the control premium. Therefore, the control premium CP exists if a subject is willing to pay a protest amount P that is higher than $P' = |CE(l) - s|$ which is the absolute difference between the elicited CE of the lottery and the safe payment of the respective choice scenario. It follows the hypothesis with regards to the control premium:

H2b: *Subjects choose irrational protest amounts $P > P'$ and therefore attach an intrinsic value to taking back control over the decision by paying a control premium $CP = P - P'$.*

If the stated hypothesis will be validated not only on the aggregate but also on the individual level needs to be investigated. Past studies show that there is heterogeneity in individuals' preferences for control (Bartling et al. (2014), Ferreira et al. (2020)) which could result in significant differences between the subjects in this experiment.

After describing the expectations regarding the valuation of control, it is also important to look closer at the decisions themselves and what are the potential reasons behind allowing or not allowing the “non-citizen” to join.

4.3 Immigration choices

Despite the motivation to protest an overruling to implement the subject's preferred choice, it is important to understand the drivers behind the initial decision itself. Similar to the arguments with regards to the protest amount, the decision to allow the “non-citizen” into the state should

be influenced by the material benefit that the “citizen” can expect but could also be impacted by the overruling treatment and the experienced control loss. At first, the general reasoning behind the decisions will be looked at more closely. Secondly, the focus will lay on the three scenarios A, B and C that might be subject to preference reversals due to the overruling treatment.

Choice preferences

In the decision sequence that the subjects face, different reasons could cause the final choice to accept the lottery or the safe payment. Social preferences play a role since the other player only receives a payoff if the deciding player accepts the lottery. Furthermore, the subject has to take risks to benefit the “non-citizen” which also includes risk preferences that have to be taken into account. Finally, the individual’s material benefit should impact the decision if the subject does not have extreme preferences like pure altruism or very strong risk aversion. In general, the higher the expected value EV of the lottery compared to the safe payment, the more likely it should be that the subject accepts the other player.

Considering the context of the experiment that provides an abstract setting of a state with “citizens” and outside subjects that are “non-citizens”, the participants might subconsciously include their attitudes towards immigration in their decisions. Haaland and Roth (2020) find that subjects have a more positive attitude towards immigration if their concerns about adverse labor market impacts from immigrants are mitigated. In an abstract sense, this would confirm the expected behavior since increased expected benefits should result in a higher acceptance rate of the “non-citizens”.

Besides the obvious material motivation, it is interesting to investigate if the overruling treatment has an impact on the likelihood to allow the other subject to join the state. Subjects might deviate from their usual behavior after experiencing a loss of control and decide differently in upcoming choices. In contrast to the protest amounts which are assumed to be higher, there is no clear expectation about how the subject’s behavior could change with regards to their choices to allow or not allow the other player to join. For example, the subject might deduce that the unpleasant overruling happened because of the particular choice it has made and therefore chooses the other option more often. This would be in line with the idea of disutility from being overruled by Fehr et al. (2013) and the preference for non-interference by Neri and Rommeswinkel (2017). On the other hand, an individual might act out of spite and choose the overruled option even more to resist against the overruling. Hence, it is not clear in

what direction the behavior could deviate but it is plausible to assume that the likelihood to allow the “non-citizen” to join might depend on overruled decisions in the past. Condensing both the material benefit and the potential impact of control loss, the hypothesis with regards to the choice preferences is stated as follows:

H3a: *The likelihood of subjects allowing the “non-citizen” to join (i) increases with the difference ($EV - s$) between the expected value of the lottery l and the safe payment s and (ii) depends on the history of overruled decisions.*

Moving away from the general view on all scenarios, the following paragraphs focus on three scenarios that might reveal inconsistent behavior after the overruling treatment.

Preference Reversals

The scenarios A, B and C appear in the first three choice situations where no overruling treatment is present and again after several decisions of the subject were overruled. Scenario B is directly attached to the first row of consecutive overruled decisions to increase a potential treatment effect. By confronting the subject with the same scenario twice, it is possible to clearly identify preference reversals that are independent of economic consequences because the payoffs do not change between the two choices. As described above, the general decision-making behavior might be influenced by the overruling treatment and lead to a different decision in an identical scenario. Therefore, it follows naturally from hypothesis H3a that subjects should exhibit the below behavior:

H3b: *Preference reversals are observed in the scenarios A, B and C.*

The purpose of hypothesis H3b is to provide clearer evidence for an impact of the overruling treatment on the decision behavior since it would be inconsistent to choose differently if the scenario is exactly the same. Summarizing both hypotheses, it is expected that subjects vary their behavior after the overruling treatment independently of economic concerns. However, the direction of the change is not obvious since subjects might have different reasons to change their preference and the literature does not provide clear predictions.

After stating the hypotheses with regards to other-regarding behavior, the control premium and choice preferences, the next chapter will describe the results from the experiment which potentially validates the assumed behavior.

5 Results

Before presenting the experimental results in accordance with the stated hypotheses in the previous chapter, a few general facts and characteristics about the participating subjects will be displayed.

5.1 Subject pool characteristics

Since the number of participants is small with a total of 40 subjects, it can't be expected to be representative of the general population or varied enough to draw conclusions based on specific demographic characteristics. Nevertheless, the major attributes of the subject pool will be outlined to give an idea of its composition.

The participants were recruited through the VCEE and consisted mostly of students from a variety of fields with 20 % studying economic subjects, 15 % medicine, 13 % natural sciences and 52 % being students from other areas. The average age was 24.6 years with 52.5 % female and 47.5 % male subjects. The participants were relatively experienced since only 12.5 % of the subjects never took part in an experiment before while 35 % did so already more than 5 times.

With regards to the task performance in part 1, almost all subjects could correctly identify all 77 images with one participant failing to do so for one image. Therefore, every subject earned the maximum amount of 20 points in the real-effort task and was able to make further payoff relevant decisions in part 2. Since the total earnings of a subject did not only depend on their decisions but also heavily on the outcomes of the lotteries and the role specification of being a “citizen” or “non-citizen”, there were significant differences between the participants. The earnings ranged from EUR 7.25 to EUR 26.75 with an average amount of EUR 16.05 and a standard deviation of EUR 4.88.

5.2 Other-regarding behavior

As described in the previous chapter, there are two main indicators of other-regarding behavior from the subjects. The certainty equivalents could reveal a preference to increase the other player's payoff and the acceptance rate could reflect a deviation from purely selfish behavior.

Certainty Equivalents

In part 3, subjects were facing the choice between a lottery and a safe payment in a neutral context. The lottery included a payment to another player while the safe payment did not. By eliciting the CE, the subjects could therefore include their valuation for the other player's payoff which should lead to an increase of the CE if other-regarding behavior is present. Hence, as derived for hypothesis H1, an average CE above the EV of the lottery would be an indicator of such behavior towards the "non-citizen".

Indeed, for all three lotteries, the average CEs have a higher value than the EV of the standalone lottery that only includes the payments to the deciding player. In Table 5.1, the three lotteries, their respective EVs and the average elicited CEs from the subjects are presented. In the last two lines, the difference between the two measures and its statistical significance are displayed. The two-tailed t-test confirms a significant deviation of the average CE from the EV for the first two lotteries with a confidence level of 90 % and 99.9 % while the last difference is not high enough to be considered statistically significant.

	Lottery 1	Lottery 2	Lottery 3
Lottery	(15, 0.8; -5, 0.2)	(10, 0.3; -25, 0.7)	(20, 0.5; -15, 0.5)
EV	11	-14.5	2.5
CE	12.4	-11.3	3.9
CE - EV	1.4	3.2	1.4
Significance	$\alpha = 0.1$	$\alpha = 0.001$	not significant

Table 5.1: Certainty Equivalents

As argued for hypothesis H1, it is a conservative assumption that $CE(l) \geq EV(l)$ should hold since most people are risk averse (Harrison and Rutström (2008)) and the potential impact of social preferences could be mitigated in the presence of risk (Brock et al. (2013), Cappelen et al. (2013), Owens (2016)). Although it might be that risk-seeking preferences are present, the results are also in line with other-regarding behavior since all three CEs surpass this benchmark with two of them being statistically significant. The fact that lottery 2 with a relatively high negative EV exhibits the most significant difference could be explained by more risk-seeking behavior in the loss domain as predicted by prospect theory (Kahneman and Tversky (2013)). Subjects are willing to take more risks to escape a certain loss and therefore show an even higher CE than for the other lotteries.

By comparing the standard deviations for the individual CEs, it is remarkable that it is the highest for lottery 3 ($s_3 = 7.2$) while it is similar for the other two lotteries ($s_1 = 4.6$ and $s_2 = 4.9$). In contrast to them, lottery 3 with its equally likely outcomes and an EV closer to 0 makes

it difficult to heuristically label it as desirable or not. It might be that subjects had difficulties in choosing an appropriate safe payment that they would prefer over the lottery which could be the reason for the statistical insignificance.

With regards to the short-coming of the experimental setup to properly differentiate between risk and social preferences, it should be mentioned that the majority of subjects would need to be labeled as risk loving if other-regarding behavior is not present. Two thirds of the individual CEs are higher than the respective EVs and 30% of the subjects chose a higher CE for all three lotteries. This would be in dramatic contrast to observed behavior in other experiments where the share of risk seeking individuals is significantly lower (Harrison and Rutström (2008), Hartog et al. (2002)). Therefore, it is more likely that subjects indeed account for the other player's payment.

Overall, the elicited CEs are in line with well-known behavior by reflecting common risk and other-regarding preferences. The results give a hint that subjects take the payment to the other individual into account while determining their CEs for these specific lotteries. One could expect that this result would directly translate into an elevated acceptance rate of "non-citizens" since the CEs indicate a high valuation for the option to allow the other player to join the state. The acceptance rate is the second indicator of social preferences and will be looked at more closely now.

Acceptance rate

As outlined in hypothesis H1, a risk neutral individual that would strictly maximize its expected earnings would only choose the "Allow"-option five times in the sequence of fifteen decisions. Therefore, the average acceptance rate over all decisions should be higher than the benchmark of 33 % such that the presence of other-regarding behavior is possible.

Overall, the willingness to accept the other player depends heavily on the scenario and results in rates ranging from 12.5 % (Scenario D) to 90 % (Scenario E' and H). Naturally, the better the conditions of the lottery compared to the safe option, the higher is the acceptance rate. This fact will be further explored at a later point when the general decision-making behavior will be analyzed. All choice situations and the associated average acceptance rates are shown in Table 5.2.

Scenario	Lottery (Allow)	EV	Safe Option (Not Allow)	Avg. Acceptance Rate
A	(10, 0.7; -2, 0.3)	6.4	7	75%
B	(6, 0.4; -4, 0.6)	0	1	50%
C	(10, 0.6; -10, 0.4)	2	3	50%
D	(10, 0.3; -25, 0.7)	-14.5	0	13%
E	(15, 0.8; -5, 0.2)	11	12	83%
F	(20, 0.5; -15, 0.5)	2.5	3	58%
B	(6, 0.4; -4, 0.6)	0	1	53%
A	(10, 0.7; -2, 0.3)	6.4	7	83%
C	(10, 0.6; -10, 0.4)	2	3	65%
D'	(10, 0.3; -25, 0.7)	-14.5	-12	63%
E'	(15, 0.8; -5, 0.2)	11	8	83%
E''	(15, 0.8; -5, 0.2)	11	4	90%
F'	(20, 0.5; -15, 0.5)	2.5	-3	83%
G	(10, 0.0; 0, 1.0)	0	0	85%
H	(11, 0.95; -2, 0.05)	10.35	8	90%

Table 5.2: Acceptance Rates

One can clearly see that many scenarios exhibit particularly high rates which results in an overall average acceptance rate of 68 % across all decisions. A two-tailed t-test shows that this is a significant difference to the threshold of 33 % a selfish EV maximizer would choose ($\alpha = 0.001$) and to a rate of 50 % a randomized choice would induce ($\alpha = 0.01$) which supports hypothesis H1. Therefore, it is reasonable to conclude that subjects actively choose to accept the other player even if the expected earnings for the subject are lower than the safe payment.

Furthermore, there are two scenarios (G and H) with special properties that can be used to identify clearer and more direct evidence for social preferences. Scenario G has no payoff consequences for the deciding player by choosing either the lottery or the safe option. This constellation should lead subjects with other-regarding preferences to the obvious conclusion to accept the other individual. Scenario H only offers a slightly better outcome than the safe payment by accepting the lottery but has a very high probability of success. Even subjects with particularly high risk aversion could be convinced to accept the lottery while individuals that do not care about the “non-citizen” might still want to choose the safe option in scenario H. With an acceptance rate of 85 % in scenario G and 90 % in scenario H, the results clearly show that most subjects pay attention to the lottery properties and are willing to enable payments for the other player even if there is only a slight or no profit involved.

Overall, combining the evidence from the elicited CEs and the average acceptance rates, the results indicate that subjects take the well-being of the “non-citizen” into account or reveal risk-seeking preferences in the abstract context of an immigration choice. Therefore, the following result is established:

Result 1: *With an average acceptance rate of 68% and the condition $CE(l) \geq EV(l)$ being true for all three average CEs, subjects' behavior is consistent with risk-seeking and other-regarding behavior.*

Hence, subjects are willing to increase the payment of the other individual even if they are exposed to more risk while doing so. This is in line with results from similar experiments that look at general sharing behavior under risk (Brock et al. (2013), Cappelen et al. (2013), Owens (2016)). As outlined before, due to the low share of risk-seeking people and the magnitude of elicited CEs, it is likely that social preferences indeed are a driver behind “Allow”-decisions.

Having discussed the social preferences of the subjects, the scope will now move to the main objective of the thesis by identifying the intrinsic value of control in the choice situations.

5.3 Control Preferences

Subjects can protest overruled decisions by inserting a protest amount they would pay to switch back to their preferred option. High protest amounts that exceed the economic profit of the desired option could reveal a control premium that subjects are willing to pay. Before analyzing the drivers behind the protest amounts more thoroughly, evidence on the existence of a control premium will be presented.

Control Premium

As described in hypotheses H2a and H2b, a control premium CP is observed when the protest amount P is higher than the rational price $P' = |CE(l) - s|$ with $CE(l)$ being the certainty equivalent which is elicited in part 3 and with s being the safe payment.

In fact, in all 7 overruled decisions the subjects chose an average protest amount $P > P'$ revealing control premiums that differ in magnitudes between scenarios. In Figure 5.1, all overruled scenarios with their respective average protest amounts and control premiums are depicted. The sum of the blue and orange bar is the full protest amount P with the orange bar representing the control premium CP while the blue bar shows the economically rational part P' . The associated error margins for a confidence level of 95 % are displayed for the CP values.

As the bars already visually indicate, scenarios D', E, E' and F exhibit particularly high control premiums that are statistically significant ($\alpha = 0.001$) while it is only slightly appearing in the protest amounts in choice situations D, E'' and F'. For an overview of the scenarios with the respective average values for P and CP , see Table 5.3. As many other studies that found a

significant average control premium (Bartling et al. (2014), Bobadilla-Suarez et al. (2017), Fehr et al. (2013), Ferreira et al. (2020), Lübbecke and Schnedler (2020), Owens et al. (2014), Pikulina and Tergiman (2020) and others), the results also show a statistically significant average premium of 2.79 points (EUR 0.7) which represents 40 % of the average protest amount of 7.06 points (EUR 1.77). Subjects are therefore willing to forego 13% of their total earnings to gain back control over their choice. This value is similar to results from Lübbecke and Schnedler (2020) with 15%, Owens et al. (2014) with premiums ranging from 8 % to 15 % and Pikulina and Tergiman (2020) with a share of 10%.

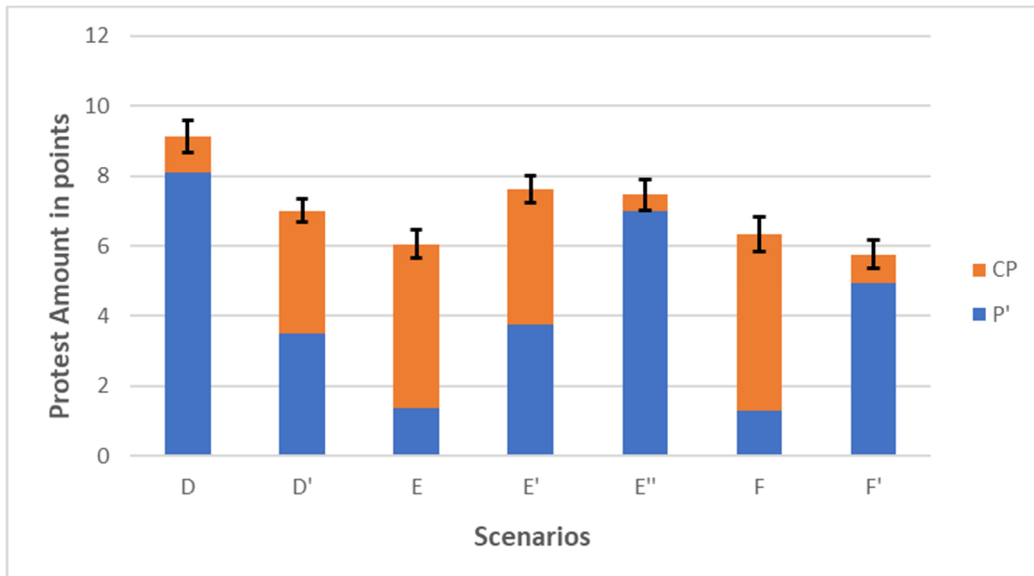


Figure 5.1: Control Premiums and Protest Amounts

Scenario	Lottery (Allow)	EV	Safe Option (Allow)	Avg. P	Avg. CP
D	(10, 0.3; -25, 0.7)	-14.5	0	9.15	1.03
D'	(10, 0.3; -25, 0.7)	-14.5	-12	7.02	3.53*
E	(15, 0.8; -5, 0.2)	11	12	6.06	4.72*
E'	(15, 0.8; -5, 0.2)	11	8	7.63	3.88*
E''	(15, 0.8; -5, 0.2)	11	4	5.77	0.82
F	(20, 0.5; -15, 0.5)	2.5	3	6.35	5.07*
F'	(20, 0.5; -15, 0.5)	2.5	-3	7.47	0.47
(*: significant at 99.9% level)				Overall average:	7.06
					2.79*

Table 5.3: Overruled Scenarios with respective P and CP

As suggested by Bartling et al. (2014), one could try to measure the consistency of the subjects' preferences for control by computing Cronbach's Alpha¹ (Cronbach (1951)) which is a coefficient that shows how consistently a group of items can measure the aspect of interest. In

¹ $\alpha = \frac{N}{N-1} \left(1 - \frac{\sum_{i=1}^N s^2(x_i)}{s^2(y)} \right)$, with $-1 < \alpha < 1$, N being the number of items (7 overruled decisions), x_i being the measured control premiums of all individuals for decision i and y being the sum of all premiums measured for one individual across all 7 games.

the context of this experiment, if a subject i would reveal higher preferences for control compared to another individual j in one decision, it should exhibit larger values for CP across all decisions relative to subject j (Bartling et al. (2014)). For this sample, Cronbach's Alpha is equal to 0.75 which shows that the measured individuals' control premiums are strongly positively correlated and therefore indicate consistent control preferences across decisions. The result is comparable to the values measured by Bartling et al. (2014) and Ferreira et al. (2020) which range between 0.47 and 0.77.

Interestingly, there seems to be no significant difference in the strength of control preferences between initial choices to allow or to not allow the “non-citizen” to join the state. The average control premium for “Allow”-choices was 2.45 points (EUR 0.61) while it was 3.47 points (EUR 0.87) in decisions where “Not Allow” was originally chosen. A two-sample t-test² cannot reject the null hypothesis ($\alpha = 0.05$) that the two values stem from the same population. This might indicate control premiums that are independent of the preference direction which was also stated by Hausfeld et al. (2020). Subjects are therefore equally interested in gaining back control over decisions where they accepted or rejected the other player. Furthermore, the results also suggest that the participants show similar behavior in the loss and gain domain which can be observed based on scenario D' which exposes the subject to high expected losses for both options. The control premium in D' is similar to the values from other scenarios that are located in the gain domain which a two-sample t-test³ verifies ($\alpha = 0.05$). This is in line with the analysis of Bobadilla-Suarez et al. (2017) who specifically designed an experiment to explore the potential difference of control preferences between losses and gains which confirmed the equal behavior in both domains.

In addition to comparable results from other studies, the novel design of the experiment also sheds some light on new aspects of the control premium. Firstly, the participants of other experiments are usually well informed about the possibility of being overruled, also because the majority of studies analyzed delegation games (Bartling et al. (2014), Bobadilla-Suarez et al. (2017), Fehr et al. (2013), Ferreira et al. (2020), Hausfeld et al. (2020)) while this design captures an element of surprise and uncertainty since it is not clear which scenarios can be overruled. Secondly, this experiment tries to frame the subject's decision in an immigration

² Both a paired and unpaired sample t-test showed no significant result. For the paired t-test, the two average control premiums for “Allow” and “Not Allow” decisions were calculated for every individual to produce comparable paired values (except for one individual that had to be excluded because it only chose “Not Allow”).

³ A similar procedure to calculate a paired test was used as described before for the comparison between “Allow” and “Not Allow” decisions. The value for CP from scenario D' was paired with the average CP from all other scenarios for every individual to enable a comparison.

context where she has to decide about the fate of her own state which has not been done before in experiments with regards to control preferences.

Hence, an average control premium does not only exist if potential interference is expected but also if there is uncertainty about which decisions might be overruled. Subjects' control preferences seem to be stable and consistent across decisions even if the majority of them are not subject to the overruling treatment. Furthermore, in the abstract sense of an immigration context, the subjects seem to care not only about the economic outcome of their choices but also if they were responsible in choosing them or if they were forced to take them. Therefore, it follows the result:

Result 2: *Subjects are on average willing to pay a significant control premium to gain back control which represents 13% of their total earnings and can be observed in either preference direction (“Allow” vs. “Not Allow”).*

Nevertheless, the results raise the question why certain scenarios exhibit high values for *CP* while others do not. For any of the three lotteries, there exists a scenario where a significant *CP* is measured while a variation of the scenario, where only the safe option is varied, exhibits no significant premium. For example, by comparing scenario D (low *CP*) to D' (high *CP*) in Table 5.3, one can see that the lottery is identical, while the safe payment is different. The same is true for the other lotteries: The scenarios E, E' and F show significant premiums, while their variations with a different safe option E'' and F' exhibit low values for *CP*.

Although significant results in all choice situations would further strengthen the evidence on the existence of a control premium, heterogeneity in the valuation for control is also found in other experiments. Pikulina and Tergiman (2020) try to classify the different desire for power among subjects into classes that consider standard, social and power preferences and observe that they differ substantially in their behavior. Therefore, the relatively small subject pool ($n = 40$) of this thesis' experiment might not represent the full range of preferences which could distort the results. Despite that, insignificant results for control premiums in single decisions were also found by other studies that had larger sample sizes. Ferreira et al. (2020) observe no significant intrinsic value of decision rights in two out of ten delegation games with subjects from Japan ($n = 132$) and Bartling et al. (2014) also identified one delegation game with a remarkably lower statistical significance compared to the other nine decisions ($n = 172$). Hence, it seems that even with a more representative sample, certain scenario characteristics might

enhance or mitigate the magnitude of CP which results in the heterogenous distribution between decisions.

One of these characteristics might be the span between the EV and the safe option that is particularly small in the significant decisions compared to the insignificant ones, as can be seen in Table 5.3. As observed by Owens et al. (2014), the control premium has a higher impact on behavior when the expected earnings of the available options only differ slightly. This is also true for all scenarios in this experiment since control premiums are decreasing when the difference between the EV and the safe option $|EV - s|$ increases. Especially scenario E and its variations show well how the control premium declines while the span $|EV - s|$ increases⁴. Hence, it seems that the control premiums are particularly small in the scenarios where the lottery and the safe payment differ significantly in their monetary payoffs.

Despite that, there is another fact that distinguishes the insignificant from the significant decisions which is the amount of negative values for CP . Due to the design of the experiment, subjects exhibit a negative control premium when their chosen P is below $P' = |CE(l) - s|$. While some error margin for CP should be considered since subjects might not be able to exactly compute the appropriate protest amount and rely on heuristic calculations, the share of negative premiums is still higher for the insignificant scenarios. By looking at the individual control premiums across scenarios, the heterogenous distribution of control preferences with negative and positive premiums can be observed more clearly. Figure 5.2 shows histograms of the distribution of individual control premiums for the different scenarios with blue bars being negative premiums (control aversion) while orange bars are positive ones. The last histogram shows the total distribution across all scenarios.

Firstly, two regularities can be observed: (i) It seems that the majority of subjects' control premiums are located near zero in the interval $[-5, 5]$ and (ii) one can clearly see that negative control premiums are elicited in all scenarios with shares from 33 % in scenario E to 63 % in scenario E''. Visually, the histograms already indicate that the significant scenarios (D', E, E' and F) show a lower share of negative values for CP than the other scenarios. On average, 36 % of subjects exhibit a negative valuation for control in the significant scenarios with an average CP of -3.32 while a higher share of 55 % do so in the insignificant scenarios D, E'' and F' with an even lower average premium of -5.68. Additionally, the share of people that exhibit negative premiums of -10 or lower is stable at approximately 10 % for the insignificant scenarios while

⁴ Comparing scenario E to E'', the span $|EV - s|$ increases from 1 to 7 while the value for CP decreases from 4.72 to 0.82

the significant ones hardly have such highly negative values. Although the overall share of control averse subjects is higher than in other similar experiments, negative values for control premiums are observed regularly for 10 to 20 % of subjects in other studies (Bartling et al. (2014), Ferreira et al. (2020), Owens et al. (2014)). Therefore, there appears to be a stronger concentration of control aversion in the scenarios that revealed insignificant average control premiums.

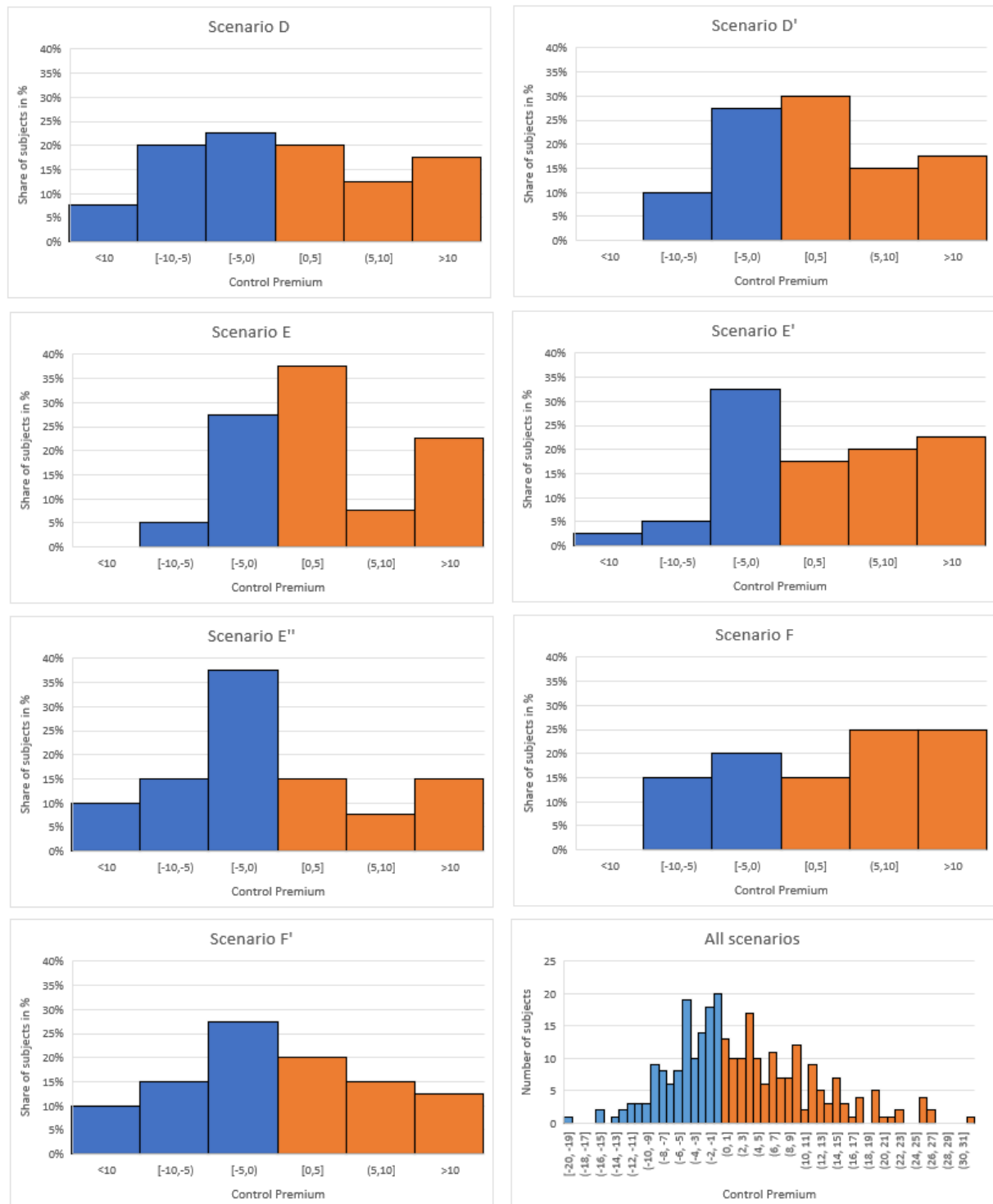


Figure 5.2: Individual Control Premiums

On the other hand, we see a significant share of subjects with $CP > 10$ across all scenarios ranging from 13 % to 25 %. Interestingly, only 35 % of subjects exhibited such a high control premium at least twice throughout the whole experiment. This group of participants (High-Group) is responsible for 78 % of all elicited premiums larger than 10 which reveals a drastic difference in control preferences between the High-Group and the rest of participants (Low-Group). While the average control premium in the High-Group is 8.05 points, it is -0.05 for the Low-Group which shows that the overall average premium is heavily skewed by individuals in the High-Group.⁵ This is further illustrated by Table 5.4 which shows the average values for the High- and Low-Group across the different scenarios.

Scenario	Low-Group Avg.	High-Group Avg.	Total Avg.
D	-1.84	6.35	1.03
D'	1.26	7.76	3.53
E	1.43	10.82	4.72
E'	0.68	9.81	3.88
E''	-1.84	4.76	0.47
F	1.09	12.45	5.07
F'	-1.12	4.43	0.82
Total	-0.05	8.05	2.79

Table 5.4: Control Premiums in High- vs. Low-Group

One can clearly see that subjects of the High-Group are constantly willing to pay high control premiums despite the different economic conditions between the scenarios. It seems that only a minority in the population expresses particularly strong control preferences while most of the subjects only have moderate or even control averse preferences. Hence, an overall measured desire for control is not due to small or medium-sized premiums from a large share of subjects but rather from high premiums produced by only a third of the sample.

Summarizing the described results, an average significant control premium can be identified although there is a substantial share of subjects that reveal control averse preferences, especially in the single decisions that did not deliver significant values for CP . It might be that other behavioral drivers play a role when the difference in expected earnings of the two options increases such that a control premium would still be present but is not easily distinguishable from other effects. On the other hand, a minority of subjects that reveals particularly strong control preferences heavily impacts the average control premium in the single scenarios and overall. Furthermore, the results are mostly in line with recent experimental studies, showing

⁵ The separation between High- and Low-Group is arbitrary but was chosen to illustrate that only a minority of the sample (High-Group) is mainly responsible for the measured significant overall control premium.

that a control premium exists also in the novel experimental design which suggests the context of an immigration choice.

After thoroughly analyzing control preferences, the following paragraphs will have a deeper look into the motives behind the individual protest amounts which could further explain the described characteristics of the control premium.

Protest Amounts

As outlined before, there is significant heterogeneity in control premiums that subjects exhibit in the different scenarios. This variety is also visible in the individual protest amounts that are depicted in Figure 5.3. For every scenario, the histograms show the share of subjects that chose a value for P in the respective bins while the last figure shows the total distribution of P across all decisions. Although there are significant differences between the scenarios, there are some regularities that can be observed.

The majority of protest amounts are located in the range between 0 and 10 which represents a total share of 78 %. Since this density in the lower range of values can be observed in all scenarios, subjects in this group might think that an appropriate value for P should stem from this interval although it is rational to protest more in some scenarios. The focus on a narrow range of values could indicate that participants follow a heuristic approach to find a direct and simple solution to the problem. In the last histogram in Figure 5.3, the total distribution reveals three spikes for the popular values 0, 5 and 10 which are also the modal values across the scenarios. Despite having the option to insert any value between 0 and 25 with one decimal, subjects focus on amounts that are multiples of 5 because they might be more salient. The described heuristics could contribute to the observed low or even negative control premiums if subjects tend to avoid a thorough analysis and comparison of the two options and rather choose lower amounts that feel suitable.

In contrast to the frequent choice of lower values, there is a stable share of 12.5 % to 20 % for values of $P \geq 15$ across all choice situations which are exclusively chosen by a group of 40 % of participants. Interestingly, this group coincides almost completely with the High-Group (subjects that exhibited $CP > 10$ at least twice) which was introduced earlier with 86 % of its members choosing $P \geq 15$ at least once. This translates into significant differences between the groups: The overall average protest amount in the High-Group is 11.6 points while it is only 4.6 for subjects in the Low-Group. While members of the Low-Group choose to not protest at all in 26 % of their decisions, subjects in the High-Group do so in only 8 % of their choices.

The main reason for the low average protest amount in the Low-Group is the major share of 66 % of choices that have a value of $P \leq 5$ while it is only 29 % for the High-Group.

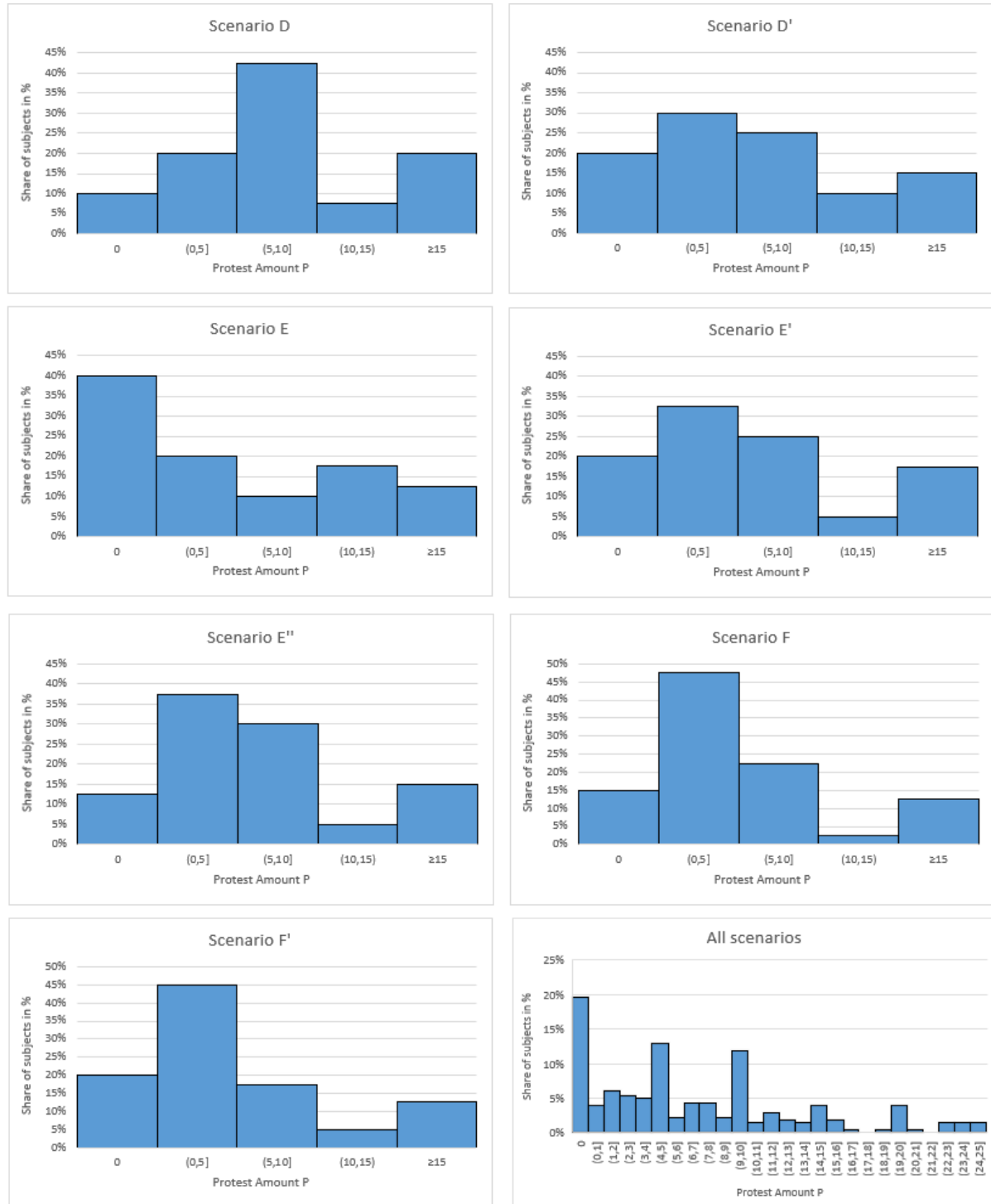


Figure 5.3: Individual Protest Amounts

The described disparity in protesting behavior between the two groups could deliver an important explanation for the phenomenon that we observe high control premiums exactly in the choice situations that offer similar payoffs for the lottery and safe option (Table 5.3) which was also observed by Owens et al. (2014). It seems that subjects in the Low-Group are too hesitant in increasing P even if the utility difference between the options increases while members of the High-Group steadily choose high protest amounts even if the scenario offers two similar options. The stable share of high protest amounts of the High-Group tip the scale in the scenarios with a small utility difference which produces particular high control premiums. For scenarios that have insignificant premiums, it seems that members of the Low-Group only increase the protest amounts slightly even if the difference between the options is significantly high.

Despite the individual differences between the two groups, the whole sample seems to react if the pecuniary difference between the options increases. The histograms in Figure 5.3 show how the mass of protest amounts moves more to the higher bins when the difference between lottery and safe option increases while the share of zero protest decreases.⁶ Only for scenario F and F' one can see no significant change in the protests which might be due to the low stakes of the safe option (Table 5.3). It could be that subjects cannot distinguish very well between the two scenarios since they do not perceive the gain or loss from the safe option as impactful. Linear regression results which are shown in Table 5.5 also confirm the expected behavior that subjects should protest more heavily the better the preferred option is compared to the imposed one.

Variables	Coefficient	Std. Error	t-value	p-value
Intercept	6.59	1.17	5.64	0.0000***
ev_out_diff	0.18	0.09	2.02	0.0443*
sequence_pos	-0.05	0.11	-0.43	0.6689
(*: $p \leq 0.05$, **: $p \leq 0.01$, ***: $p \leq 0.001$)				

Table 5.5: Regression Results on Protest Amount P

The underlying data for the regression contains 280 single decisions ($40 \text{ participants} \times 7 \text{ overruled decisions} = 280$) with the protest amount P being the independent variable while the explanatory variables are (i) *ev_out_diff* which is the absolute difference $|EV(l) - s|$ between the expected value of the lottery and the safe payment and (ii) *sequence_pos* which is the position of the scenario in the sequence. The positive and significant coefficient of *ev_out_diff* shows that subjects tend to increase their protest amount when they observe a higher

⁶ The utility difference increases from scenario D' to D, from scenario E to E'' and from scenario F to F'. See Figure 5.3 for details.

disparity between their desired option and the other one. Therefore, one part of hypothesis H2a is supported and well-behaved preferences that reflect selfish economic interests are observed. This result is not trivial since it indicates that subjects understood the tasks and know how to act in their best interests. It might support the assumption that subjects are well aware of their valuation for control by putting in high protest amounts which is also described by Bobadilla-Suarez et al. (2017).

The variable *sequence_pos* serves the purpose to test the other part of hypothesis H2a which is concerned with the potential impact of past overruled decisions on the magnitude of P . The position in the sequence directly reflects how many times the subject was exposed to the overruling treatment. It could be that subjects increase their efforts to protest when frequent overruling stirs up the desire to avoid future interference. The coefficient is not significant and therefore does not suggest that such an effect exists.⁷ There seems to be no correlation which is also reflected in a correlation coefficient of $r = -0.05$ that does not show any significant relation. Hence, subjects seem to care only about current interference and do not exhibit any spillover effects from past decisions while choosing a protest amount.

Since the described facts only partially confirm the hypothesis, the following result is formulized:

Result 3: *Selfish interests generally increase the protest amount while subjects with strong control preferences (High-Group) contribute disproportionately more to the control premium.*

By looking at the presented results, one can clearly see that the control premium on an aggregate level is composed of a vastly heterogenous distribution of protest amounts that might also be based on heuristic approaches. The analysis of the amounts on an individual level confirmed that a minority with strong control preferences might be the deciding factor behind measured significant premiums. In general, well-behaved preferences with regards to selfish interests can be observed.

While this chapter showed that a desire for control exists in either preference direction (“Allow” vs. “Not Allow”), we did not pay much attention to the initial decision itself. Moving away from the focus on the overruled decisions and the control premium, the following chapter’s scope will be on general decision behavior and the drivers behind allowing or not allowing the “non-citizen” to join the state.

⁷ Other variables like the number of overruled decisions in the whole history or in the last three rounds also do not show significant results.

5.4 Immigration choices

Since subjects are confronted with a variety of different scenarios, it is interesting to have a deeper look at factors that potentially impact the final decision to accept the other player into the state. The protesting behavior already suggests that selfish economic interests play a role in the initial decision while the overruling treatment could have an effect as well. Firstly, the drivers behind an “Allow”-decision are analyzed which is followed by evidence about potential preference reversals.

Choice Preferences

As outlined before, subjects are willing to accept “non-citizens” significantly more often than a random choice or pure self-interest would predict. Hence, social preferences seem to impact the decision which was described in one of the previous chapters. It follows naturally to ask if selfish interests play a role as well which could be analyzed in the same way as for the protest amounts by looking at the pecuniary difference $|EV(l) - s|$. The probability to accept the other player should increase with the monetary value of the lottery compared to the safe payment. In Table 5.6, the results of a probit model that predicts the binary outcome (“Allow” vs. “Not Allow”) based on the already presented explanatory variables *ev_out_diff* and *sequence_pos* are presented.⁸ The data contains all 600 single decisions ($40 \text{ participants} \times 15 \text{ decisions} = 600$).

Variables	Coefficient	Std. Error	z-value	p-value
Intercept	0.32	0.12	2.56	0.0105*
ev_out_diff	0.11	0.02	6.86	0.000***
sequence_pos	0.03	0.01	2.04	0.0417*
(*: $p \leq 0.05$, **: $p \leq 0.01$, ***: $p \leq 0.001$)				

Table 5.6: Probit Model Results on “Allow”-Probability

The coefficient for *ev_out_diff* is highly significant and shows that the probability of allowing the other player to join is impacted by the economic attractiveness of the lottery compared to the safe option. Subjects are much more likely to accept another participant if their selfish interests are met than when both options have similar pecuniary consequences. This is also supported by the acceptance rates that were presented in Table 5.2 that show significant differences between the scenarios. While scenario D with the most unattractive lottery shows a

⁸ Since the probability to allow the “non-citizen” to join is the scope of the probit model, the variable *ev_out_diff* represents the difference $EV - s$ to enable a correct interpretation of the coefficient. It therefore does not show the absolute difference $|EV - s|$ as presented before in the linear regression for the protest amount.

low acceptance rate of 12.5 %, other scenarios with obviously desirable lotteries (e.g. E' and E'') are accepted by up to 90 % of the subjects. Lotteries that are not clearly better than the safe payment (e.g. B, C and F) show more divided behavior with rates of 50 % to 65 %.

Interestingly, the variable *sequence_pos* seems to have a significant positive impact on the “Allow”-probability.⁹ Hence, the later the subject is exposed to a certain scenario, the more likely it will choose “Allow”. It might be that the continuous interference with the subject’s choices leads to an increased tendency to let other players join their state since more overruled decisions in the past are reflected in a higher value of *sequence_pos*. Although the significant coefficient might be a first indicator of such an effect, the visualized data shows a rather mixed picture.

In Figure 5.4, the upper plot shows the acceptance rate (y-axis) for the scenarios depending on what position in the sequence (x-axis) the subject faced the scenario. For example, it could be that the acceptance rate is lower for a scenario when it was placed at position 8 compared to position 15. If the timing of the scenario has no impact on the decision behavior, we would expect a horizontal line for every scenario since the acceptance rate would stay the same independently of the position. As one can clearly see, the upper plot shows a different picture with varying rates across the positions which do not form a clear pattern. The lower plot in Figure 5.4 shows the linear trend lines based on the previously described rates to visualize a general direction and confirms the impression. Some trends are positive while others are negative which does not indicate that there is a clear positive impact on the acceptance rate when moving forward in the sequence. Only 8 out of 15 scenarios are depicted in the graphs since the remaining ones are not randomized in their position.¹⁰

It is important to note that the comparison of acceptance rates for scenarios across different positions in the sequence relies heavily on a sufficient randomization such that the frequencies for every position are equal.¹¹ This is definitely a problem in the presented sample of only 40 subjects which does not enable an adequate comparison or further statistical tests. Hence, the depicted graphs must be analyzed with a grain of salt.

⁹ Other variables like the number of overruled decisions in the whole history or in the last three rounds do not show significant results.

¹⁰ The first three scenarios (A,B,C) are only randomized on the first three spots and the following four scenarios (D,E,F,B) have fixed positions and are therefore excluded.

¹¹ For example, if one sequence position only has a few observations, the acceptance rates are heavily influenced by only a little number of subjects and are most likely not comparable to other positions.

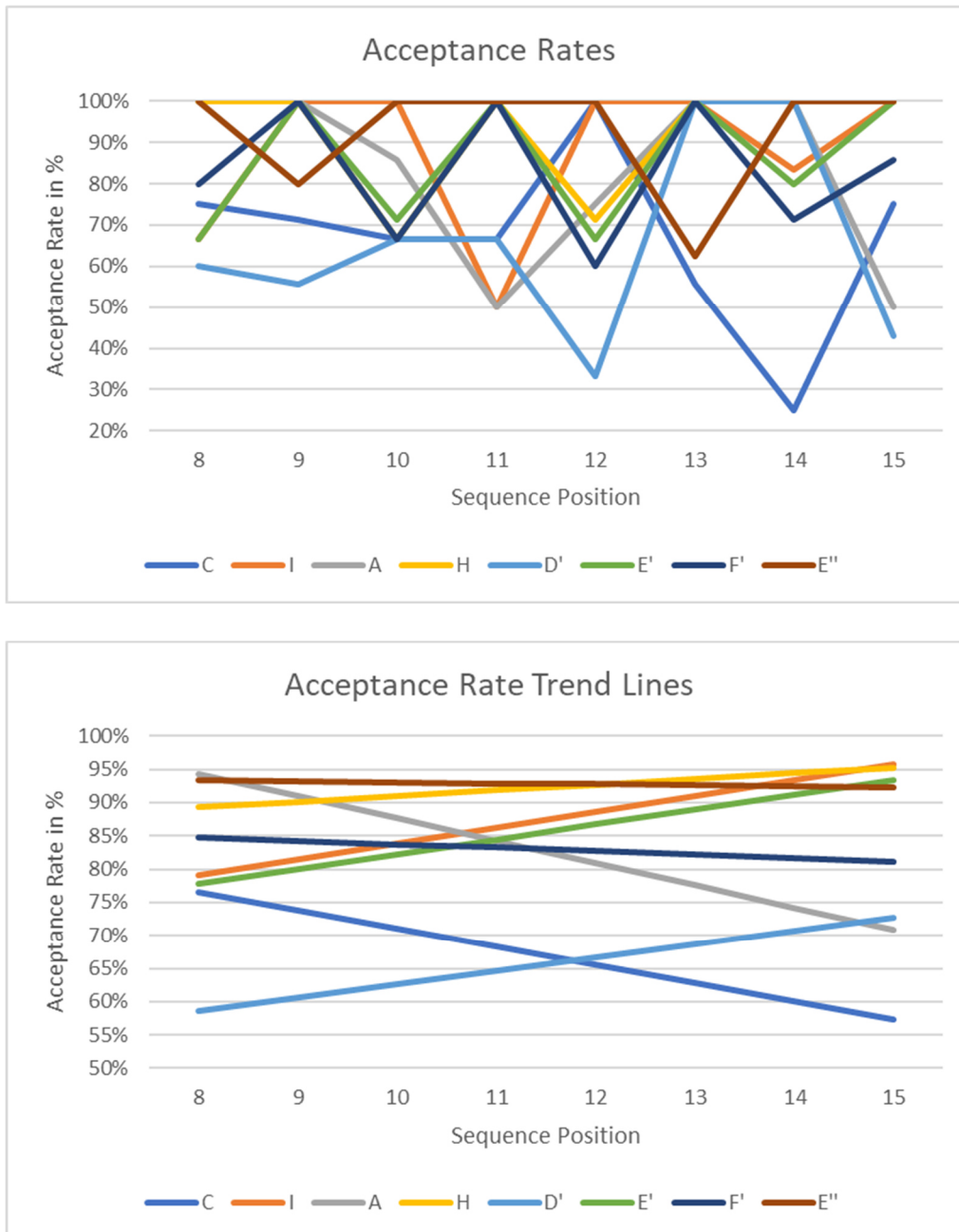


Figure 5.4: Acceptance Rates depending on sequence position

Nevertheless, the mixed results do not give a clear answer to the question if repeated overruling has an impact on the decision behavior. They give a first hint that there might be some effect whose direction is unclear as described in the hypothesis. The only significant evidence is the probit model that suggests a greater likelihood to accept the “non-citizen” when the scenario has a later position. Since the acceptance rate is generally high across all scenarios, choices to accept the “non-citizen” got overruled more frequently than decisions to reject. This might have led subjects to choose “Allow” even more in an act of defiance to make up for unpleasant overruling in the past. It is difficult to identify if this particular reason is the main driver or if other explanations might be more plausible.

For example, the trend of increasing relative attractiveness of lotteries in the second half of the experiment is an important factor. Since the first seven decisions are mostly fixed and only the eight remaining scenarios are randomized, there is a significant difference in the payoff structure between both halves. While in the first half of the sequence the average difference ($EV - s$) amounts to -2.8 points, the lotteries are relatively more attractive in the second half with an average of 1.72 points. Therefore, it seems intuitive to expect an increased acceptance rate in the later stage of the experiment because the lotteries simply get better over time. This could explain the observed positive trend in the “Allow”-probability. On the other hand, the relative value of the lottery should already be captured by the coefficient *ev_out_diff* in the probit model which still results in a significant coefficient for *sequence_pos*.

Overall, a slight trend towards more acceptance of “non-citizens” in the later stage of the experiment is visible but it is not clear if this is caused by the changing payoff structure in the decisions or by the intended treatment effect. Since there are three scenarios that are presented twice, one could compare them directly and draw further conclusions by observing potential preference reversals in the subjects’ behavior.

Preference Reversals

Scenarios A, B and C come up twice in the sequence by presenting them once in a randomized order at the start of the sequence and a second time after the first three overruled decisions. Scenario A and C are randomized between positions 8 to 15 while scenario B is fixed at position 7 (Table 3.1) and they are not subject to the overruling treatment. Naturally, one would expect that subjects stick to their initial choices since the conditions of the scenarios are identical. While the first decisions for these three scenarios were made without the knowledge of possible interference, the second choice happens after the overruling treatment. In Table 5.7, the share of reversed decisions is shown in the first line while the second and third line show in which direction the subjects chose to switch their initial decision. The last two lines depict the acceptance rate in the first and second occurrence of the scenarios.

Scenario	A	B	C
Reversed decisions	23%	43%	30%
Switch to "Not Allow"	8%	20%	8%
Switch to "Allow"	15%	23%	23%
Acceptance Rate (First)	75%	50%	50%
Acceptance Rate (Second)	83%	53%	65%

Table 5.7: Preference Reversals

A substantial fraction of subjects seems to change their mind and decide differently after witnessing interference in their other choices. Remarkably, scenario B which was directly attached to the first sequence of overruled decisions exhibits the highest share of reversed decisions with 42.5 % compared to 22.5 % and 30 % for scenarios A and C. While more subjects tend to switch from choosing “Not Allow” in the first appearance to “Allow” in the second one, it is an almost even split for scenario B. The net effect of switched decisions leads to an overall increase of the acceptance rate in all scenarios. Scenario C shows a relatively high increase with 15 % while B shows only a slight increase due to the even mix of reversed decisions.

Interestingly, the same shift towards more acceptance can be observed as previously described in the general analysis of all scenarios. Since the presented scenarios are identical, the discussed increasing relative value of the lotteries in the second half of the sequence cannot be the reason behind the preference reversals. Hence, these results slightly strengthen the impression that subjects are more open to accept the “non-citizen” as the sequence progresses and more of the overruling treatment took place.

Of course, there might be other reasons for the preference reversals that should be addressed briefly. First of all, since the monetary value of the lotteries and safe payments for these scenarios do not differ much, it could be that subjects are indifferent and choose randomly which is supported by the divided decisions for B and C in the first round with both having 50 % acceptance rate. On the other hand, A has a clear tendency towards acceptance which is further increased in the second round which does not indicate a random choice. Additionally, if subjects indeed do not care about which option gets implemented, they should follow the same approach in the second round. This might be the case for B but probably not for C since the acceptance rate climbs from 50 % to 65 % which is a significant shift. It might be that subjects were indifferent between the two options in scenario B but not for A and C where we see the more significant shifts towards more “Allow”-choices. Hence, it cannot be reasonably concluded that a randomized choice pattern led to the increase in acceptance in the second round.

Despite that, it might be that subjects were not yet familiar with the choice situations and did not decide based on correct reasoning in the first cycle. Learning effects could lead to a more rational decision in the second cycle which could reflect their true preferences. The experimental setup does not allow to test this specifically but the instructions were carefully formulated such that misunderstandings should be avoided. Especially in the first three

decisions where the overruling treatment was not present yet, it is not likely that subjects were very confused about the task. Furthermore, the acceptance rates in the three scenarios do not show shifts in an extraordinary manner that would suggest a vastly different opinion about the identical scenario in the second cycle.

Summarizing the results, preference reversals can be observed and are directed towards a slightly higher acceptance of “non-citizens” after the overruling sequence. This gives some support for the previously observed tendency to choose “Allow” more often when the sequence progresses which is potentially caused by the treatment effect of interference. Therefore, the final result is established:

Result 4: *Subjects tend to allow “non-citizens” to join more frequently after repeated interference of the overruling treatment and when the economic benefits of such a choice increase.*

After describing the experimental results with regards to other-regarding behavior, the control premium and the general decision behavior of subjects, a thorough discussion should follow which addresses new and previously described problems with the conclusions and potential need for further research in this area.

6 Discussion

The main goal of this thesis is to provide laboratory evidence on the impact of control loss on immigration attitudes and choices. A control premium – which is a well-established concept and is described by many other studies – is also observed in the abstract context of immigration that this experiment tries to suggest. Of course, an experiment cannot replicate big-scale immigration dynamics and its political and societal implications. Nevertheless, if one would dare to draw conclusions from this thesis for the real world, perceived control over immigration flows seems to play a role in the decision-making process. Although there is significant heterogeneity in control preferences, people seem to care not only about economic consequences of immigration but also about the control instruments to steer it.

Interestingly, the results suggest that the direction of immigration attitudes does not play a role for control preferences. Hence, natives that favor strict immigration rules could see potential migrants as acceptable only if they bring a net profit to the economy that exceeds the control premium. On the other hand, people with positive attitudes towards immigration might be willing to accept migrants who are economically weak just to overcome restrictive immigration policies.

Besides control preferences, the decisions were designed to capture the main components of an immigration choice which are (i) selfish concerns about economic impacts (ii) uncertainty about the negative or positive direction of such impacts and (iii) social considerations about immigrants. Subjects showed expected selfish behavior and exhibited most likely a compound of social and risk preferences that were drivers behind their choices. Generally, the majority of subjects were willing to frequently accept the other participants if the economic consequences were not too severe. Therefore, at least in this sample, most subjects had positive attitudes towards welcoming the other players into their state.

Despite that, potential spillover effects from the constant interference were observed that were pointing in the direction of increasing acceptance rates. Although this result has to be considered with a grain of salt, it could give a first hint that people might change their behavior when they are exposed to authorities that overrule their preferred outcomes. For example, governments that enact new immigration policies might be confronted with increasing protests from people with opposite views that goes way beyond economic concerns. Control signals from the government could transport the feeling that immigration flows are overseen and managed properly which potentially lowers such protests (Wright (2016)).

Besides the confirmation of some results of the existing literature surrounding control preferences, the experiment also sheds some light on new aspects of the control premium. A minority in the sample is responsible for high average premiums while a significant share of subjects is willing to accept the overruling. Surprisingly, there is stronger resistance to the overruling treatment than Lübbecke and Schnedler (2020) observe in their study where only 20 % of subjects resist interference and are willing to pay a control premium. Hence, it seems that arbitrary overruling from the computer in our experiment produces similar or even stronger protesting behavior than human interference that is targeted to correct the subject's choice like in the design of Lübbecke and Schnedler (2020). We would have expected that resistance would be higher in situations where another participant wants to patronize the deciding subject. It could be that the framing of the overruling in our design contributed to the observed behavior which seems stronger than in other studies.

Another interesting result is the slight tendency to allow more “non-citizens” into the state after repeated overruling. Spill-over effects from interference were also described and observed by Fehr et al. (2013) which stems from the willingness to avoid future disutility from being overruled and translated into lower delegation rates. In our experiment, it is not as clear why people increase their acceptance rates throughout the choice sequence and would probably need an extension to the design to investigate it properly.

Although interesting insights can be derived from the novel experimental design, there are some short-comings that need to be discussed in the following paragraphs.

Design aspects

First of all, the main feature of the experiment – the computation of the control premium by comparing the economically rational protest amount P' with the actual amount P – needs to be discussed since it relies heavily on the correct interpretation of the CE. Although other studies used similar methods (Bartling et al. (2014), Ferreira et al. (2020)), it is not clear if subjects indeed value options based on the elicited CEs. If subjects do so, one would expect that throughout the whole experiment, they always choose a lottery l over a safe payment s if $CE(l) > s$ holds. By looking at the data of all overruled decisions, 74 % of the individual decisions were consistent with this strict rule.

It does not seem that only a small number of subjects is responsible for the majority of inconsistent decisions. Only two individuals stand out since they were acting consistent only in two out of seven of the overruled decisions. It rather looks like that two scenarios caused

particularly low shares of consistent decisions: Scenario E with 60 % and scenario F with 55 %. As outlined before, these two scenarios belong to the ones with a low monetary difference between the lottery and safe payment. It could be that the options seem equal to the subjects when they made decisions in part 2 while there was a more nuanced valuation in the longer elicitation process later in part 3. When we would allow an error margin of 30 %¹², the share of consistent decisions increases to 84 % which is significantly better but still leaves room for distortion of the results.

Hence, there seems to be some inconsistent behavior especially in two decisions where a significant control premium was measured. Since the quantification of the premium is based on the correct valuation of the lottery, it might pose a problem in the measurement method. On the other hand, the other two scenarios with significant control premiums D' and E' show a large share of consistent decisions with 83 % and 78 % (98 % and 88 % with error margin applied). Therefore, it does not seem that a measurement error is the driver behind the significant control premiums although one would need to analyze this problem more thoroughly with a larger sample size.

In general, the measurement of control premiums could be enhanced by enabling a higher variation of protest amounts between the scenarios. Although there is significant variance in the individual amounts, the average values for P are similar in the different scenarios ranging between 5.77 points in scenario E'' and 9.15 in scenario D. If a control premium could be observed in scenarios that provoke very strong protesting reactions and in choice situations that exhibit low protest amounts, it would provide stronger evidence for consistent control preferences. Additionally, a stronger variation in scenario conditions could be used to identify situations where control plays a significant role or has no effect at all.

Another design factor which could be improved to gather more clear insights is the identification of spillover effects. The mostly randomized sequence of choices allowed the investigation of positional effects only on an individual level but lacks a sound analysis on an aggregate level. It cannot be answered distinctly if the observed tendency to allow more “non-citizens” towards the end of the sequence is simply due to the different payoff structure in the second half. Furthermore, it would be interesting to extend the design such that more preference reversals can be observed since this experiment only looked at three scenarios. If the sequence

¹² For example, if $CE = 5$ and $s = 4$, a subject would act inconsistently if it would choose the safe payment since it valued the lottery higher. With an error margin of 30 %, the subject does not act inconsistently since the safe payment is now assumed to be perceived by the subject in the interval (2.8, 5.2) and the upper bound would be higher than the CE .

would contain more identical choice situations, more convincing observations on behavioral changes from control loss could be gathered.

Finally, the small sample size poses a problem for some statistical analysis and tests. For example, the exploration of the reasoning behind increasing acceptance rates after repeated overruling was limited due to lacking randomization across the different positions in the sequence. A larger sample size would provide stronger statistical reliance to draw conclusions from and allows a more detailed analysis. It would also provide a better representation of demographic characteristics that are especially interesting in the context of immigration. One could get insights from attributes like immigration background, beliefs about immigrants and other factors.

After identifying potential room for improvements in the design, additional aspects of interests will be outlined that could be investigated by future research.

Further research

The novel design of the experiment provided first evidence of existing control preferences in the domain of immigration attitudes. Nevertheless, the results raise some interesting follow up questions that should be further investigated.

First of all, the experiment is not designed to distinguish the framing effect of the overall immigration context from the overruling treatment impact. Due to the within-subject design, the influence of control loss is only observed in the immigration context and has no comparable reference point without framing. It is true that the CEs are elicited in a neutral context in part 3 but only for decisions that are also subject to the overruling treatment. Hence, it is not possible to disentangle the effects of framing and control loss that might lead to a control premium being paid. Ideally, one would design a control group that faces the same choice sequence in a neutral context which is similar to the one in part 3. Then it would be possible to observe if the treatment effect of control loss is amplified, reduced or remains stable between a neutral and an immigration context.

Additionally, one need to account for the framing of the control loss in this experiment which is distinct from the immigration setting. The first overruling treatment comes as a surprise and the alert messages in case of interference are formulated in an alarming tone that is emphasized by red and bold text. As outlined earlier, the strong framing as well as its arbitrary nature might have an impact on the subject's behavior. It would be useful to think about other ways to let subjects experience control loss to explore the differences between overruling framings. Since

this experiment tries to frame control loss as something unwanted and forceful, one could think of contexts that are perceived as neutral or more acceptable.

A more neutral framing could describe an overruled decision as random and not dependent on the subject's decision. Potential feelings of defiance, independence from others or the fear of overruling due to mistakes would only play a minor role. A different framing might be seen as beneficial if it inherits some "correction of behavior" that is described as a desired attitude. Suggestions that the deciding subject was not acting "politically correct" or being "morally wrong" by rejecting the immigrant could reduce the desire to gain back control after interference. On the other hand, it might still provoke resistance due to the patronizing nature of the overruling.

Hence, to explore different framing effects and to disentangle them from the isolated impact of control loss, one could design 6 treatments which differ in the framing of choices and overruling, as illustrated in Table 6.1.

Treatments	Framing of Choices	Framing of Overruling
A	Neutral	Forceful
B	Neutral	Neutral
C	Neutral	Beneficial
D	Immigration	Forceful
E	Immigration	Neutral
F	Immigration	Beneficial

Table 6.1: Treatments with Framing Variations

While this experiment corresponds to treatment D, a control group like treatment A would already be a great addition to see potential differences stemming from the framing as immigration choices. Treatments D, E and F could reveal important factors that determine the severeness of people's resistance to immigration policies. Decision makers might want to design policies and political communication in a way that (i) it seems strict and authoritarian (forceful treatment) or (ii) strong messages and opinions are avoided (neutral) or (iii) that a certain policy is framed as the "right thing to do" (beneficial). Finally, by conducting treatments with a neutral framing of choices, one could explore if control preferences are amplified in an immigration context.

Beyond the framing of control loss, it would be interesting to design different settings to measure control preferences by varying status quo distributions of control instruments. In this experiment, the subjects hold the decision rights until they are taken away by the computer. However, one could change the status quo and let subjects start without control over their

decisions and elicit a willingness to pay for control like it is done in other studies (Neri and Rommeswinkel (2017), Pikulina and Tergiman (2020)). One could compare utility gains and losses from different redistributions of control and draw useful conclusions from it. Although Bobadilla-Suarez et al. (2017) observed that control preferences stay stable while experiencing economic gains or losses, it could be different for the valuation of control itself. Subjects might experience a loss of control as more severe than an equally worthy gain of control. Hence, reference points can play a role and determine the severeness of people's reaction to control redistributions and the needed economic compensation to make them indifferent about it.

Such a compensation could be the measured control premium which indicates how much people value decision rights in monetary terms. However, it would be interesting to see if other aspects besides pecuniary interests could be convincing to "calm people down" if they experience a loss of control. Following the insights from other studies, people's desire for control might be reduced if they do not perceive a control loss as an interference from others (Fehr et al. (2013), Neri and Rommeswinkel (2017)) or as ignoring their authorship (Lübbecke and Schnedler (2020)). In the context of real-world immigration, Wright (2016) describes that governments need to send control signals which communicate that attention is paid to immigration and proper controls are in place. Perhaps one could investigate if control signals also work on an individual level: One subject that holds the decision rights signals competence to another participant who in turn is willing to forego control over its own outcomes.

The thought of willingly giving control to other parties is especially interesting with regards to this experiment's results. The high share of negative control premiums that were measured in the different scenarios raise the following question: Are there situations in which people prefer being controlled? Is there something like "benevolent overruling"? For example, Hausfeld et al. (2020) suggest that there might be control aversion towards unfair outcomes. Hence, people prefer to avoid responsibility over decisions that might have negative consequences for others. With regards to immigration, it could be that economically weak immigrants are unwanted but still accepted due to social concerns. If the decision is reversed by an authority, it might be that people happily accept the overruling and do not protest it since they are not in control and not responsible for the impact on the immigrant. It would be interesting to create scenarios that try to replicate such a situation and further investigate it.

Finally, the general study of control preferences and control premiums lacks some observations from the field that should be a focus of future research. Experimental studies with large samples and different variations of games have been conducted and consistently find significant control

premiums. Therefore, it would be a great enrichment of current research to replicate such findings in non-laboratory environments. More external validity could be achieved and more conclusions with regards to real-world implications could be drawn. Nevertheless, it is probably difficult to find natural conditions where control preferences in economic decisions can be properly observed. Especially in the context of immigration choices, it might be problematic to elicit control desires on an individual level since these decisions are typically not decided by single individuals but in political processes that are far more complex.

7 Conclusion

This thesis investigated the relevance of control preferences in a choice sequence that is framed as an immigration setting. Subjects could decide between the risky option to let “non-citizens” join their state or to reject them and receive a safe payment. A feeling of control loss was induced by overruling the subject’s decisions and implementing the unwanted option. Subjects were able to resist the interference and pay a price to reverse the overruling.

The main finding is the existence of a control premium that subjects are willing to pay to gain back control over their choice. After controlling for social and risk preferences, the observed protest amounts were significantly higher than the economic value of the desired options. However, individual control preferences differ notably and lead to a heterogenous distribution of control premiums. While a remarkable share of subjects exhibits control averse behavior, a minority of one third consistently chooses high protest amounts signaling a strong desire for control. Hence, the average control premium is heavily skewed by a small group of control-loving subjects outweighing control averse preferences.

Furthermore, the results indicate that subjects generally care about the well-being of the “non-citizen” who profits from entering the state. The acceptance rates are generally high and increase further if the economic interests of the deciding and the passive player are aligned. Interestingly, it seems that subjects tend to choose the “Allow”-option more often the further the experiment progresses. It might be that repeated interference leads to an increased acceptance of “non-citizens” which is supported by observed preference reversals after the overruling treatment.

The framing of the experiment suggests that the feeling of lost control can influence immigration attitudes. This study contributes to the existing literature with regards to control preferences and mostly confirms its results. Control premiums are equally strong if subjects decide to allow or not allow another participant to join their state. Hence, control preferences might not only play a role for anti-immigration campaigns like it was expressed by the Brexit campaign but also for people who would like to welcome more immigrants. Economic interests, risk and social preferences obviously play a role in the forming of immigration views. However, decision makers must take into account that people like to feel responsible for their own fate and despise the feeling of losing control. Immigration policies need to signal that proper controls are in place even if they are designed to attract more immigrants.

Further research should focus on the exploration of different settings of control redistributions and gathering observations in the field. The discussed treatments could help to disentangle framing effects from control preferences and to better understand the mechanisms behind strong or weak desires for control.

References

- Alesina, A., and Tabellini, M. (2022). The Political Effects of Immigration: Culture or Economics? (No. w30079). *National Bureau of Economic Research*.
- Bartling, B., Fehr, E., and Herz, H. (2014). The intrinsic value of decision rights. *Econometrica*, 82(6), 2005-2039.
- Bobadilla-Suarez, S., Sunstein, C. R., and Sharot, T. (2017). The intrinsic value of choice: The propensity to under-delegate in the face of potential gains and losses. *Journal of Risk and Uncertainty*, 54(3), 187-202.
- Bouacida, E., and Foucart, R. (2022). Rituals of Reason: A Choice-Based Approach to the Acceptability of Lotteries in Allocation Problems (Working Paper).
- Brock, J. M., Lange, A., and Ozbay, E. Y. (2013). Dictating the risk: Experimental evidence on giving in risky environments. *American Economic Review*, 103(1), 415-37.
- Cappelen, A. W., Konow, J., Sørensen, E. Ø., and Tungodden, B. (2013). Just luck: An experimental study of risk-taking and fairness. *American Economic Review*, 103(4), 1398-1413.
- Charness, G., and Gneezy, U. (2010). Portfolio choice and risk attitudes: An experiment. *Economic inquiry*, 48(1), 133-146.
- Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16(3), 297-334.
- d'Ancona, C., and Ángeles, M. (2016). Immigration as a threat: Explaining the changing pattern of xenophobia in Spain. *Journal of International Migration and Integration*, 17(2), 569-591.
- Dustmann, C., and Preston, I. (2006). Is immigration good or bad for the economy? Analysis of attitudinal responses. In *The Economics of Immigration and Social Diversity*. Emerald Group Publishing Limited.
- Dustmann, C., and Preston, I. P. (2007). Racial and economic factors in attitudes to immigration. *The BE Journal of Economic Analysis and Policy*, 7(1).
- Dustmann, C., and Frattini, T. (2014). The fiscal effects of immigration to the UK. *The Economic Journal*, 124(580), F593-F643.
- Eckel, C. C., and Grossman, P. J. (2008). Men, women and risk aversion: Experimental evidence. *Handbook of Experimental Economics Results*, 1, 1061-1073.
- Fehr, E., Herz, H., and Wilkening, T. (2013). The lure of authority: Motivation and incentive effects of power. *American Economic Review*, 103(4), 1325-59.
- Ferreira, J. V., Hanaki, N., and Tarroux, B. (2020). On the roots of the intrinsic value of decision rights: Experimental evidence. *Games and Economic Behavior*, 119, 110-122.
- Florio, E. (2022). Contact vs. Information: What shapes attitudes towards immigration? Evidence from an experiment in schools. *Journal of Behavioral and Experimental Economics*, 96, 101790.

- Gawn, G., and Innes, R. (2019). Who delegates? Evidence from dictator games. *Economics Letters*, 181, 186-189.
- Haaland, I., and Roth, C. (2020). Labor market concerns and support for immigration. *Journal of Public Economics*, 191, 104256.
- Hainmueller, J., and Hopkins, D. J. (2014). Public attitudes toward immigration. *Annual Review of Political Science*, 17.
- Harrison, G.W. and Elisabet Rutström, E. (2008), "Risk Aversion in the Laboratory", Cox, J.C. and Harrison, G.W. (Ed.) *Risk Aversion in Experiments* (Research in Experimental Economics, Vol. 12), Emerald Group Publishing Limited, Bingley, pp. 41-196.
- Hartog, Joop, Ada Ferrer-i-Carbonell, and Nicole Jonker. "Linking measured risk aversion to individual characteristics." *Kyklos* 55.1 (2002): 3-26.
- Hausfeld, J., Fischbacher, U., and Knoch, D. (2020). The value of decision-making power in social decisions. *Journal of Economic Behavior and Organization*, 177, 898-912.
- Hjerm, M., and Nagayoshi, K. (2011). The composition of the minority population as a threat: Can real economic and cultural threats explain xenophobia?. *International Sociology*, 26(6), 815-843.
- Jasso, G., Massey, D. S., Rosenzweig, M. R., and Smith, J. P. (2000). The New Immigrant Survey Pilot (NIS-P): Overview and new findings about US legal immigrants at admission. *Demography*, 37(1), 127-138.
- Kahneman, D., and Tversky, A. (2013). Prospect theory: An analysis of decision under risk. In *Handbook of the Fundamentals of Financial Decision Making: Part I* (pp. 99-127).
- Khadjavi, M., and Tjaden, J. D. (2018). Setting the bar-an experimental investigation of immigration requirements. *Journal of Public Economics*, 165, 160-169.
- Koch, C., Dezső, L., and Tyran, J.-R. (2022). Populist Propaganda and Anti-Migration Sentiments. Mimeo
- Langer, E. J. (1975). The illusion of control. *Journal of Personality and Social Psychology*, 32(2), 311.
- Lübbecke, S., and Schnedler, W. (2020). Don't patronize me! An experiment on preferences for authorship. *Journal of Economics and Management Strategy*, 29(2), 420-438.
- Neri, Claudia, Rommeswinkel, Hendrik, 2017. Decision Rights: Freedom, Power, and Interference (Working paper).
- Owens, D. (2016). Lotteries in Dictator Games: An Experimental Study. *Eastern Economic Journal*, 42(3), 399-414.
- Owens, D., Grossman, Z., and Fackler, R. (2014). The control premium: A preference for payoff autonomy. *American Economic Journal: Microeconomics*, 6(4), 138-61.
- Pikulina, E. S., and Tergiman, C. (2020). Preferences for power. *Journal of Public Economics*, 185, 104173.
- Sloof, R., and von Siemens, F. A. (2017). Illusion of control and the pursuit of authority. *Experimental Economics*, 20(3), 556-573.

Storesletten, K. (2003). Fiscal implications of immigration—A net present value calculation.
The Scandinavian Journal of Economics, 105(3), 487-506.

Wright, C. F. (2016). Control signals and the social policy dimensions of immigration reform.
In *Handbook on Migration and Social Policy* (pp. 161-179). Edward Elgar Publishing.

Appendix

A1 Instructions

The following pages show the instructions that were identical for every subject. The only differences arise in the outcome which depends on the randomized role of a “citizen” or a “non-citizen” and in the overruling screen which depends on the decision in the overruled scenario. The variations on the overruled screen are shown with a blue text. An example for the outcomes are shown from the perspective of the deciding player. Personal data which was meant as contact information was removed from the welcome page.

A1.1 First Page (Welcome)

Welcome to our Experiment!

Now, you are participating in an economic experiment where you may earn a significant amount of money. However, your actual earnings will depend on your choices, the choices of others, luck and your ability to perform in a task (which we will describe later on). In order to maximize your earnings, you must clearly understand your options, which requires reading the instructions very carefully.

Confidentiality:

During the experiment, you will make decisions **anonymously**. This means that we will not link your name or any personal information to your choices. Your responses and choices in the experiment will be recorded, stored in computers and will not be given to any third party. Your participation is voluntary and you can withdraw your participation at any point in time before the experiment is over. In this case, however, we will not be able to pay you.

Rules of conduct:

During our experiment, **you are not allowed to use electronic devices (other than the computer on which the experiment is running) or to communicate with other subjects**. If you have a question or want to say something, please raise your hand. Please do not ask questions publicly. We will come to you and answer your question in private.

Earnings:

The currency in which all experimental earnings are charged are called *points*. During the experiment, all earnings and costs will be expressed in points. While you are

expected to mostly gain points, in some parts of the experiment you may also lose some. At the end of the experiment, your earnings and losses will be summed up and converted into Euros at the following exchange rate:

4 points = 1€

At the end of the experiment, you will be paid in Euros which you will receive in cash. Notably, even if you are unlucky and make some losses, the minimum amount you will receive is 7.50 EUR (but in all likelihood you earn considerably more).

Please click on Next to start the experiment.

A1.2 Second Page (Introduction)

What is the experiment about and how will it proceed?

This experiment is about **building an effective state** and **the consequences of risky and safe choices for the citizens of that state**.

In our experiment, you are a **CITIZEN** which means that you **can invest in your state** and **that you have the decision power whether other people will join and benefit from your state or not**. Later, we will elaborate on what these choices are and explain it in more detail.

The experiment consists of **three parts**. At the beginning of each part, you will get detailed instructions.

Once you have read these instructions and do not have any questions, you can start Part 1 of the experiment. Please click on Next if you are ready to proceed.

A1.3 Real Effort Task (Introduction)

Part 1: Working for your state

As a **CITIZEN**, you will have the **opportunity to work for your state** in Part 1 and make it more efficient for Part 2. In addition, you will have the opportunity to **earn some points for yourself**.

More precisely, you face **77 image identifying tasks**. For each task, you will be presented with **3 images** and **one label**. Your task will be to click on the image that best fits the label.

A short example: Assume that the label says „cat“ and you see three images showing a „desk“, a „cat“ and a „spoon“. To complete the task correctly, you have to click on the image of the „cat“.

Overall, you have **300 seconds** to complete all 77 tasks. Your goal should be to correctly answer as many tasks as you can. Should you make a mistake, you will receive an error message and will be prompted to try again. Should you give an incorrect answer at your second try, your answer will be marked as incorrect.

Implications and earnings from Round 1:

Your **first goal** is to correctly solve at least **50** image identifying tasks. Only if you manage to do that you 'make your state more effective' in Part 2. More precisely: only then your state will generate an *income stream of 10 points per round* and you will be able to make some additional earnings in Part 2. (Notably, there will also be a little twist but we will tell you more about it later on).

Your **second goal**: Any number of correctly solved images above this threshold will give you some personal **Part 1** earnings as summarized in the table below:

Correct tasks	Part 1 Earnings	Investment in your state
75-77	20 points	YES
70-74	15 points	YES
65-69	10 points	YES
60-64	5 points	YES
50-59	0 points	YES
0-49	0 points	NO

Please raise your hand if you have questions. If you have no question, click on Next.

HINT: To identify images faster: Select the correct image with the mouse and click "enter" on the keyboard!

Notably, the task (and your 300 seconds) will not start immediately. There will still be another announcement by the experimenter after everyone has clicked 'Next'.

A1.4 Real Effort Task (Feedback)

Part 1: Feedback

You have correctly identified **77 images**. Consequently, your **Part 1** earnings are **20 points** (out of a maximum of 20 points).

You have completed Part 1 and will start Part 2 on the next page. In Part 2, you will face a sequence of choices where you have to **decide if another individual is allowed to join your state or not**. On the next page, we will give you detailed instructions.

A1.5 Immigration Decisions (Introduction)

Part 2

You are about to start Part 2. In this part, due to your performance in Part 1, your state will generate an **earnings stream** for 15 periods.

Crucially, as a **CITIZEN** of your state, you will also have to **decide whether and under which conditions** another participant – the **NON CITIZEN** - **would be able to join you** and benefit from your state.

The **NON CITIZEN** is **another participant in this session** that would **always benefit from joining your state**. If you don't allow him/her to join, s/he will receive nothing and only keep the earnings from Part 1.

Who is the NON CITIZEN?

As indicated before, there is a little twist: Indeed all participants have read the same instructions so far but half of today's subjects will be assigned to **NON CITIZENS** in the end of the experiment. Put differently, this means that everyone completes the whole experiment as a **CITIZEN** and take decisions that have payoff consequences for both groups. However, **at the end of the experiment** – with a 50% chance – you will be assigned to one of the groups and receive their payoffs.

If the payoffs for the **CITIZEN** are chosen for you, your own decisions determine your payoffs. If the payoffs for the **NON CITIZENS** are chosen for you, the decisions of other participants determine your payoffs.

Overall, it is rational to assume that you are the CITIZEN and decide that way!

Please click Next to see detailed instructions for the decisions you will take.

A1.6 Immigration Decisions (Example)

Part 2: Decisions and Example

How you will decide if other subjects can join your state or not?

You will be given **15 choices** one after another. These choices revolve around whether you want to **"allow"** the **NON CITIZEN** to join your state or whether you do **"not allow"** him/her to join. Each choice has different payoff consequences for you and the other subject, depending on your decision.

Note, however, that the effective state **you created due to your success in Part 1** always generates an *earning stream* for you, which means that every period you will **receive a payoff of 10 points, independent of your decision**.

Allow: If you allow someone to join, **you will accept a lottery** with outcomes that happen with certain probabilities, which are clearly displayed. (This lottery is in addition to your income stream of 10.) The **NON CITIZEN** will receive a payoff of 10.

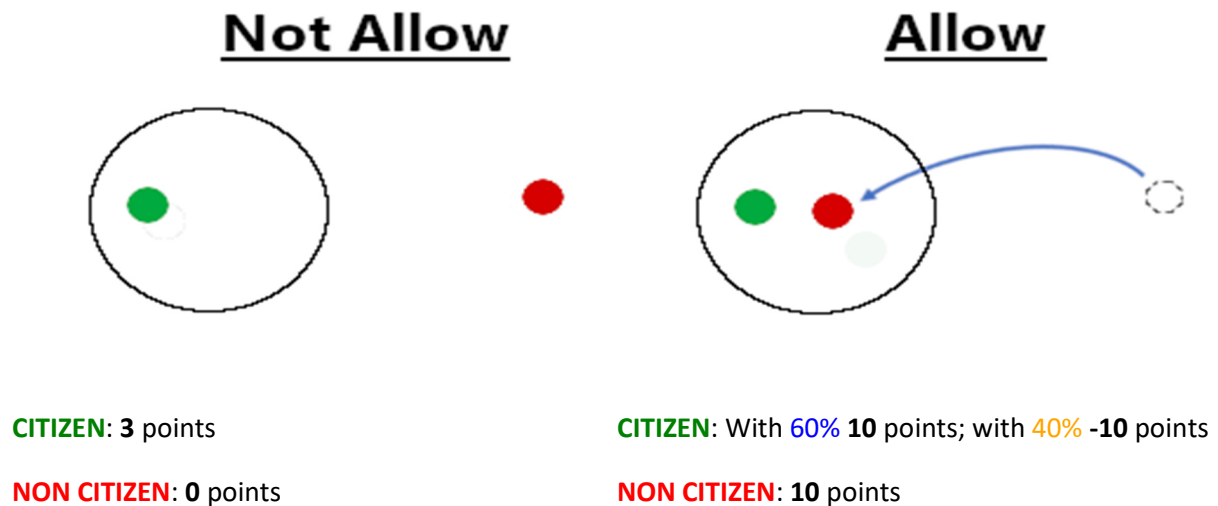
Not Allow: If you do not let someone join, you will receive a safe payoff (in addition to your income stream of 10). The **NON CITIZEN** will get a payoff of 0.

Please note, that only **3 randomly** chosen choices out of the 15 choices will eventually be implemented and determine the payoffs in Part 2. At the very end of the experiment, the computer will **randomly** select these three choices. Please take all 15 choices seriously, since all of them could potentially be chosen!

Example choice situation

Let us consider an example of the choices you will have to make. Remember that everyone earned some payoff in Part 1 (you earned 20). Any **gains** or **losses** you make in Part 2 will be **added** or **subtracted** to the Part 1 earnings, which will be your final payoffs (once Part 3 payoffs are also included):

Effective State Earnings for **CITIZEN: 10** points (independent of choice)



The figure illustrates the two possible choices: On the left side, you can see the situation if you do **not allow** the other subject to join. On the right side, the case is shown if you **allow** the other subject to join and the respective payoff consequences. Let's look at them in more detail:

Not allow:

If you choose "Not Allow", the **NON CITIZEN** is not allowed join:

Please note that the payoffs in this situation are always **certain**. The **CITIZEN** gets a safe payoff (varying for different situations) and the **NON CITIZEN** always gets 0, as s/he cannot benefit from your state. In this particular example the payoffs are:

CITIZEN: Safe payoff of **3 points**

NON CITIZEN: Safe payoff of **0 points**

Allow:

If you choose "Allow", the **NON CITIZEN** joins you. This has a positive effect on his/her earnings (10 points) but it makes your earnings as a **CITIZEN uncertain**. In this particular example, the situation is the following:

CITIZEN: **10 points** with a probability of **60%**

or

-10 points with a probability of **40%**

NON CITIZEN: Safe payoff of **10 points**

Put differently, you as a **CITIZEN**, would receive a payoff which is determined by a **lottery**. With a probability of 60%, you will get the "good" lottery outcome (i.e. 10 points) and with a probability of 40%, you will receive the "bad" lottery outcome (i.e. 10 points will be subtracted from your earnings). The **pie chart** illustrates the relevant probabilities. By choosing "Allow", you accept this lottery and if this decision is

selected in the end, the lottery will be played out. The other subject, as a **NON CITIZEN**, would receive a certain payoff of **10 points**.

Please note, that the details of the lottery differ in the 15 periods. Independent on whether the **CITIZEN** decides to "allow" the **NON CITIZEN** to join or not, s/he will always receive 10 points on top of either the lottery or the safe payoff.

To proceed, please click Next (you will still be able to read all previous instructions).

A1.7 Comprehension Questions

Comprehension Questions

Below you can find some comprehension questions. Please answer them by ticking the correct answer. Further below, you can find the instructions again

- Question 1)

In case you choose to allow the other participant to join. What are the implications?

- ☒ The NON CITIZEN will get a safe payoff. But as a CITIZEN your payment will be a lottery as it is uncertain how "good" or "bad" the impact of the joining is.
- ☐ As a citizen you will receive a safe payoff and the NON CITIZEN will get nothing.

- Question 2)

In case you choose Not to allow the other participant to join. What are the implications?

- ☐ The NON CITIZEN will get a safe payoff. But as a CITIZEN your payment will be a lottery as it is uncertain how "good" or "bad" the impact of the joining is.
- ☒ As a citizen you will receive a safe payoff and the NON CITIZEN will get nothing.

- Question 3)

Which decisions will be relevant for payment?

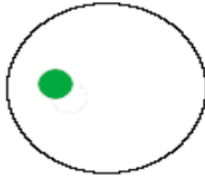
- ☐ All decisions will be paid out
- ☒ 3 out of 15 decisions are randomly chosen and paid out

A1.7 Choice Situation (Example from the sequence)

Choice Situation 4

Effective State Earnings for **CITIZEN**: 10 points (independent of choice)

Not Allow

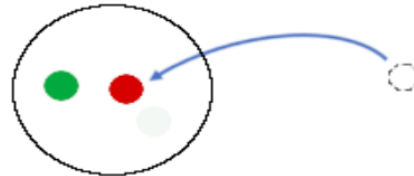


Payoffs

CITIZEN: 0 points

NON CITIZEN: 0 points

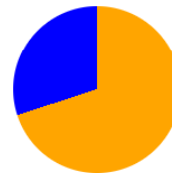
Allow



Payoffs

CITIZEN: With 30% 10 points; with 70% -25 points

NON CITIZEN: 10 points



The pie chart illustrates the probabilities

'Allow' or 'not allow' the other to join?

☐ Not allow

☐ Allow

Next

A1.8 Overruled Screen (Introduction)

ATTENTION Your decision was overruled!

In the last choice situation, you decided **NOT to allow (to allow)** the **NON CITIZEN** to join (i.e. you went for the safe payoff). However, your decision was overruled. While this overruling is arbitrary (i.e. it does not imply that your decision was wrong in any way), it means that now the **NON CITIZEN** is **allowed (not allowed)** to join your state (i.e. you will get the lottery), which is the opposite of what you decided initially.

From now on, all of your upcoming decisions can potentially be overruled (but **only some of them will**) and, if this is the case, we will inform you about it right after you made your decision. (If there is no such message, your preferred choice is valid). In case one of your decisions was overruled, you will always be able to **protest** against it.

How does the Protest work?

You can protest against the overruling and potentially **change the decision back to your preferred option** (i.e. NOT to allow **(to allow)** the other to join in Choice Situation 4). For that, you have to choose the **maximum amount of points** you would be willing to invest in your protest. More precisely, your investment has to be higher than a certain price which is randomly determined by the computer.

Successful Protest: If your chosen (maximum) amount is higher than the (random) price, you switch back to your original choice and get its benefits and **only have to pay the (lower) price** (instead of your stated amount).

Failed Protest: If your stated investment is lower than the (random) price, your decision remains overruled and you stay with the new decision but you do **NOT** have to pay the price.

Important This procedure implies that you should really state your **personal maximum** to stop the overruling as - in case your protest is successful - you only have to pay the potentially lower price (and not the full amount you state).

By clicking Next, you will get to the decision that just got overruled (Choice Situation 4).

A1.9 Overruled Screen (Example from the sequence)

ATTENTION Your decision was overruled!

In the last choice situation, you decided **NOT to allow (to allow)** the **NON CITIZEN** to join (i.e. you went for the safe payoff). However, your decision was overruled. This means, that now the **NON CITIZEN** is **allowed (not allowed)** to join your state (i.e. you will get the lottery), which is the **opposite of what you decided initially**.

As indicated before, you have to option to **protest** the overruling. What is the **maximum amount** of points you would be willing to invest to switch back to your preferred option "Not Allow" ("**Allow**")?

[You have to choose points between 0 and 25 and can use non-integer values such as 4.5 or 17.3]

*(If you type in a price of 0, so no protest, the allow/lottery-situation (**not-allow/safe payment-situation**) (and not your original choice) will be implemented for sure. A random price mechanism will be used to determine whether your protest is successful. This gives you an incentive to state your personal maximum as you only have to pay the potentially lower price or nothing at all in case your protest is unsuccessful.)*

Maximum amount you would like to invest to go back to Not Allow (Allow)?

A1.10 Part 3: Certainty Equivalent Elicitation (Introduction)

Part 3

You are about to start Part 3. In this part, you will have to do a task (described below) **three times**, where you can potentially earn some additional money. Notably, you are again matched with someone else in this experiment and your choices can affect the other person as well. This other person will also do the same choices. At the end, we will randomly decide whether your choices or the other person's choices will be implemented.

The Task: Each of the three tasks consists of 6 steps. In these steps, you repeatedly choose between Option A and Option B, which have different consequences for you and the other person.

Option A means that you get a lottery with a chance to get a **positive or negative payment** with certain probabilities. This lottery will always stay the same within the task. The other person gets a payoff of 10.

Option B: means that you get a sure payment. This sure payment varies across the steps in the task. The other person gets a payoff of 0.

Note that the person whose choices will finally be implemented will receive an additional 10 points in each of the three tasks (similar to Part 2). Therefore, your potential payoff for one task is the payment of the option you choose plus the 10 points.

Your goal is to always choose the option that you would prefer, since it may be paid out in the end.

A1.11 Part 3: Certainty Equivalent Elicitation (Example decision)

Your Decision - Task 1

Page 1 of 6

What do you prefer? (Independent of your choice, you will get an additional 10 points)		
Option A:	a 80.0% chance of winning 15 points and a 20.0% chance of winning -5 points The other participant will get 10 points for sure.	Option A
Option B:	an amount of 5.0 points as a sure payment The other player will get 0 points for sure.	Option B

A1.12 Outcome (Certainty Equivalent Elicitation)

Results for Part 3

In this part, YOUR decision counts

Your choice (for the randomly chosen step) in Task 1:

Option A: a 80.0% chance of winning 15 points and
a 20.0% chance of winning -5 points

Option A

Option B: an amount of 16.25 as a sure payment

Option B

As indicated above, you decided for "Option A" in this choice.

For "Option A", one of the two possible outcomes has been randomly realized based on the corresponding probabilities.

Your payoff in this task equals

-5.0

Your choice (for the randomly chosen step) in Task 2:

Option A: a 30.0% chance of winning 10 points and
a 70.0% chance of winning -25 points

Option A

Option B: an amount of 3.125 as a sure payment

Option B

As indicated above, you decided for "Option A" in this choice.

For "Option A", one of the two possible outcomes has been randomly realized based on the corresponding probabilities.

Your payoff in this task equals

-25.0

Your choice (for the randomly chosen step) in Task 3:

Option A: a 50.0% chance of winning 20 points and
a 50.0% chance of winning -15 points

Option A

Option B: an amount of 14.0625 as a sure payment

Option B

As indicated above, you decided for "Option A" in this choice.

For "Option A", one of the two possible outcomes has been randomly realized based on the corresponding probabilities.

Your payoff in this task equals

-15.0

Overall, your payoff from Part 3 (including the $3 * 10$ points you get independent of your choice) is: -15.0 points

Next

A1.12 Outcome (Immigration Decisions)

Payment for decisions in Part 2:

For the choices in Part 2, the computer **randomly selected you** as the one deciding

The following 3 scenarios were randomly selected: **1) Scenario M , 2) Scenario A , 3) Scenario E** (see table below for details)

For these three scenarios, the computer randomly draw the following lottery outcomes:

- 1) **good outcome!**
- 2) **bad outcome!**
- 3) **bad outcome!**

The following decision were made:

- 1) **NOT allow**. This decision was **overruled** though (your protest was unsuccessful)
- 2) **NOT allow**
- 3) **NOT allow** This decision was **overruled** though (your protest was unsuccessful)

This implies that you earn:

- 1) **15 points**
- 2) **7 points**
- 3) **-5 points**

Overall, your payoff from Part 2 (including your income stream of 10 points for the 3 periods selected) is: **0**

This is actually due to the fact that you did not identify enough correct images in Part 1 to invest in an effective state. Hence, you get no earning from Part 2.

For a more detailed summary of what happened during Part 2, please find below the decisions made by yourself and the other participant (The order in which you saw the scenarios may deviate from the order in this table):

Choice Situation	"Allow" Lottery Good	"Allow" Lottery Bad	Prob. Good %	"Allow" NON CITIZEN pay	"Not allow" safe pay	Your decision	Your Max	Other's decision	Other's Max	Drawn Price
A	10	-2	70	10	7	not allow		not allow		
B	6	-4	40	10	1	not allow		not allow		

C	10	-10	60	10	3	not allow		not allow		
D	10	-25	30	10	0	not allow	0	not allow	0	2.4
E	15	-5	80	10	12	not allow	0	not allow	0	6.4
F	20	-15	50	10	3	not allow	0	not allow	0	7.3
G	6	-4	40	10	1	not allow		not allow		
H	10	-10	60	10	3	not allow		not allow		
I	10	0	0	10	0	not allow		not allow		
J	10	-2	70	10	7	not allow		not allow		
K	11	-2	95	10	8	not allow		not allow		
L	10	-25	30	10	-12	not allow	0	not allow	0	3.8
M	15	-5	80	10	8	not allow	0	not allow	0	6.5
N	20	-15	50	10	-3	not allow	0	not allow	0	6.5
O	20	-15	50	10	-3	not allow	0	not allow	0	2.9