

Guber, Raphael; Kocher, Martin; Winter, Joachim

Working Paper

Does Having Insurance Change Individuals Self-Confidence?

Discussion Paper, No. 80

Provided in Cooperation with:

University of Munich (LMU) and Humboldt University Berlin, Collaborative Research Center Transregio 190: Rationality and Competition

Suggested Citation: Guber, Raphael; Kocher, Martin; Winter, Joachim (2018) : Does Having Insurance Change Individuals Self-Confidence?, Discussion Paper, No. 80, Ludwig-Maximilians-Universität München und Humboldt-Universität zu Berlin, Collaborative Research Center Transregio 190 - Rationality and Competition, München und Berlin

This Version is available at:

<https://hdl.handle.net/10419/185750>

Standard-Nutzungsbedingungen:

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

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

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

Terms of use:

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

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

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

Does Having Insurance Change Individuals Self-Confidence?

Raphael Guber (Munich Center for the Economics of Aging)
Martin Kocher (University of Vienna)
Joachim Winter (LMU Munich)

Discussion Paper No. 80

March 12, 2018

Does having insurance change individuals' self-confidence?*

Raphael Guber ^{†1}, Martin G. Kocher^{2,3,4}, and Joachim Winter⁵

¹Munich Center for the Economics of Aging, Max Planck Society, Germany

²Department of Economics, University of Vienna, Austria

³Institute for Advanced Studies, Austria

⁴Department of Economics University of Gothenburg, Sweden

⁵Department of Economics, University of Munich, Germany

March 9, 2018

Abstract

Recent research in contract theory on the effects of behavioral biases implicitly assumes that they are stable, in the sense of not being affected by the contracts themselves. In this paper, we provide evidence that this is not necessarily the case. We show that in an insurance context, being insured against losses that may be incurred in a real-effort task changes subjects' self-confidence. Our novel experimental design allows us to disentangle selection into insurance from the effects of being insured by randomly assigning coverage after subjects revealed whether they want to be insured or not. We find that uninsured subjects are underconfident while those that obtain insurance have well-calibrated beliefs. Our results suggest that there might be another mechanism through which insurance affects behavior than just moral hazard.

JEL-Codes: D84, D82, C91

Keywords: Overconfidence, insurance choice, underplacement

*We thank participants at the CEAR/MRIC Behavioral Insurance Workshop for helpful comments. Financial support by Deutsche Forschungsgemeinschaft through CRC TRR 190 is gratefully acknowledged. Guber acknowledges funding through the International Doctoral Program "Evidence-Based Economics" of the Elite Network of Bavaria.

[†]Corresponding author. E-mail: guber@mea.mpsoc.mpg.de.

1 Introduction

Self-assessments and beliefs matter in decision making and contract design. Optimal decisions depend on correct self-assessments and well-calibrated beliefs. One important example is self-confidence in own ability and performance. In particular, overconfidence has been established as a relevant aspect in individual's economic behavior. For example, overconfidence has been found to predict excess market entry of entrepreneurs (Camerer and Lovallo, 1999), risky investment decisions of CEOs (Malmendier and Tate, 2005), and speculative trading (Scheinkman and Xiong, 2003). In the context of insurance, Sandroni and Squintani (2007) consider the Rothschild and Stiglitz (1976) model in the presence of overconfident individuals. They find that if the share of overconfident types in the population is large enough, compulsory insurance is not Pareto-optimal anymore. It follows that overconfidence as a behavioral inclination has important implications for contract design in many settings (see for example Sautmann, 2013, De la Rosa, 2011 and Santos-Pinto, 2008).

While the majority of papers focuses on the case of overconfidence, situations in which individuals are underconfident are just as researched (Hoelzl and Rustichini, 2005; Clark and Friesen, 2009; Moore and Cain, 2007; Sautmann, 2013; Benoît et al., 2015; De la Rosa, 2011; Sandroni and Squintani, 2007; Sautmann, 2013). Imperfect self-confidence calibration relates to many effects observed in human decision making. However, a general interpretation of the literature on self-confidence is that over- or underconfidence are comparably stable traits, at least within a certain decision environment. That is, one can be overconfident when driving and underconfident with math tasks, but overconfidence when driving should not be affected by the color of the car.

This paper provides evidence for self-confidence to be malleable in a setting that has relevant implications. Our focus here is on over- and underplacement. Larrick et al. (2007) define the degree of an individual's overplacement as the difference between her perceived and actual percentile in the distribution of performance within a group. It differs from other concepts of overconfidence in that it depends on the believed performance of others. We show in a laboratory experiment that self-placement

depends on whether people acquire insurance or not. While insurance in our setup partially covers potential losses from bad performance in a real-effort task, it should be unrelated to self-placement for rational decision makers. At the same time, we find no evidence for more overconfident individuals choosing more or less insurance in the first place.

More specifically, we implement an experimental design that allows us to cleanly disentangle effects from the incentives provided by the insurance contract from effects coming from selection into the contract. Before attempting the real-effort task, individuals are given the choice to buy the insurance contract. Conditional on this choice, actual insurance status is randomized, i.e. whether one obtains insurance or not is based on a random draw, and individuals are informed about true their insurance status throughout the experiment. Our design is similar to the one used in a credit market field experiment by Karlan and Zinman (2009). Their idea is to attract borrowers with an advertised interest rate and, conditional on showing up in the lenders office, to randomize the actual interest rate. However, Karlan and Zinman (2009) are not able to impose an interest rate that is higher than the one advertised, as borrowers could simply walk out of the experiment. In a laboratory experiment, by design there is no attrition. This allows us to assess whether the effect of insurance on relative self assessment only comes from feeling (un-)lucky when actually (not) receiving it - remember, insurance status is based on a random draw - or whether there is another mechanism that is able to explain the effect. A related design is used by Bó et al. (2010), who let individuals vote on a policy that allows punishment for defection in a prisoners dilemma, but then randomize the actual implementation of the policy (see also Sutter et al., 2010).

Our real-effort task involves the forecasting of numbers with the help of two cue values (Brown, 1998; Vandegrift and Brown, 2003; So et al., 2017). This task fulfills two requirements for our purpose of creating a realistic insurance setting. First, the ability for forecasting, which might in the present case be related to math skills, varies sufficiently in the sample to create different levels of confidence in own ability. Second, the participant's effort can influence the precision of their forecasts and thus their relative performance. Schram and Sonnemans (2011) also consider insurance

choice by varying various parameters such as the number of available contracts. However, in their setting, losses occur without a subject's influence, which may not be realistic for some insurance contracts such as car insurance. Previous experiments studied insurance choice with exogenous loss in various settings, see for example Ganderton et al. (2000) and Laury et al. (2009). Our design naturally exhibits features of insurance markets outside the laboratory such as adverse selection and moral hazard.

Selfplacement is measured as the difference in an individual's self-assessed and true performance rank among all participants within the experimental session. The elicitation of the self-assessed rank is incentivized by rewarding accuracy. We find that, on average, insured individuals have well-calibrated beliefs about their ability relative to others, while those individuals that do not have insurance underplace themselves. These results are in line with experiments by Clark and Friesen (2009) and Murad et al. (2016), who argue that the use of real-effort tasks and incentivized confidence elicitation leads to a lack of overconfidence which is generally observed in "better-than-average" predictions. Moore and Cain (2007) and Hoelzl and Rustichini (2005) find that subjects tend to underplace themselves in tasks that are perceived as difficult and where performance is low in absolute terms, which is in line with our setup.

Our contribution is threefold. First, we show that individuals' self-confidence can be affected strongly by contracts. While in its generality, this result is probably not too surprising, its impact on our insurance application bears relevant implications – just imagine that drivers become relatively more overconfident after being insured. While contract design has started to take behavioral biases into account (Kőszegi, 2014), we are not aware of any existing model that would be consistent with our main finding. Second, we experimentally study assumptions made on the selection mechanism into contracts based on presumably stable traits such as self-confidence calibration (see for example Sandroni and Squintani (2007, 2013)). This paper thus speaks to a broader literature that studies sorting into contracts based on behavioral biases and preferences (Larkin and Leider, 2012; Dohmen and Falk, 2011). Finally, we add experimental evidence to decision making in a behavioral insurance context

in which own effort instead of a random device determines losses (Browne et al., 2015). We believe that such a setup adds to the external validity of our results for certain insurance classes.

2 Experimental Design

We start by describing the general procedure in our experiment, the real effort task and then the insurance decision. Monetary payoff was based on points, converted to euros at a fixed and pre-announced exchange rate. Participants received an endowment of 100 points, equal to EUR 10. The show-up fee for participants was EUR 4. The experiment was computerized with the help of z-tree (Fischbacher, 2007), and participants were invited with the organizational software ORSEE (Greiner, 2015).

2.1 Experimental Procedure

All steps in the experimental setup were known in advance and common knowledge among participants. However, we did not announce that we would elicit assessment of own relative performance after the real-effort task and insurance decision. The experiment consisted of three parts, and participants were aware of the existence of the three parts from the start of the experiment. They did, however, not know anything about the content of the following part until the end of the previous part. In the following, we just report results from the first part.¹ The experimental procedure for the relevant stages is illustrated in Figure 2.1, along with the variables generated at each stage. We explain the details for each stage below and in the subsequent sections.

In the first stage, subjects received a sheet of paper with ten examples of solutions in the real-effort task. The real-effort task was a forecasting task, and participants saw realized values of Y , W_1 and W_2 , which could be studied for five minutes, on the example sheet. A pen was provided, and participants were allowed to take notes,

¹The second part consisted of a set of lottery decisions; the third part was a short survey on relevant experience with insurance. Experimental instructions for the first part are provided in Appendix 5, and screenshots of steps 2 to 6 of the procedure can be found in Appendix 5.

which was done frequently. The second stage consisted of five practice rounds (five forecasts) with feedback on individual performance. These practice rounds were not incentivized, but there was an implicit incentive in the form of a potential information gain regarding one’s own ability in this task. In the third stage, individuals had to decide whether they wanted to buy the insurance for the upcoming payoff-relevant rounds or not. An on-screen calculator could be used at this point. The fourth stage randomized actual insurance receipt, and the choice made in stage 3 was realized with 70% chance. Thus, if a subject did not want to buy insurance, there was still a 30% chance that she got the insurance and that she had to pay the premium. Conversely, there was an equally large chance to not receive insurance, although the subject wanted to buy it. This creates a 2 by 2 matrix of possible outcomes shown in table 2.1. The probability of 70% was chosen trading-off incentive-compatibility and statistical power. A message informed participants about the realized insurance status. The message stayed on the screen throughout the following ten payoff-relevant rounds of the real-effort task in stage 5.

After the ten rounds of the real-effort task were completed, we elicited self-assessed performance in stage 6. Remember that this stage was not announced in the instructions. Individuals were asked to think about their average performance in the previous ten rounds and should indicate which rank they think that they hold in their respective session. The person with the lowest average forecasting error would take the first rank, the one with the second-lowest the second rank, and so on. At this point, subjects had not received any feedback on their or other participants’ performance. Guessing the rank correctly earned 10 additional points, and a deviation of plus or minus one from the realized rank earned 5 additional points. We chose to measure confidence in performance after the task, instead of before the task, in order to avoid hedging behavior and possible priming effects. Asking individuals about their relative performance to others before the task could give the wrong impression of a competitive environment, which we neither consider in this paper, nor is it common in an insurance context. We are well aware of the fact that linear incentives when eliciting beliefs have their limitations (see, Gächter and Renner, 2010; Trautmann and Kuilen, 2015), but for our case it seems a good com-

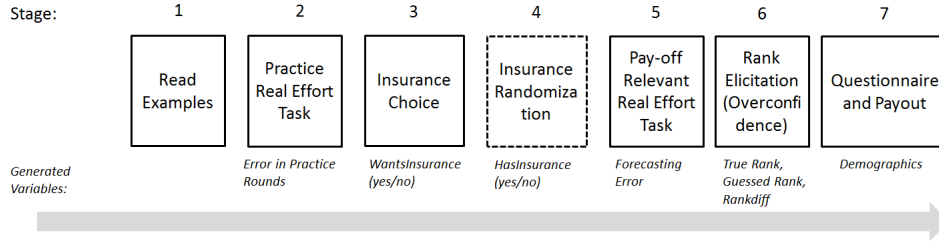


Figure 2.1: Experimental procedure and definition of variables.

promise between validity and straightforward implementation. Between stages 6 and 7, the second and third parts of the experiment took place. In stage 7, one of the ten real-effort task rounds was randomly drawn by the computer, and subjects were informed about their performance and earnings in this round. They also learned how much they earned from the ranking guess. At the end of the experiment, individuals answered a standard demographic questionnaire and were paid out in private.

Insurance status		actual		Total
		yes	no	
choice	yes	68	41	109
		41%	25%	
	no	13	45	58
		8%	27%	
Total		81	86	167

Table 2.1: Sample distribution

2.2 More Information on the Real-effort Task

We used the forecasting task by Brown (1998), Vandegrift and Brown (2003), and So et al. (2017). Participants are asked to enter the price Y of a fictitious stock whose price they had to predict from two cue values W_1 and W_2 . The true relationship of

Y and the two cues was given by

$$Y = 50 + 0.3W_1 + 0.7W_2 + e,$$

where $W_1, W_2 \sim U(0, 250)$ and $e \sim N(0, 5)$. Y was rounded to the nearest integer. Individuals knew that there was a potential constant, but did neither know that the function was linear, that the weights added to one, nor that there was a random error term e . During the task, individuals were shown W_1 and W_2 on the screen and had 60 seconds every round to enter their forecast $\hat{Y}$ into a box and click OK (see figure B.4 in the Appendix). The remaining time was always displayed on screen. There were no incentives for speed, but after 60 seconds without any input the program would skip to the next round, automatically creating a no-input. We introduced a penalty to avoid this, and the details are described in the next section. From the forecasting input we derived the error in each forecast, which is given by the absolute difference between the true and the predicted value of Y :

$$error = |Y - \hat{Y}|$$

2.3 Insurance

Based on a pilot of the real-effort task, we set the insurance premium to 22.5 points, with a coverage rate of 65%. Remember that only one round was payoff-relevant, i.e. the insurance was valid for all rounds. Earnings from the task are

$$earnings_{no} = 100 - error$$

for individuals that did not get the insurance and

$$earnings_{in} = 100 - error \times (1 - 0.65) - 22.5$$

for those that did. Thus, insurance covered 65% of the loss from the absolute difference between the true and the predicted value of Y . Notice that we capped losses at

the zero earnings boundary. As a consequence, there were no losses from this part of the experiment unless a participant had not entered any forecast at all for the randomly chosen round and was insured. In that case, the participant would have to pay the insurance premium of 22.5 points from her show up fee. This happened only once.

2.4 Experimental Participants

We conducted seven sessions in November 2015 in the MELESSA laboratory at the University of Munich. In total, 167 subjects participated and earned on average EUR 12.50 in a bit more than one hour per session. Participants were mainly students from various fields of study, with 33% from economics or business, 18% from life sciences or engineering and 13% from humanities. Almost 60% of participants were female, and age ranged from 18 to 43, with an average of 22.

3 Results

3.1 Descriptive Results on Self-Placement and Insurance Choice

We first look at a set of descriptive results. Our variable of interest is *rankdiff*, the difference between the individual’s actual and guessed ranks as entered in stage 6 of the experiment:

$$Rankdiff = TrueRank - GuessedRank.$$

A positive value indicates overplacement, where higher values imply stronger overplacement. A similar variable has been applied by Sautmann (2013), who uses the difference between predicted and actual scores in trivia quizzes as her measure for overconfidence. The mean of *rankdiff* in our study is -1.37 (which is significantly different from zero at the 5% level), indicating slight underplacement, on average. The distribution of *rankdiff* is shown in figure 3.1. The average underconfidence result is in line with Hoelzl and Rustichini (2005) and their task-specific explana-

tion. However, there exists considerable variation of *rankdiff* in our sample on the individual level and when comparing treatments. An alternative measure is a simple indicator variable for overconfidence. It takes on the value one if *rankdiff* is larger than zero, and the value zero otherwise. The entire sample has a share of 38.32% overconfident individuals according to this measure.

Remember that we can distinguish between four insurance outcomes, indicated by the variables *HasInsurance* and *WantsInsurance*. The variable *HasInsurance* describes the true insurance status of an individual in the real-effort task, and it is randomized. The variable *WantsInsurance* describes the individual's initial choice for or against insurance, and it is endogenous in the sense that it may correlate with any observed or unobserved individual characteristics such as gender, age and risk attitude. Conditional on insurance choice ($=WantsInsurance$), *HasInsurance* identifies the incentive effects of the insurance contract. Conditional on actual insurance status ($=HasInsurance$), *WantsInsurance* identifies selection effects, i.e. differences between individuals who wanted insurance and those who did not.

Table 3.1 displays means and standard deviations of *rankdiff* by insurance outcome. Table A.1 in the Appendix contains p-values of t-tests within every cell of table 3.1 whether the mean of *rankdiff* is significantly different from zero. In addition, table A.2 displays p-values of pairwise, two-sided Wilcoxon-Mann-Whitney tests for differences in *rankdiff* between all experimental groups. We observe strong and highly significant underplacement without insurance. There is, however, also significant underplacement for those who did not want insurance, when we pool observations for those who ended up with insurance and those who did not.

Two-third (109 out of 167) of individuals wanted to buy the insurance. We can investigate which individual characteristics predicted insurance choice. Table 3.2 shows mean values of these variables by insurance choice status and in the full sample. Individuals who made larger errors in the practice rounds were more likely to want insurance, which is in line with standard predictions of adverse selection models. *Insurance pays off* is a dummy equal to one if the forecasting error in a practice round was larger than $22.5/0.65=34.62$, which is the break-even point (error) of the insurance for a fully rational risk-neutral decision maker. There is a

	Wants Insurance=1	Wants Insurance=0	Total
Has Insurance=1	0.088 (7.39)	-0.46 (6.00)	-2.01 (7.67)
Has Insurance=0	-2.88 (6.99)	-2.46 (7.96)	-1.03 (7.41)
Total	-2.66 (7.56)	0.00 (7.23)	-1.37 (7.50)

Table 3.1: Mean and standard deviation of Rankdiff

large difference (20%-points) between those who wanted insurance and those who did not. However, buying insurance would still have paid off in 40% of rounds for those that did not want to buy insurance. Females more frequently wanted insurance than males and so did younger individuals.

	Did not want insurance	Wanted insurance	Total
Error in practice rounds	41.52	57.81***	52.15
Insurance pays off	0.40	0.60***	0.53
Female	0.36	0.67***	0.56
Age	23.33	21.42***	22.08

Insurance pays off is a dummy equal to one if the forecasting error in a practice round was larger then 22.5/0.65=34.62, which is the break-even point (error) of the insurance for a fully rational risk-neutral decision maker. Stars indicate mean differences significant at 1% (***), 5% (**), and 10% (*) level. Standard errors clustered at individual level in rows 1 and 2.

Table 3.2: Insurance choice

3.2 Regression analysis

We now turn to the effect of insurance on self-confidence and selection into insurance based on self-confidence by using parametric models. All regressions in table 3.3 use OLS estimations and include session fixed effects.² We start with performance in the real effort task in the first column. We find that having the insurance increases

²Ordered logit (for rank outcomes) and logit (for the overconfident dummy) models yield very similar results. The results are available on request.

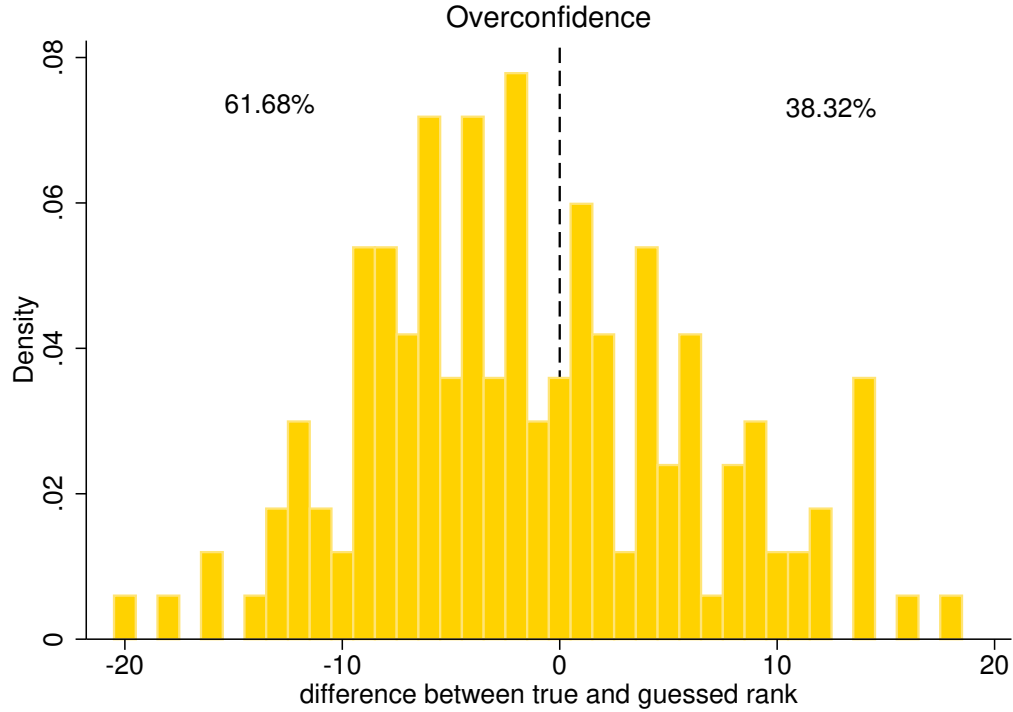


Figure 3.1: Distribution of variable rankdiff.

the absolute forecasting error by 4 points (or 0.15 standard deviations). The same difference is found between individuals who wanted and did not want insurance. The first effect is moral hazard and the second adverse selection, two classic elements in insurance markets (Shavell, 1979; Rothschild and Stiglitz, 1976). Column 2 shows the direct consequence of a lower performance in the task: both incentive and selection effects lead to a higher (i.e. worse) ranking within a session. Column 3 concerns the rank that individuals guessed they are taking. Individuals who ultimately got the insurance do not rank themselves worse or better than those who did not. In contrast, the pure selection effect in guessed ranks equals the one in true ranks. It follows in column 4 that insurance increases the difference between individual's guessed and actual rank by 2.7 ranks. Conditional on actual receipt, there exists no significant difference in self-confidence between those subjects that wanted and did not want

the insurance. This is in contrast to Sandroni and Squintani (2007), who assume that overconfident individuals are less likely to buy insurance, because they perceive their risk to be lower than is actually the case. We find that, on average, individuals anticipate their performance in the task based on their skill level and adjust their rank accordingly, but independent of the actual insurance status.

In the following we investigate if other biases specific to the experimental environment drive our results. One explanation could be that not getting the insurance despite wanting it leads to what is called "choking", a sudden decline of concentration and performance when individuals feel under pressure (Baumeister, 1984). This could lead to a severe underestimation of own performance, independent of its true level. Conversely, individuals receiving the insurance might feel lucky and thus rank themselves better than they actually are. These two confounding factors imply that the effect of the insurance on self-confidence should be larger among those individuals who also wanted it. In our 2 by 2 design, we can test for this possibility. Column 5 of table 3.3 shows that the interaction term between wanting and actually receiving the insurance is positive, but far from significant. The main effect of the insurance is not significant anymore, but the point estimate is similar to that in the columns before.³ Column 6 includes gender and age as explanatory variables to check if these explain the non-significant selection effect. Although the coefficient turns positive, it is not statistically significant and only one-third of the insurance effect. Columns 7 and 8 replicate columns 4 and 6 with a dummy equal to one if *Rankdiff* is positive as outcome variable and we get qualitatively similar results. The occurrence of overconfidence in ranking is increased by one-quarter under the insurance contract.

³This could also be due to lack of power, as the main coefficient of *HasInsurance* now refers to the insurance effect in the group that did not want the insurance and this group comprises only one-third of the sample. The insurance effect in the group that wanted the insurance is still significant at the 10% level.

Outcome:	(i) Error	(ii) True rank	(iii) Guessed rank	(iv)	(v) Rankdiff	(vi)	(vii) $\mathbb{1}\{\text{Rankdiff} > 0\}$	(viii)
HasInsurance	4.088** (1.729)	2.311** (1.147)	-0.649 (0.872)	2.960** (1.235)	2.443 (2.137)	3.157** (1.254)	0.240*** (0.082)	0.251*** (0.083)
WantsInsurance	4.032*** (1.544)	3.081*** (1.177)	3.303*** (0.893)	-0.222 (1.262)	-0.473 (1.710)	0.925 (1.400)	-0.016 (0.084)	0.042 (0.091)
Has $\times$ Wants Insurance Female					0.729 (2.709)	-1.651 (1.329)		-0.016 (0.080)
Age						0.391** (0.171)		0.031*** (0.010)
Constant	18.171*** (2.407)	9.368*** (1.730)	11.341*** (1.118)	-1.974 (2.174)	-1.943 (2.187)	-11.268** (4.793)	0.296** (0.114)	-0.475* (0.263)
Session f.e.	yes	yes	yes	yes	yes	yes	yes	yes
N	1,670	167	167	167	167	167	167	167
Adj. R-squared	0.017	0.056	0.053	0.000	-0.006	0.028	0.032	0.074

Rankdiff is the difference between the true and guessed rank of performance in the task. Individuals were incentivized to guess their rank among all participants in their session with respect to their average performance in the 10 payoff-relevant rounds of the forecasting task. No feedback on performance was provided. Robust or clustered (column 1) standard errors in parentheses. *** $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table 3.3: Insurance and overconfidence

4 Discussion

One major concern when trying to elicit self-assessment biases is to detect what Benoît and Dubra (2011) call apparent overconfidence. If individuals are Bayesian updaters and receive only a limited number of noisy signals on their performance, they might rationally rank themselves better than others, while this is interpreted as overconfidence by the researcher. This is less of a concern in our experiment, as individuals do not receive any signal on their (or others') performance in the payoff-relevant rounds. Their ranking should therefore solely be based on the perceived difficulty of the task over the ten rounds and an idiosyncratic component, which on average is the same between those that get and do not get the insurance, conditional on choice. Furthermore, Merkle and Weber (2011) demonstrate that the extent to which apparent overconfidence poses a problem in the laboratory is limited.

Another concern may be an insurance-induced change in a potential hedging motive when confidence levels are elicited. Since insurance reduces the downside risk in the real-effort task, the hedging motive in the elicitation loses importance. As a result, insured individuals could understate their performance less strongly than non-insured. However, this would imply that the insured place themselves at better ranks than the non-insured, which is not the case, as can be seen in column 3 of table 3.3. Another change in placement behavior arises if participants anticipate the lower performance of others, potentially induced by having insurance. Knowing that others will perform worse, they can place themselves better in the confidence elicitation. However, such higher order thinking applies to both treatment groups and should therefore be averaged out.

5 Conclusion

In this paper, we reported results of a laboratory experiment in which losses from a real effort task could be reduced by purchasing an insurance. After subjects revealed whether they want to be insured or not, insurance coverage was randomized. This novel design allows us to disentangle selection from incentive effects.

Self-confidence in the form of self-placement is measured as the difference between an individual's true and self-assessed performance rank. We find that, on average, uninsured individuals underplace themselves, while those individuals that obtain insurance have well-calibrated beliefs about their ability relative to others. While the previous literature is concerned about selection, we are the first to demonstrate that incentives irrelevant in standard economic models can change self-confidence ex-post. Moreover, we find no evidence for selection into insurance based on self-confidence.

Why does insurance coverage make individuals relatively less underconfident in their ability than uninsured individuals? One possible explanation suggested by our regression analysis is that individuals do not anticipate the moral hazard that is introduced by the insurance. Subjects do however anticipate their skill level and adjust their rank estimate accordingly. Put differently, the effect of the insurance is not reflected in an adjusted ranking, while the selection effect is. Another explanation involves the perception of the difficulty of the task. Under insurance, the task could appear easier, although in fact only the loss that subjects can incur in the real-effort task is lowered. As a consequence, underplacement is reduced. One can imagine alternative psychological explanations: for instance, insurance could let individuals focus more strongly on potential gains and thus the expected performance could appear more gloomy.

Our results have implications for insurance markets. Take car insurance as an example. Outside the laboratory it is next to impossible to distinguish between potential moral hazard effects and potential self-confidence effects. If both are present, the optimal policy of the insurer should take both into account. Remedies against moral hazard would not be enough to minimize unwanted behavioral tendencies, when we assume that biased self-confidence has negative consequences on driving. The experiment in this paper also has its limitations. For reasons explained above we do not have measures of self-confidence before randomization of the insurance. Further, we have no information on whether the induced self-confidence translates to other tasks and situations without insurance or on whether it is persistent or not. Ultimately answering this puzzle will require further research on why individuals

become overconfident in the first place.

References

- Baumeister, Roy F**, “Choking under pressure: self-consciousness and paradoxical effects of incentives on skillful performance.,” *Journal of Personality and Social Psychology*, 1984, *46* (3), 610–620.
- Benoît, Jean-Pierre and Juan Dubra**, “Apparent overconfidence,” *Econometrica*, 2011, *79* (5), 1591–1625.
- , —, and **Don A Moore**, “Does the better-than-average effect show that people are overconfident?: Two experiments,” *Journal of the European Economic Association*, 2015, *13* (2), 293–329.
- Bó, Pedro Dal, Andrew Foster, and Louis Putterman**, “Institutions and behavior: Experimental evidence on the effects of democracy,” *American Economic Review*, 2010, *100* (5), 2205–2229.
- Brown, Paul M**, “Experimental evidence on the importance of competing for profits on forecasting accuracy,” *Journal of Economic Behavior & Organization*, 1998, *33* (2), 259–269.
- Browne, Mark J, Christian Knoller, and Andreas Richter**, “Behavioral bias and the demand for bicycle and flood insurance,” *Journal of Risk and Uncertainty*, 2015, *50* (2), 141–160.
- Camerer, Colin and Dan Lovallo**, “Overconfidence and excess entry: An experimental approach,” *American Economic Review*, 1999, *89* (1), 306–318.
- Clark, Jeremy and Lana Friesen**, “Overconfidence in forecasts of own performance: An experimental study,” *The Economic Journal*, 2009, *119* (534), 229–251.
- Dohmen, Thomas and Armin Falk**, “Performance pay and multidimensional sorting: Productivity, preferences, and gender,” *American Economic Review*, 2011, *101* (2), 556–590.

- Fischbacher, Urs**, “z-Tree: Zurich toolbox for ready-made economic experiments,” *Experimental Economics*, 2007, *10* (2), 171–178.
- Gächter, Simon and Elke Renner**, “The effects of (incentivized) belief elicitation in public goods experiments,” *Experimental Economics*, 2010, *13* (3), 364–377.
- Ganderton, Philip T, David S Brookshire, Michael McKee, Steve Stewart, and Hale Thurston**, “Buying insurance for disaster-type risks: experimental evidence,” *Journal of Risk and Uncertainty*, 2000, *20* (3), 271–289.
- Greiner, Ben**, “Subject pool recruitment procedures: organizing experiments with ORSEE,” *Journal of the Economic Science Association*, 2015, *1* (1), 114–125.
- Hoelzl, Erik and Aldo Rustichini**, “Overconfident: Do you put your money on it?,” *The Economic Journal*, 2005, *115* (503), 305–318.
- Karlan, Dean and Jonathan Zinman**, “Observing unobservables: Identifying information asymmetries with a consumer credit field experiment,” *Econometrica*, 2009, *77* (6), 1993–2008.
- Kőszegi, Botond**, “Behavioral contract theory,” *Journal of Economic Literature*, 2014, *52* (4), 1075–1118.
- la Rosa, Leonidas Enrique De**, “Overconfidence and moral hazard,” *Games and Economic Behavior*, 2011, *73* (2), 429–451.
- Larkin, Ian and Stephen Leider**, “Incentive schemes, sorting, and behavioral biases of employees: Experimental evidence,” *American Economic Journal: Microeconomics*, 2012, *4* (2), 184–214.
- Larrick, Richard P, Katherine A Burson, and Jack B Soll**, “Social comparison and confidence: When thinking you’re better than average predicts overconfidence (and when it does not),” *Organizational Behavior and Human Decision Processes*, 2007, *102* (1), 76–94.

- Laury, Susan K, Melayne Morgan McInnes, and J Todd Swarthout**, “Insurance decisions for low-probability losses,” *Journal of Risk and Uncertainty*, 2009, 39 (1), 17–44.
- Malmendier, Ulrike and Geoffrey Tate**, “CEO overconfidence and corporate investment,” *Journal of Finance*, 2005, 60 (6), 2661–2700.
- Merkle, Christoph and Martin Weber**, “True overconfidence: The inability of rational information processing to account for apparent overconfidence,” *Organizational Behavior and Human Decision Processes*, 2011, 116 (2), 262–271.
- Moore, Don A and Daylian M Cain**, “Overconfidence and underconfidence: When and why people underestimate (and overestimate) the competition,” *Organizational Behavior and Human Decision Processes*, 2007, 103 (2), 197–213.
- Murad, Zahra, Martin Sefton, and Chris Starmer**, “How do risk attitudes affect measured confidence?,” *Journal of Risk and Uncertainty*, 2016, 52 (1), 21–46.
- Rothschild, Michael and Joseph Stiglitz**, “Equilibrium in Competitive Insurance Markets: An Essay on the Economics of Imperfect Information,” *Quarterly Journal of Economics*, 1976, 90 (4), 629–649.
- Sandroni, Alvaro and Francesco Squintani**, “Overconfidence, insurance, and paternalism,” *American Economic Review*, 2007, 97 (5), 1994–2004.
- and —, “Overconfidence and asymmetric information: The case of insurance,” *Journal of Economic Behavior & Organization*, 2013, 93, 149–165.
- Santos-Pinto, Luís**, “Positive self-image and incentives in organisations,” *The Economic Journal*, 2008, 118 (531), 1315–1332.
- Sautmann, Anja**, “Contracts for agents with biased beliefs: Some theory and an experiment,” *American Economic Journal: Microeconomics*, 2013, 5 (3), 124–156.

- Scheinkman, Jose A and Wei Xiong**, “Overconfidence and speculative bubbles,” *Journal of Political Economy*, 2003, 111 (6), 1183–1220.
- Schram, Arthur and Joep Sonnemans**, “How individuals choose health insurance: An experimental analysis,” *European Economic Review*, 2011, 55 (6), 799–819.
- Shavell, Steven**, “On moral hazard and insurance,” *Quarterly Journal of Economics*, 1979, pp. 541–562.
- So, Tony, Paul Brown, Ananish Chaudhuri, Dmitry Ryvkin, and Linda Cameron**, “Piece-rates and tournaments: Implications for learning in a cognitively challenging task,” *Journal of Economic Behavior & Organization*, 2017, 142, 11–23.
- Sutter, Matthias, Stefan Haigner, and Martin G Kocher**, “Choosing the carrot or the stick? Endogenous institutional choice in social dilemma situations,” *The Review of Economic Studies*, 2010, 77 (4), 1540–1566.
- Trautmann, Stefan T and Gijs Kuilen**, “Belief elicitation: A horse race among truth serums,” *The Economic Journal*, 2015, 125 (589), 2116–2135.
- Vandegrift, Donald and Paul Brown**, “Task difficulty, incentive effects, and the selection of high-variance strategies: an experimental examination of tournament behavior,” *Labour Economics*, 2003, 10 (4), 481–497.

Appendix A: Additional Figures and Tables

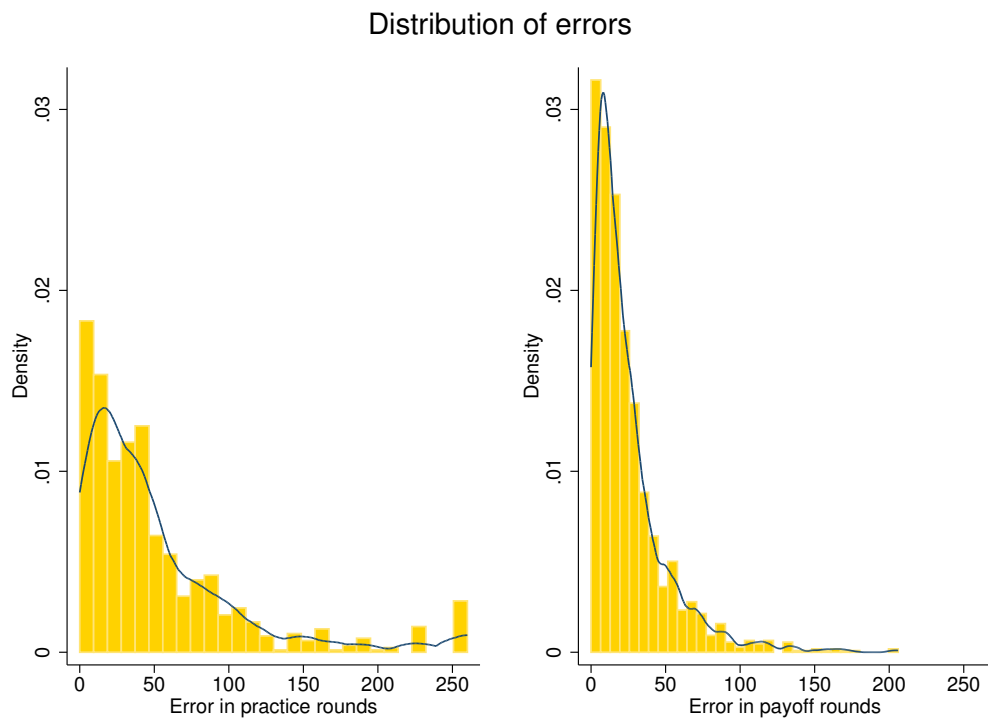


Figure A.1: Distribution of forecasting errors in practice and payoff-relevant rounds.

	Wants Insurance=1	Wants Insurance=0	Total
Has Insurance=1	0.922	0.794	1.000
Has Insurance=0	0.013	0.046	0.002
Total	0.151	0.050	0.019

Notes: Table shows p-value from t-test with the Null hypothesis that the mean of rankdiff equals zero within the respective cell.

Table A.1: P-values for zero mean t-test of rankdiff

Group 1	Group 2	p-value
has=1	has=0	0.021
wants=1	wants=0	0.445
has=1	has=0 wants=1	0.051
has=1	has=0 wants=0	0.287
wants=1	wants=0 has==1	0.862
wants=1	wants=0 has==0	0.839

Notes: Table shows p-value from Wilcoxon-Mann-Whitney test of a difference in rankdiff between experimental groups.

Table A.2: P-values from Wilcoxon-Mann-Whitney test of pairwise difference in rankdiff

Appendix B: On-screen instructions

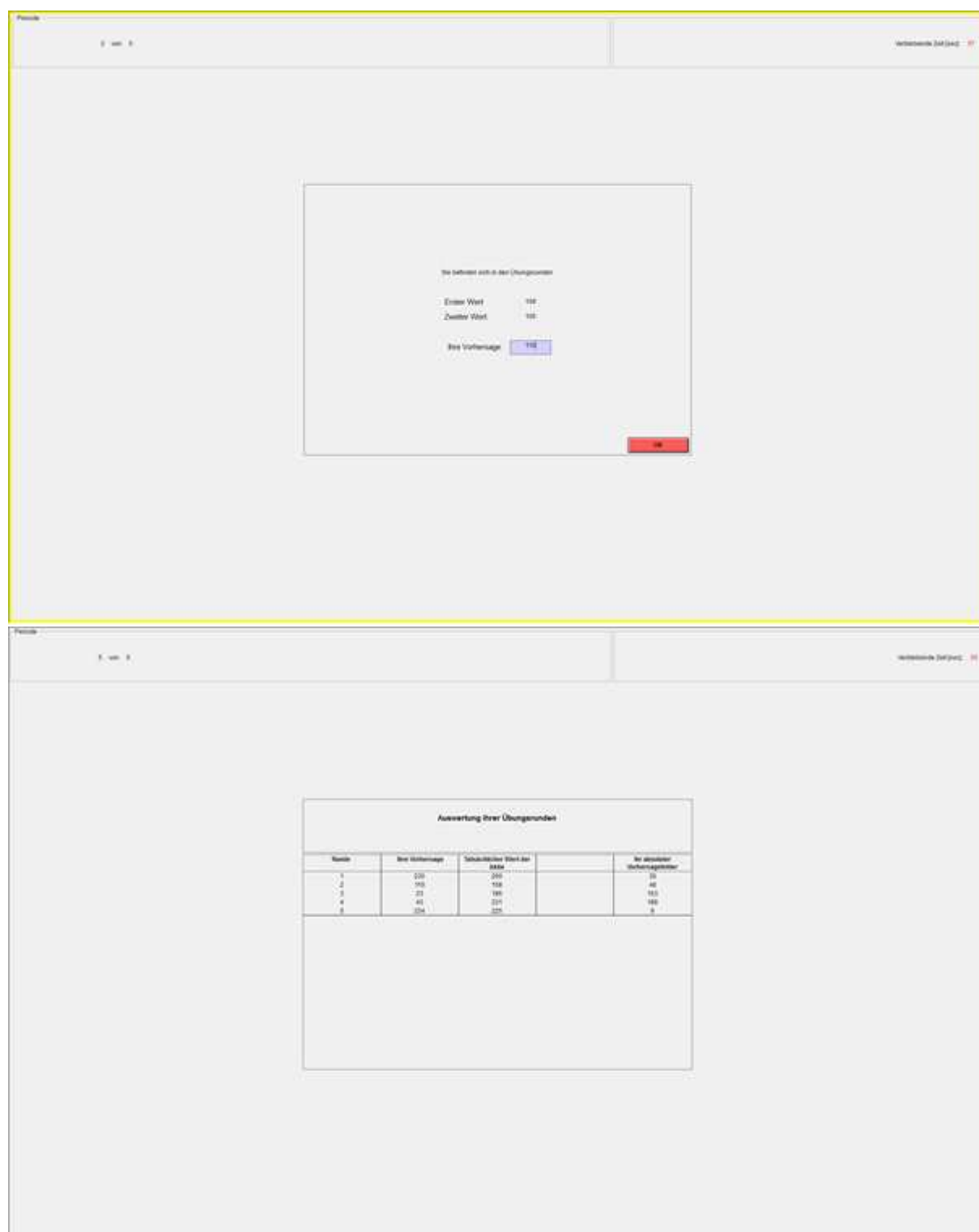


Figure B.1: Stage 2: The real effort task²⁵ in practice rounds and following feedback.

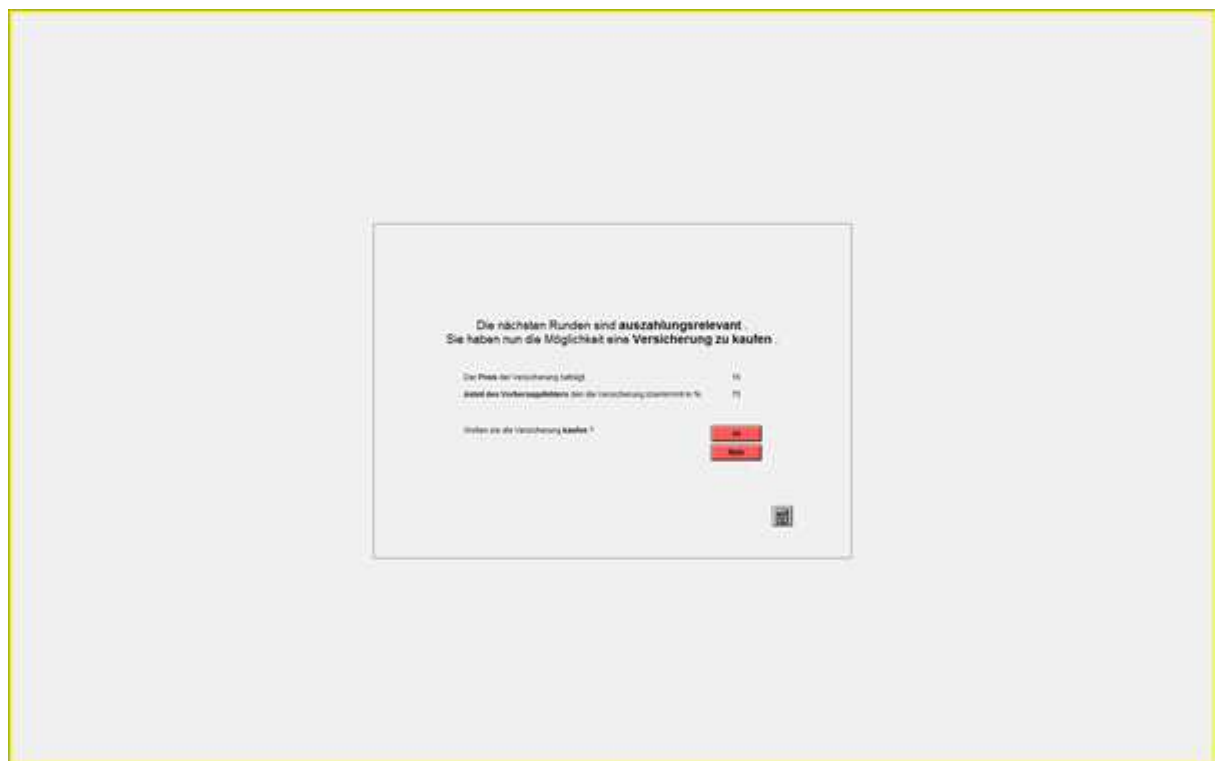


Figure B.2: Stage 3: Decisions whether to buy the insurance.

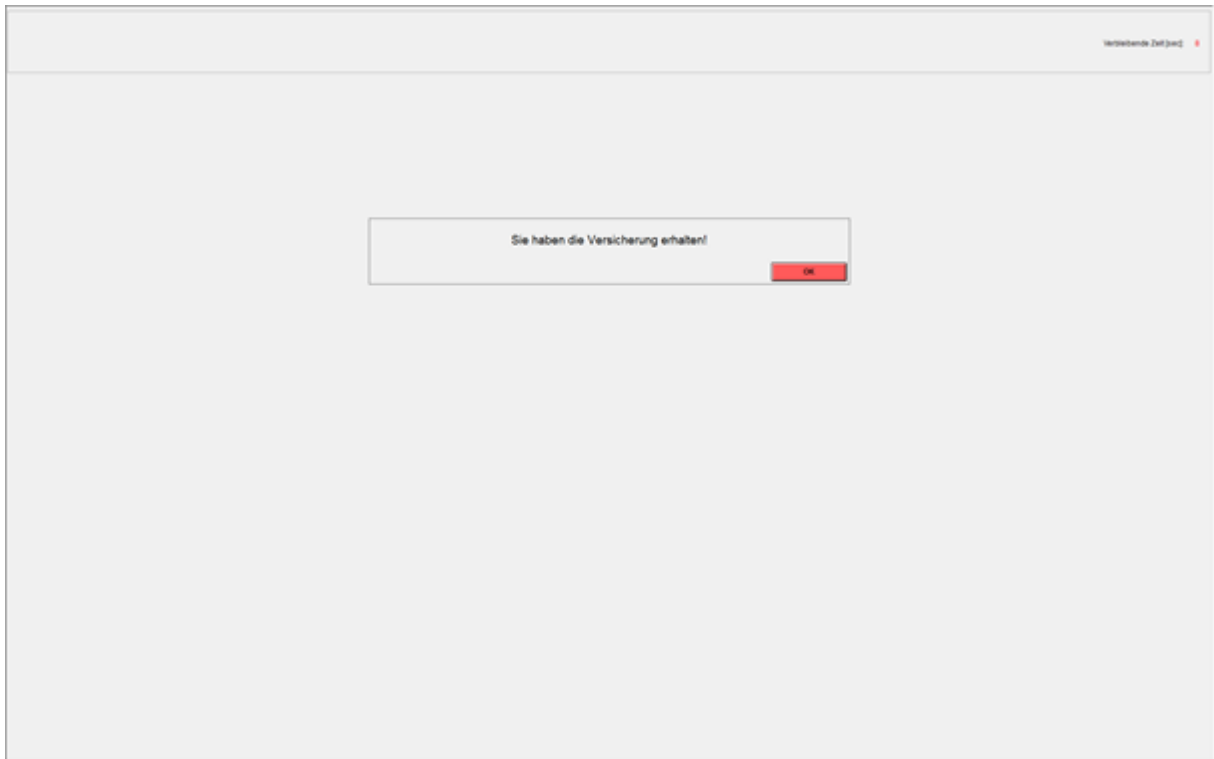


Figure B.3: Stage 4: Message on realized insurance status.

Problemlösung

5. von 10

Verbleibende Zeit (min): 10

Sie befinden sich in den entscheidungsrelevanten Runden.

Erster Wert	110
Zweiter Wert	100
Ihre Vorhersage:	

Wiederholungsrate

Info: Sie haben die Versicherung erhalten.

Figure B.4: Stage 5: The real effort task in payoff-relevant rounds.

Proble...

14 von 14

Verbleibende Zeit (s): 00

Bitte denken Sie an Ihren **durchschnittlichen Vortragsfehler** in den letzten 10 Runden.

Unter allen Teilnehmern im Raum, **welchen Rang** glauben Sie zu belegen?

(Der Person mit dem niedrigsten durchschnittlichen Vortragsfehler im Raum beträgt der erste Rang, die mit dem zweitniedrigsten der zweite Rang und so weiter.)

Geschätzter Rang

☐ 1	☐ 2	☐ 3	☐ 4	☐ 5	☐ 6	☐ 7	☐ 8	☐ 9	☐ 10	☐ 11	☐ 12
☐ 13	☐ 14	☐ 15	☐ 16	☐ 17	☐ 18	☐ 19	☐ 20	☐ 21	☐ 22	☐ 23	☐ 24

OK

Figure B.5: Stage 6: Ranking of own performance within session.

Appendix C: Experimental Instructions

Instructions are translated from German. Instructions were identical for all participants. Instructions from the second part of the experiment are not shown here.

Welcome to the experiment and thank you for your participation!

Please stop talking with the other participants now

General procedures

In this experiment we study economic decision making. You can earn money by participating. The money you earn will be paid to you after the experiment privately and in cash.

The experiment takes about 1 hour and consists of three parts. At the beginning of each part you will receive detailed instructions. If you have any questions about the instructions or during the experiment, please raise your hand. An instructor will then come to you and answer your questions privately.

Payment

Your profit will be denoted in points, where 10 points = EUR 1. In part I and II you will have to solve multiple rounds. Which round of a part is payout relevant will be randomly and with equal probability decided at the end of the experiment (part III). Since you do not know which round will be drawn, it is optimal to behave as if every round is payout-relevant.

At the end of the experiment your points will be converted into Euro and immediately paid out to you in cash. For showing up on time you receive EUR 4 in addition to what you will earn in the experiment.

Anonymity

The analysis of the experiment will be anonymous. That is, we will never link your name with the data generated in the experiment. You will not learn the identity of any other participant, neither before nor after the experiment. Also the other participants will not learn your identity. At the end of the experiment, you have to sign a receipt to confirm the payments you received. This receipt will only be used for accounting purposes.

Part I

Task

In this part, we ask you to forecast the price Y of a fictitious stock. To do this, you receive two values W_1 and W_2 , which underlie the price of the stock. You will not learn how exactly the price of the stock is formed out of the two values and a possible constant. However, you will receive examples for this relation, which **will not change** throughout the experiment. Please enter the predicted price of the stock into the respective window on the screen and click on OK. You have 60 seconds for this task. There are no advantages or disadvantages if you enter your solution faster than 60 seconds. You cannot change your input after clicking on OK. You can enter integer values between 1 and 500.

Procedure

At the beginning of the experiment you receive 100 points. 10 points are equal to EUR 1. To get a feeling for the relationship of the stock with the two values, you will once receive 10 examples at the beginning of the experiment on a piece of paper. You then have 5 minutes to study these examples. You can keep them for the rest of the experiment, but may not leave with them.

Next, you have the possibility to practice the task. There are 5 practice rounds with 60 seconds time each. After the five practice rounds you will be shown the true price of the stock, your forecast and the deviation of your forecast. The practice rounds do not influence your payout, but should help you in estimating your abilities for this task.

After the practice rounds the task will be done ten more times. This time, the accuracy of your forecast influences your payout. Every unit that your forecast deviates from the true value leads to a reduction of 1 point.

At the end of the experiment, one out of the 10 rounds will be chosen randomly and with equal probability. The forecasting error from this chosen round will be deducted from your 100 points. If the error is larger or equal to 100 points, you receive no payout from this part.

Insurance

Before solving the task, you have the possibility to buy an insurance. This insurance costs you once 22.5 points and is valid for all 10 rounds. The insurance reimburses 65% of your forecasting error. This means that, if you own the insurance, only 35% of your forecasting error will be deducted from your points.

However, it is not sure if you receive the insurance. In a first step you have to indicate if you want to buy the insurance. If you want to buy the insurance, you will actually receive it with a probability of 70%. With a probability of 30% you will not receive it. In this case you also don't need to pay 22.5 points. The reverse holds, if you indicate that you do not want to buy the insurance. With a probability of 70% you will not receive it, and with a probability of 30% you will receive it nevertheless and you have to pay 22.5 points.

After you decided for or against the purchase of the insurance, you will be informed if you received it or not. Then the 10 rounds start. Only at the end of the experiment will you know the correct value, your forecast and the deviation of your forecast. None of the other participants will ever be informed about your forecast, your choice or receipt of the insurance.

When choosing the insurance, you can activate a calculator by clicking on it symbol in the lower right corner on the screen.

Payment

The payout-relevant round will be drawn at the end of the experiment. If you did not receive an insurance, profit from this part of the experiment will be

$$(100 - |PriceStock - Forecast|) \times 0.1\text{EUR}.$$

If you did receive the insurance your profit will be

$$(100 - |PriceStock - Forecast| \times 35\% - 22.5) \times 0.1\text{EUR}.$$

If you do not enter any forecast within 60 seconds in a round and if this round

is chosen as payout-relevant you do not receive any profit from this part of the experiment, even if you have the insurance.

Let's look at some examples.

Example 1

After the practice rounds you decide against buying the insurance. You receive the message that you actually did not get the insurance. Now you perform the task 10 times. At the end of the experiment a random draw decides that round 7 is payout relevant. The true price of the stock in this round was 122. Your prediction was 170. The absolute difference of 48 will be deducted from your 100 points. Converted to euros you will receive $(100 - 48) \times 0.1 = 5.2$ Euro.

Example 2

After the practice rounds you decide to buy the insurance. You receive the message that you actually did get the insurance. Now you perform the task 10 times. At the end of the experiment a random draw decides that round 2 is payout relevant. The true price of the stock in this round was 99. Your prediction was 105, so your forecasting error equals 6. The insurance reimburses 65% of your error, or 3.9 points which will be rounded to 4. Hence, only 2 points will be deducted from your 100 points. However the price of the insurance of 22.5 points will also be deducted. Converted to euros you will receive $(100 - 6 \times 35\% - 22.5) \times 0.1 = 7.6$ Euro.

Example 3

After the practice rounds you decide to buy the insurance. However you receive the message that you did not get the insurance. Now you perform the task 10 times. At the end of the experiment a random draw decides that round 10 is payout relevant. The true price of the stock in this round was 150. Your prediction was 100. Since you did not get the insurance a full 50 points will be deducted from your 100 points. Converted to euros you will receive $(100 - 50) \times 0.1 = 5$ Euro.

Example 4

After the practice rounds you decide against buying the insurance. However you receive the message that you did get the insurance. Now you perform the task 10

times. At the end of the experiment a random draw decides that round 3 is payout relevant. The true price of the stock in this round was 175. Your prediction was 125, so your forecasting error equals 50. The insurance reimburses 65% of your error, or 32.5 points which will be rounded to 33. Hence, only 17 points will be deducted from your initial 100 points. However the price of the insurance of 22.5 points will also be deducted. Converted to euros you will receive $(100 - 50 \times 35\% - 22.5) \times 0.1 = 6.1$ Euro.

Examples for Part I

Here you find 10 examples on the relation of the fictitious stock Y and the two values W_1 and W_2 . The exact form of this relationship is identical in the examples, the practice rounds and the payoff-relevant rounds.

Y	W_1	W_2
137	73	95
160	152	85
175	79	152
151	100	87
115	76	49
85	27	37
212	219	139
129	244	7
203	14	217
90	69	25

Please leave this paper on the table when you exit the room.