



THURGAU INSTITUTE
OF ECONOMICS
at the University of Konstanz

Irenaeus Wolff
Dominik Folli

Why is Belief-Action Consistency so Low? The Role of Belief Uncertainty

Research Paper Series
Thurgau Institute of Economics
and Department of Economics
at the University of Konstanz

Member of



thurgauwissenschaft

www.thurgau-wissenschaft.ch

Why Is Belief-Action Consistency so Low? The Role of Belief Uncertainty[§]

Irenaeus Wolff

Dominik Folli

*Thurgau Institute of Economics, Hafenstrasse 6, 8280 Kreuzlingen, Switzerland &
University of Konstanz, Universitätsstraße 10, 78457 Konstanz, Germany*

wolff@twi-kreuzlingen.ch dominik.3.bauer@uni.kn

Forthcoming in: *Journal of Economic Behavior & Organization*

Abstract: Experimental research typically shows that best-response rates are below what plausible error rates would suggest. We experimentally test the conjecture that observed action-belief inconsistencies are related to belief uncertainty. We rely on a belief-sampling model that has been highly successful in explaining behaviour in multi-armed bandit problems and aggregate outcomes in games, markets, and surveys. Our data shows that inducing higher belief uncertainty leads more frequently to choices that are inconsistent with stated beliefs and—in an experiment directly testing the mechanism—to stochastic belief reports. The uncertainty-inconsistency relationship continues to hold when we control for error costs econometrically in several ways.

Keywords: Best Response, Belief Elicitation, Discoordination Game, Knightian Uncertainty, Errors, Elusive Beliefs

1 Introduction

Best-responses to one’s beliefs are a central building block in economic theory. However, empirically-measured best-response rates in economic experiments have been surprisingly low (see Table 1 for a non-exhaustive overview). Even

[§]We thank Fabian Dvořák, Thomas Hattenbach, Jordi Brandts, Miguel Costa-Gomes, Georg Weizsäcker, Ian Krajbich, Wieland Müller, Tomasz Strzalecki, Wolfgang Luhan, Alexander K. Wagner, Simon Gächter, Roberto Weber, Marie Claire Villeval, Dirk Sliwka, Nick Netzer, Felix Mauersberger, the research group at the Thurgau Institute of Economics, participants of the microeconomics seminar at the University of Konstanz, the Thurgau Experimental Economics Meeting (theem) 2017, and the ESA European Meeting 2017 for helpful comments. We are grateful to Miguel Costa-Gomes, Georg Weizsäcker, and Pedro Rey-Biel for sharing their data with us, and to Katrin Schmelz and Mo’ath Hussien for their help in the data analysis even when being sick. This paper has been circulated previously under the title “Elusive Beliefs: Why Uncertainty Leads to Stochastic Choice and Errors” (Wolff & Bauer, 2018). Declarations of interest: none.

if we allow for standard error rates in the range of 5-10%, the reported best-response rates seem excessively low. At the same time, there is substantial variance in reported best-response rates.

One of the dimensions on which prior studies have differed from one another is the complexity of the games used, where complexity would be related to the number of actions, to the number of equilibria, to whether games are dominance solvable, to whether payoff structures are symmetric, and, if they are not, to how asymmetric they are.¹ *A priori*, there is no reason for why more complex situations would make people respond less to their beliefs, unless more complex situations make people more uncertain about their beliefs, and unless this belief uncertainty is related to belief-action consistency in some way. This is the explanation we will be focusing on in this paper.

We introduce an experimental design targeted at exogenously varying participants' belief uncertainty by providing different levels of information. After observing that the observed best-response rate does not respond to the amount of information while simultaneously remaining around 68% only (in a four-action game), we re-invented a belief-sampling model that has been used in other parts of the literature for a long time, and which had recently also been imported into the economics literature by Celen, Geng & Li (2019) and Mauersberger (2021). The model accounts for the data from the experiment well.

To provide more thorough tests of its mechanism—to our knowledge, the first individual-level tests in the literature, we add two experiments. The model clearly predicts the results of the first additional experiment, and the second one provides further support for the assumption that best-response rates are negatively related to belief uncertainty.

Reference	Best response rate	Type of game
Costa-Gomes & Weizsacker (2008)	52% (59%)	3x3 games (beliefs before actions)
Danz, Fehr & Kübler (2012)	63%	variable-sum 3x3 games
Hyndmann et al. (2012)	66% (56%)	3x3 (4x4) games
Ivanov (2011)	68% (57%)	3x3 (2x2) games
Manski & Neri (2013)	89%	2x2 games
Nyarko & Schotter (2002)	75%	2x2 games
Rey-Biel (2009)	69% (65%)	const. sum (var. sum), 3x3 games

Table 1: Best-response rates to first-order beliefs in the literature.

¹We are grateful to an anonymous referee for making this point explicit.

Let us provide a little more detail on each of the elements above. To start with, let us motivate the idea of ‘uncertain beliefs’ by a simple example: We know for sure what the odds of a fair coin flip are. When offered a bet on this coin flip, it is easy to see whether it is worth accepting the bet or whether the odds-maker tries to trick us. Now imagine a colleague offers you a bet over a bottle of wine on your favorite football team winning the next match. It is the final match for the championship, your team is the home team, and your team performed better overall during the season. But then, one of the top scorers of your team is injured.

So, you start thinking about your belief on how likely your team is to win the match. In the heat of the conversation with your colleague, you are optimistic and think the odds are 60:40 that your team wins, so you accept the bet. However, on your way home, you think again of the top scorer’s injury and judge the odds to be down to 33:67. Over the next few days, you keep re-considering and re-“adjusting” your beliefs (without any new information coming in, and without you coming any closer to a stable estimate). If we took the average estimate, we might come to a winning probability of 50%.

For both the coin flip and the football match it would have been sensible to report a fifty-fifty belief when asked for it. However, there is a difference. For the coin flip, there is no sensible answer other than fifty-fifty. For the football-match, 50% is just one of many possible answers, and this assessment changed when re-thinking about the match. Arguably, people face this kind of uncertainty about the probability of events very often.

In the belief-sampling model, people have a (Dirichlet-distributed) prior over possible probability distributions. We call the resulting distribution over probabilities a *belief distribution*, because it is a distribution over different possible beliefs. In our examples, these would be distributions over the winning probabilities (with a degenerate distribution for the coin flip).

The difference to the standard rational-choice model (and to popular models of stochastic choice) is that agents do not reduce their belief distribution to a single probability (distribution), potentially because they have no direct conscious access to the belief distribution. Instead, agents draw from the belief distribution whenever they need to act on their belief, reacting to the randomly drawn probabilities as if they were the true probabilities. If the underlying belief distribution is spread out and many different beliefs are likely to be drawn—such as in the football-match example—the agent is *uncertain* about what the ‘true’, that is, the reduced probability is. We call the variance of the belief distribution its *belief uncertainty*.²

²See Pouget, Drugowitsch & Kepecs (2016) for a neuroscience perspective on uncertainty. Just like we do, the authors define uncertainty about some proposition as the variance of a posterior distribution (p. 369).

The model has been popular in operations research as well as some areas of psychology. In these literature strands, the model goes under the name of *posterior sampling* or *Thompson Sampling*. It has been shown to account well for how people deal with an exploration-exploitation tradeoff in multi-armed bandit tasks (e.g., Schulz, Konstantinidis, and Speekenbrink, 2015; Gershman, 2018).³ Mauersberger (2021) shows that the sampling model predicts also forecasting behaviour in a macroeconomic setting. In a recent working paper, Mauersberger (2020) further shows that the model predicts aggregate behaviour well also in other economic contexts. In particular, he shows that it clearly outperforms four competing models (Nash equilibrium, quantal-response equilibrium, Bayesian learning, and uniform randomisation) in a number of different games, markets, and surveys. Celen, Geng & Li (2019) also use the model to make sense of the decisions of their experimental participants, discarding quantal-response equilibrium for their setup.⁴

We will apply the model to a decision situation that resembles a game. Intuitively, when people are convinced they know what their interaction partners will do, they will report a belief and a choice that are consistent with each other because the belief draws underlying both decisions come from a narrow belief distribution. On the other hand, when people are uncertain about what they should expect their opponents to do, they may sample very different beliefs when prompted to come up with a belief report and when having to choose an action.

As pointed out above, our paper contributes an additional puzzle. In our study, participants play against historical data and we provide them with different amounts of information about this data before their decisions. Contrary to our expectations, however, the amount of information provided does not correlate with participants' observed best-response rates.

So, how would the model account for this lack of a correlation? Loosely speaking, the important insight is that more information on an uncertain event leads to more consistent behavior (i.e., a higher share of measured best-responses) *only when* the new information and the agent's prior belief are aligned. In contrast, when the new information contradicts the person's prior belief, the uncertainty in that person's posterior belief will increase in the amount of information (unless the amount of information is large). Intuitively, the person is less convinced of her mistaken prior belief.

³In the typical multi-armed bandit task, a participant has to draw cards from several stacks (with infinite supply of cards). Each draw yields an unknown reward, with unknown distribution. Participants thus have to explore the stacks and, after some sampling, face an exploration-exploitation tradeoff: to keep exploring several stacks or to exploit the information acquired by drawing only from the stack that has yielded the highest rewards so far.

⁴The authors argue that quantal-response equilibrium either does not fit the data (under one of two possible conceptualizations) or leads to a circular-reasoning problem (under the other).

Hence, more information may increase or decrease belief uncertainty, depending on the relationship between the prior and the additional information. On average, the two effects exactly cancel each other in our study. If we separate the data by a plausible proxy for the congruency of participants' prior and the information they receive, we observe both effects in our experiment.

After making their choice, we ask for participants' probabilistic beliefs. In our 'mechanism treatment', this procedure is followed by a repetition of the belief-elicitation task after a pause of seven seconds, without additional information being displayed.⁵ In the third experiment, we add an elicitation of the prior (according to our model, also merely a draw from a prior belief distribution) and an elicitation of participants' belief certainty with respect to their reported posterior.

The data generally support our research hypothesis: higher belief uncertainty leads to lower belief-action consistency. Importantly, this finding continues to hold when we control econometrically for the costs of making a mistake (which form the basis of, *e.g.*, quantal-response models) in a number of different ways. Moreover, the 'mechanism treatment' provides evidence for the model's main mechanism: participants often come up with different beliefs when asked twice, and the model predicts *when* they will do so. Furthermore, the differing beliefs in the 'mechanism treatment' rule out a large class of stochastic-choice models (all that respect first-order stochastic dominance).⁶ Finally, our third experiment provides additional evidence that it is indeed belief uncertainty that affects best-response rates.

A short history of the paper's development

As pointed out above, our initial idea was to show a relationship between the amount of information given to participants and the observed best-response rate. Not seeing such a relationship in a pilot experiment with a display of up to 40 historical data points, we ran our main experiment with up to 360. Still failing to see a direct relationship of displayed information and best-response rates, we sat down to model the situation and came up with the belief-sampling model. For

⁵In order not to push participants in any direction, the second elicitation was introduced as follows: "Today, we are interested in whether estimates change over time. Therefore, we would like to ask you once again to tell us your (current) estimate of the probabilities with which the participant of the previous experiment randomly allotted to you has chosen the individual boxes of the current arrangement."

⁶We incentivize all belief-elicitations using a binarized scoring rule which is incentive compatible as long as agents respect first-order stochastic dominance (Hossain & Okui, 2013). Thus, different measurements should not produce different results under the above class of stochastic-choice models, because there is no reason for underlying beliefs to change between measurements.

ease of exposition, we nonetheless will present also the main experiment in the standard hypothesis-results format. For the second and third experiments, the model then provided the basis for true *ex-ante* predictions.

2 Main Experiment

In our main experiment, we manipulate belief-uncertainty exogenously by giving varying amounts of information about the decisions of the relevant target population of other players. To be able to do so, we let participants play against historical data from an earlier experiment (which they know). Thus, participants in our experiment face a task that is *de facto* an individual-decision-making task, but that is mimicking a typical experimental game.

Our participants face a series of 24 discoordination tasks in each of which they have to choose one out of four labeled boxes. If they choose a different box than a randomly-selected participant of the earlier experiment, they receive 7€ and nothing otherwise. The randomly-selected ‘opponent’ and the labels of the four options vary across tasks and we use a large variety of letters, numbers or symbols as labels. For example, we start with labels “1,2,3,4” in task 1 and “1,x,3,4” in task 2. The complete list of all labels is depicted in Figure C1 in Appendix C. The order of the tasks is the same for all participants.

What was important to us was that participants did not need to worry about how their opponents would react to the information the participants were getting. We also needed a task in which participants would be facing different data-generating processes in each game, so that learning across tasks would be limited. However, we were looking for a task in which we were confident that participants would have their own private ideas of how most others would play the respective games. These requirements led us to having participants play against (individually) randomly-selected decisions from the study of Folli & Wolff (2022). In that study, 360 participants played the actual discoordination games on the same series of labels.⁷

⁷The two experimental series were run about one year apart and the only information participants received on the earlier sample was that “the experiment has been conducted earlier, here in the LakeLab.” Participants were recruited from the LakeLab database using the same criteria, making sure that nobody would participate in both studies. In Folli & Wolff (2022), we also had participants play a different participant in each of the 24 games with no feedback in between. In the main condition, we would incentivise participants for reporting their probabilistic beliefs after each game, in a control condition before each game, and in a further (unreported) condition in a separate block of tasks after having played all the games. Under any of the conditions, participants would be paid for one randomly selected game and the belief in a *different* randomly-selected game. We did not see systematic differences between conditions, and so we pool the data for our purposes here.

Before choosing an option and reporting a belief, participants enter an information stage in which they receive varying numbers of observations from the choice distribution they are playing against.⁸ The within-participant treatment is the number of observations that we sample and display to the participants, which we call n . The amount of information ranges from 0 to 360 with four periods of zero information.⁹ We randomize the order of the different n for $n < 360$ across participants and inform them that the decision of “the other player” is not contained in the displayed information. In the last period, $n = 360$ for all participants and thus, the information contains the relevant decision (which participants know).

In the experiment, we use a four-option setup. We do so for experimental reasons, acknowledging that there is a drawback in linking our experiment to the two-option setup that we use to present the belief-sampling model in Section 3: for the four-option case, we have to rely on simulations to support our predictions. We nonetheless prefer the four-option-experiment, because in a two-option setup, even a randomly clicking person would produce a best-response rate of 50%. Hence observed consistency will be generally high in this case, which makes it likely that we would face ceiling effects. In the four-option setting, the random best-response rate is reduced to 25%. Also, with four options, we can create much more variance in the label patterns than with two options. Thus, we can keep up participants’ interest for more rounds.

Along with the choice in each ‘game’, we elicit probabilistic beliefs after the action for every period. Participants have to report a set of four probabilities, one for each box. We incentivise the belief reports via a Binarized Scoring Rule (McKelvey & Page, 1990, Hossain & Okui, 2013). The Binarized Scoring Rule uses a quadratic scoring rule to assign participants lottery tickets for a given prize, in our case another 7€. The lottery procedure accounts for deviations from risk neutrality and, under a weak monotonicity condition, even for deviations from expected utility maximization (Hossain & Okui, 2013).

For the belief question, we use the opponent frame: “*What is the [respective] probability with which the participant of the preceding experiment you were randomly matched to chose the individual boxes of the current set-up?*” At the end of

⁸Note that the different labels induce ‘interesting’ data patterns in Folli & Wolff’s (2022) study. Choice distributions are significantly different from uniformity at a 5%-level in 15 out of 24 settings (χ^2 -test). 15 out of 24 settings are clearly more than the expected 1.2 settings under equilibrium behavior. See Table C1 for the data, and Figure D1 for an example screen of how the sample of earlier choices was shown to participants.

⁹The full set of information levels n is $\{0, 9, 12, 15, 18, 36, 64, 92, 120, 148, 176, 204, 232, 260, 288, 316, 345, 348, 351, 354, 360\}$. Note that we ran a pilot with n going from 0 to 40. Given that we did not see a correlation between the number of observations and the best-response rate, we wanted to make sure we had enough leverage to discover a potential effect. We thus extended the range of n as far as we could.

the experiment, we randomly select two different periods for payment. In one period, the outcome of the ‘game’ is paid and in the other period, the belief task is paid.

Procedures

The experiment was programmed using z-tree (Fischbacher, 2007). We use data of 55 participants recruited with ORSEE (Greiner, 2015). All sessions took place in the LakeLab at the University of Konstanz and lasted for approximately 75 minutes, including a short questionnaire at the end of the session which paid 5€. The last item of the questionnaire was a *reliability-of-answers* measure which gives participants the opportunity to indicate how reliable their data is in their opinion. The average payment was 13.27€.

3 The belief-sampling model

In this section, we present the belief-sampling model. Then, we relate it to observed best-response rates and present consequences of information updating for error rates at the end. For ease of exposition, we present the model for a two-option setting, before we derive the predictions for our four-option experiment through simulations in Section 4.

3.1 Belief sampling in a two-option discoordination task

Assume an agent is facing a two-option variant of the discoordination task from our experiment, in which the set of alternatives is $\mathcal{A} = \{L, R\}$. There is a probability $\phi, \phi \in [0, 1]$, that the current opponent chooses L (and a corresponding probability $1 - \phi$ that the opponent chooses R). However, the true probability ϕ^* of L -choices in the population is unknown and hence, the agent has to form a belief $\hat{\phi}$ about ϕ . In our discoordination setting, the agent strictly prefers R over L if and only if $\hat{\phi} > 0.5$.

In this paper, we assume that the belief is a non-degenerate probability distribution over all possible values of ϕ . For example, the player might assign a probability mass of 40% to ϕ being between 0.7 and 0.8, and distribute the remaining 60% of the probability mass over $[0, 0.7] \cup [0.8, 1]$. In general, the belief is a probability distribution $\phi \sim (\mu_q, \sigma_q)$ with continuous density function $q(\phi)$ where $\int_0^1 q(\phi) d\phi = 1$ and $q(\phi) > 0, \forall \phi \in [0, 1]$. Considering this belief distribution, the player faces a compound lottery: with density $q(\phi')$ the probability that R is better than L is ϕ' . However, in standard theory this subtlety does not play a role, as the best-response depends only on the *expected* probability the agent

assigns to R being the better option, denoted by:

$$E_q[\phi] = \int_0^1 \phi \cdot q(\phi) d\phi = \mu_q \quad (1)$$

In standard theory, the player will choose R whenever $\mu_q > 0.5$.

Stochastic choice and errors

In the belief-sampling model, the agent draws one value ϕ^r from $q(\phi)$ whenever the belief is consulted. This might be because players do not have direct access to $q(\phi)$, or simply because they are not able to compute μ_q . The draw ϕ^r is then used to determine the optimal action $a^*(\phi^r)$ instead of $a^*(\mu_q)$. Hence, not only the mean but also the whole distribution $q(\phi)$ matters for players' predicted choices.

In contrast to standard theory, players will make errors in our model. We define the error rate as the probability that the player draws a ϕ^r that does not indicate the same optimal action as μ_q , that is, the probability that $a^*(\mu_q) \neq a^*(\phi^r)$. Consider the example distributions in Figure 1 with $\mu_q > 0.5$ so that $a^*(\mu_q) = R$. Then the error rate ε is characterized by the probability mass of $q(\phi)$ on all $\phi < 0.5$, that is $\varepsilon = Q(0.5)$, and indicated by the shaded areas.

In our model, two characteristics of the belief distribution determine the error rate. First, as the mean μ_q approaches 0.5, the error rate ε increases (holding the shape of the distribution constant). The closer the belief is to indifference, the more errors are made, due to the shift of probability mass across the critical threshold. However, the model also allows for the case that a belief with an expected-utility difference (ΔEU , where $\Delta EU = |\mu_q - 0.5|$) close to zero produces little or no errors if the variance of $q(\phi)$ approaches zero. The model therefore also provides an intuition for *when* people will violate first-order stochastic dominance (FOSD) in their actions. Violations of FOSD are one of the greatest challenges for stochastic-choice models: people often violate FOSD when dominance is not obvious (because their belief over which option is better is uncertain). On the other hand, people respect FOSD when dominance is obvious (and thus, they know the best option exactly).

Second, for the error rate ε to increase, it is sufficient that the variance of $q(\phi)$ increases (holding μ_q constant). Consider again Figure 1. The shaded areas are the values of ε for two belief distributions with the same mean ($\mu_q^1 = \mu_q^2$) and hence $\Delta EU^1 = \Delta EU^2$, but different variances ($\sigma_q^1 \neq \sigma_q^2$). The more variance $q(\phi)$ has around its mean, the more likely the agent commits an error. When drawing from the (blue) high-variance belief, it is more likely that $a^*(\phi^r) \neq a^*(\mu_q)$ compared to a draw from the (red) low-variance belief.

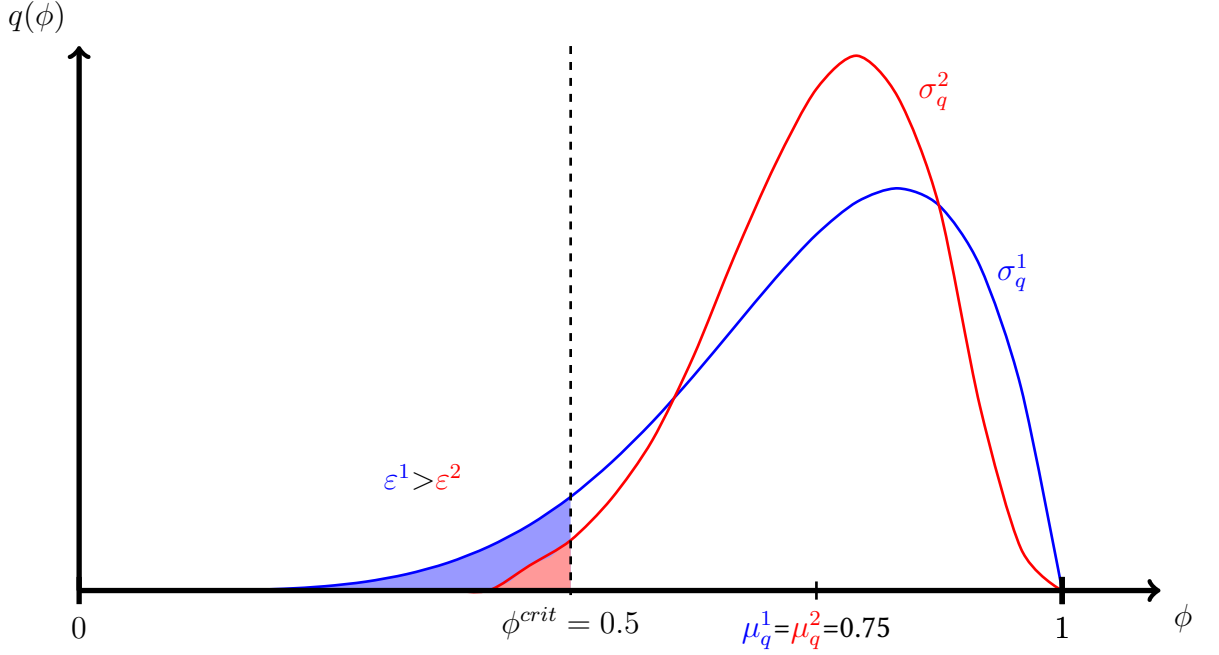


Figure 1: Two belief distributions with identical means but differing variances. The shaded areas indicate the respective error rates ε .

Stochastic beliefs

The notion of stochastic choice has consequences also for belief reports. In the usual experiment, choosing an action and reporting a belief are two separate decisions with different incentives. The reported beliefs are usually assumed to approximate μ_q and used to explain behavior. They are interpreted as the true cause of an action. Following Costa-Gomes and Weizsäcker (2008), we assume that not only the actions but also the belief reports are stochastic. Instead of calculating and reporting μ_q as a belief, the player also reports one draw ϕ^r as a belief. We assume that players use two (three) different and independent draws from $q(\phi)$ for the two (three) tasks of the main (‘mechanism’) experiment.¹⁰ Denote by ϕ_A^r the draw used for the action and by ϕ_B^r the draw for the belief report. Below, we will discuss the consequences of the combination of stochastic choice and stochastic belief reports for consistency in the main experiment.

So far, we have introduced the key idea that when making decisions and when reporting beliefs, agents draw realizations from their inner belief distribution.

¹⁰If a single draw were to determine both action and (both) belief(s), we would predict a 100% best-response rate (and only mutually consistent reports in the ‘mechanism treatment’) which definitely is rejected by the evidence in the literature as well as in our experiment.

We have characterized the error rate and demonstrated that knowing the reduced belief is not enough to predict the error rate. We now turn to the implications of the model for observed behavior in experiments.

3.2 Observable belief-action consistency and the error rate

We assume both choices and belief reports to be stochastic. Hence, the true belief distribution $q(\phi)$ and therefore also the true best-response rate and the true error rate are unobservable-in-principle in experiments.

For the experimenter to *observe* consistent behavior, that is, an action that is a best-response to the reported belief, the two draws from the belief distribution have to ‘fit together’. A best-response is observed only if $BR(\phi_A^r) = BR(\phi_B^r)$. In our example above, this is the case whenever both $\phi_A^r, \phi_B^r > 0.5$ or both $\phi_A^r, \phi_B^r < 0.5$. The expected observed best-response rate $\widehat{BR}$ is directly connected to the error rate ε defined earlier and can be characterized by:

$$\widehat{BR} = Prob\left[BR(\phi_A^r) = BR(\phi_B^r)\right] = \varepsilon^2 + (1 - \varepsilon)^2 \quad (2)$$

A best response is observed if an error occurs in either none or both of the draws ϕ_A^r, ϕ_B^r . We show in Appendix A that under relatively mild assumptions, the observed best-response rate decreases in the error rate ε in a symmetric game such as ours.

We next look at a possible determinant of belief uncertainty. A natural source of variation in the belief distribution—and hence also in belief uncertainty—is the integration of new information into the belief. To pave the ground for the hypotheses for our experiment, we will explore the influence of information integration on the error rate in the following section.

3.3 Updating based on new information

We assume that agents update their prior belief distribution in a Bayesian way when they get new information. This assumption needs some discussion given that agents may not have access to their belief distribution. To us, the assumption of Bayesian updating is a convenient technical assumption that simply is a (perhaps overly-)concrete specification of the assumption that agents’ beliefs will ‘move towards’ the information and not ignore new information altogether. This technical assumption makes the model much more tractable. It also makes the model much more comparable to the standard model because this way, belief sampling is the only deviation from the standard model (see also Mauersberger, 2020, for a discussion of ‘non-Bayesian decision-makers with Bayesian brains’).

It is left to future work to generalize the model to a more basic characterization of updating that is likely to include Bayesian updating as a special case.

From now on, let $q(\phi; \alpha, \beta)$ denote the participant's prior belief distribution. The mean of the Beta-distribution and hence the prior mean is $\mu_q = \frac{\alpha}{\alpha+\beta}$. The hyperparameter $\alpha = n_L^{Prior} + 1$ can be interpreted as the number of prior observations of L -choices in a sample of $n^{Prior} = n_L^{Prior} + n_R^{Prior}$ choices and $\beta = n_R^{Prior} + 1$ as the number of prior observations of R -choices.

Suppose the player observes a new sample of $n = n_L + n_R$ decisions from the population of N other players, where n_L denotes the number of L -choices in the sample. Because of conjugacy, the posterior is Beta-distributed as well. Hence now $\phi \sim \text{Beta}(\alpha + n_L, \beta + n_R)$. The posterior's mean can then be written as:

$$\mu_p = \frac{\alpha + n_L}{(\alpha + n_L) + (\beta + n_R)} = \underbrace{\frac{\alpha + \beta}{\alpha + \beta + n}}_{1-w} \cdot \underbrace{\frac{\alpha}{\alpha + \beta}}_{\mu_q} + \underbrace{\frac{n}{\alpha + \beta + n}}_w \cdot \underbrace{\frac{n_L}{n}}_{\mu_s} \quad (3)$$

The posterior's mean is hence a weighted combination of the sample- and the prior-mean. The weights are determined by the relative number of observations in the respective distribution where w denotes the relative weight of the sample. Further note that $\lim_{n \rightarrow \infty} \mu_p = \mu_s$.

The posterior's variance can be expressed as $\sigma_p = \frac{\mu_p(1-\mu_p)}{\alpha+\beta+n+1}$. It has two important properties. First, as $\frac{\partial \sigma_p}{\partial n} < 0$ the variance decreases *ceteris paribus* in n , the number of observations in the sample. Second, the variance is inverse U-shaped with a maximum at indifference, at $\mu_p = 0.5$. Hence, the variance decreases *ceteris paribus* in the distance of the belief mean to indifference $|\mu_p - 0.5|$.

The error rate of the posterior

As described above, the sample- and prior means as well as their relative weight determine the location and shape of the posterior belief distribution. In this section we derive predictions for the posterior's error rate ε based on characteristics of the prior and the observed sample. In the following, we continue to assume $BR(\mu_q) = R$ for simplicity, but all predictions hold symmetrically for priors with $BR(\mu_q) = L$. The most important characteristic is the location of μ_s relative to μ_q and to the critical threshold, in our case, to 0.5. There are three cases:

I) Congruent sample: The sample mean is the same or greater than the prior mean: $0.5 < \mu_q \leq \mu_s$.

- i) If $0.5 < \mu_q < \mu_s$ then ε decreases as the posterior mean is shifted to the right and hence, probability mass is shifted away from 0.5.

- ii) If $0.5 < \mu_q = \mu_s$ then ε decreases as the posterior variance decreases.

In both of these subcases, an increase of the relative weight of the sample w leads to an additional decrease of posterior variance and, hence, a larger decrease of ε .

II) Sample in between: The sample mean is less extreme than the prior but favors the same action: $0.5 < \mu_s < \mu_q$. In this case, the prediction depends on the relative weight.

- i) For a sufficiently small relative weight of the sample, ε will increase due to the shift of the mean towards 0.5 which is stronger than the minor decrease of variance.
- ii) For a sufficiently large relative weight of the sample, ε will decrease as the decrease of variance of the posterior will outweigh the effect of the shift towards 0.5.

III) Incongruent sample: If $\mu_s < 0.5 < \mu_q$, that means, if the sample mean is completely different from the prior mean and the two suggest different best-responses, it is *a priori* unclear which action the posterior will favor. The prediction depends again on the relative weight:

- i) For a sufficiently small relative weight of the sample, ε increases as long as the posterior mean μ_p is such that $\mu_s < 0.5 \leq \mu_p < \mu_q$. This means, the posterior mean approaches 0.5 from the right and probability mass is shifted to the left.
- ii) If the relative weight is large enough, the prior is ‘overturned’ by the information. Then, μ_s outweighs μ_q and $\mu_p < 0.5$. From then on, ε decreases in relative weight (from $\varepsilon^{max} = 0.5$ at $\mu_p = \lim_{\epsilon \rightarrow 0} 0.5 - \epsilon$).

Note that in Cases I and II, the posterior will always favor the same action as the sample because both $\mu_q, \mu_s > 0.5$. This also holds for the overturned beliefs in case III ii). However, if the belief is not overturned, the posterior will favor a different action than the sample in case III i).

4 Predictions

In this section, we specify the hypotheses for our experiment. We base the hypotheses on our model predictions in section 3.3 and on PROPOSITION 1 in Appendix A which states that the *observed* best-response rate decreases in the error

rate ε . We derive our hypotheses using the two-option model first. We then simulate the model for the four-option case to show that all predictions continue to hold for this more complex setting. Note also that we have to use proxies for some of the variables that are relevant in the model, because we cannot observe them directly by definition. In the following paragraphs, we discuss these proxies.

Approximating congruence of prior and sample

Our theory predictions and hypotheses mostly rely on the relationship of the sample mean μ_s to i) the sample's relative weight w and ii) the prior mean μ_q . Neither i) nor ii) can be observed in our experiment. First, the particular strength of participants' prior beliefs ($\alpha + \beta$) is unobservable, so we do not know w . Second, as the core idea of our theory, participants are not able to report μ_q , let alone $q(\phi)$. Fortunately, we have suitable proxies we can use.

We proxy the relative weight $w = \frac{n}{\alpha + \beta + n}$ by our treatment variable n , the number of provided observations. This proxy works well for weak priors and loses accuracy in the strength of the prior ($\alpha + \beta$). It hence could be that a participant by chance gets a high number of observations whenever her prior is particularly strong, and a low number of observations when she has only a weak prior. However, we randomize the treatment n across participants and games. Therefore, there is no reason to expect that such cases will systematically occur or dominate our data.

Second, we would like to be able to distinguish Cases I, II ii), and III ii) (more information leads to less errors) from Cases II i) and III i) (more information leads to more errors). The proxy we use essentially separates Case III i) from all other cases. This means that our analysis will underestimate the effect of an increase in w on belief-action consistency for observations of the first type, because observations for Case II i) are still 'in the mix' with the other non-Case-III-i) observations.

In particular, we proxy the relationship of the sample to the prior mean μ_q by the relationship of the sample to the *reported* belief ϕ_B^r . We compare what the best-response to both entities separately would be. Hence, we compare on which of the four options the participant places the lowest probability in her reported belief to where the minimum number of observations is in the sample.

If the reported belief, ϕ_B^r , has a different minimum than the information and hence also a different best-response, it is highly likely that the information favored a different response than the prior mean, μ_q (Case III i), and was not enough to 'overturn' the (reduced) prior. In particular, if participants were able to report their true posterior μ_p , it would *have to be* that the sample contradicted the participant's prior.

In contrast to that, if the reported belief has the same minimum as the information (that is, $BR[\phi_B^r] = BR[\mu_s]$) it is unlikely that the information differed completely from the reduced prior (Case I & II), unless the information ‘overturned’ the prior (Case III ii). We will further discuss the influence of ‘overturned beliefs’ on consistency below when we present our hypotheses.

As a summary, we proxy the relative weight of the information by its number of observations n . The relationship between prior- and sample-mean is approximated by a dummy which compares the *reported* belief to the information. $Belief-min = Info-min$ proxies Cases I, II and III ii). Both proxies should work well on average.

Hypotheses

In situations where the sample information favors the same action as the mean prior belief, the expected observed best-response rate virtually always increases in the relative weight (except in Case II i). The same is true when the differing prior is ‘overturned’. Given there is no reason to believe Case II i) will dominate in the experiment, we formulate

HYPOTHESIS 1: In cases where participants report beliefs such that $Belief-min = Info-min$, the observed best-response rate increases in the sample size n .

Note that not being able to observe priors poses a second challenge to Hypothesis 1: beliefs that have just been ‘overturned’ often enter the category $Belief-min = Info-min$ but have a high variance, which would speak against our Hypothesis 1. We nevertheless expect Hypothesis 1 to hold because we expect these cases to be rare enough not to dominate the data, either. In any case, not separating these cases from Cases I and II goes against our Hypothesis, so that we should have even more confidence in the effect in case we find it.

Case III i) is indicated by $Belief-min \neq Info-min$. Whenever participants report a belief with $Belief-min \neq Info-min$, it is highly likely that the belief was not overturned by the sample. This indicates a strong prior. However, because the provided sample differs from the prior, the sample shifts the posterior towards the critical threshold. Hence, in these cases belief uncertainty is generally higher, compared to cases with $Belief-min = Info-min$.

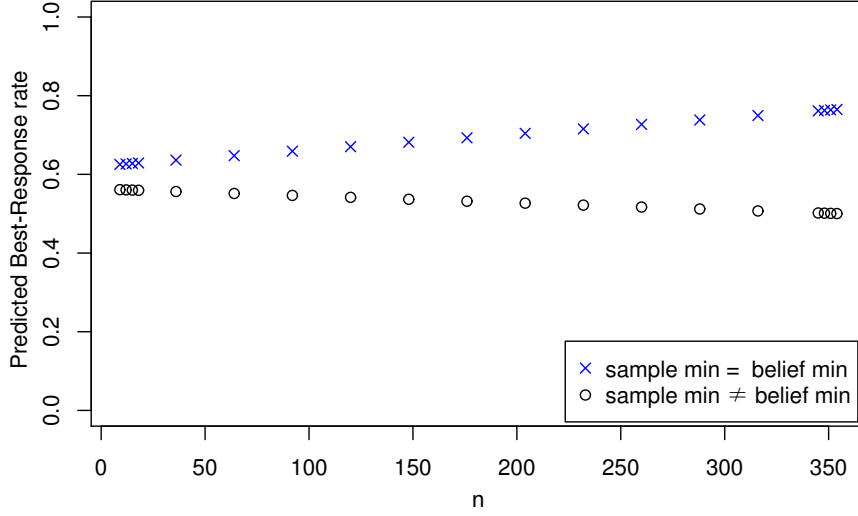


Figure 2: Predicted best-response rates in the four-option game according to our simulation.

- HYPOTHESIS 2A: In cases where participants report beliefs such that $Belief-min \neq Info-min$, the observed best-response rate is lower on average, compared to situations with $Belief-min = Info-min$.
- 2B: If $Belief-min \neq Info-min$, the observed best-response rate decreases in the sample size n .

For Hypothesis 2B, consider cases with $Belief-min \neq Info-min$ and relatively large sample sizes. In these instances, even the large n was not sufficient to overturn the prior. We hypothesize that in these cases the belief uncertainty must be particularly high, because posteriors will be close to the critical threshold. One could think that we should predict a U-shaped pattern for $Belief-min \neq Info-min$, because infinitely many observations should make players be certain in their beliefs. However, in this case, the posterior should have (close to) all of its probability mass on the belief that the probabilities in the given information are correct. Hence, cases with “infinitely many” observations almost always would fall into the $Belief-min = Info-min$ category. In the rare cases in which they would have $Belief-min \neq Info-min$, the action will be based on a second draw from the distribution which would virtually always result in a best-response to the information, and thus, not a best-response with respect to the reported belief. And hence, with “infinitely many” observations, the observed best-response rate in the $Belief-min \neq Info-min$ category will be very low. Consequently, there will not be a U-shaped pattern (which would require a high best-response rate for such cases).

So far, we derived hypotheses for a two-option setting. Our experiment uses a four-option setting, however, to reduce the ‘random best-response rate’ to 25%. To test whether the hypotheses carry over to our four-option setting, we simulated the model for the four-option setting as we were not able to derive the hypotheses for the four-option setting analytically. Hypotheses 1, 2A and 2B bear out in the simulation, as shown in Figure 2. We describe the details of the simulation in Appendix B. In short, the simulation is a one-to-one implementation of the model for the experimental setup and a large number of observations, once we replace the Beta-distribution for its multivariate generalization, the Dirichlet-distribution.

5 Results

Before we present our main analysis, let us give a brief overview of the data. On average, 70.2% of all observations are best-responses. 81.4% of the reported beliefs have a unique best-response, while 3.0% are uniform beliefs.¹¹ For ease of interpretation, we will restrict our analysis to those beliefs that have a unique best-response in the rest of the paper. Amongst these observations, 68.4% are best-responses.

If we run a simple unconditional linear regression of best-response rates on our treatment variable n , the coefficient is essentially zero ($0.006, p = 0.864$; this remains unchanged if we add participant random effects).¹² In terms of (non-)best-response types, we have only about 5% of participants who best-respond in 5 periods or less, 40% who best-respond in 17 periods or less, about 25% between 18 and 21 periods, about 25% who best-respond in 22 periods, and only about 10% who best-respond in 23 or 24 periods (almost all of them in 23). In the end, we have variation on the individual level for all but a single participant.

With respect to the information, 60% of the beliefs have the minimum at the same action as the provided information. This percentage increases slightly over time, by about 0.5 percentage points per round. It also increases in n as soon as we include random effects for individual subjects, by about 9 percentage points as we go from 0 to 360 observations, but there is no clear evidence that it would

¹¹One third of the uniform beliefs come from a single participant who reported uniform beliefs in 13 out of the 24 rounds; the remaining 26 uniform beliefs are well-spread over 15 different participants.

¹²The variation over games also is limited: there are only 2 games with best-response rates of more than 80%, and only four with rates below 60% (all four within the first six periods). Of the remaining 18 games, 13 have rates between 68% and 75%.

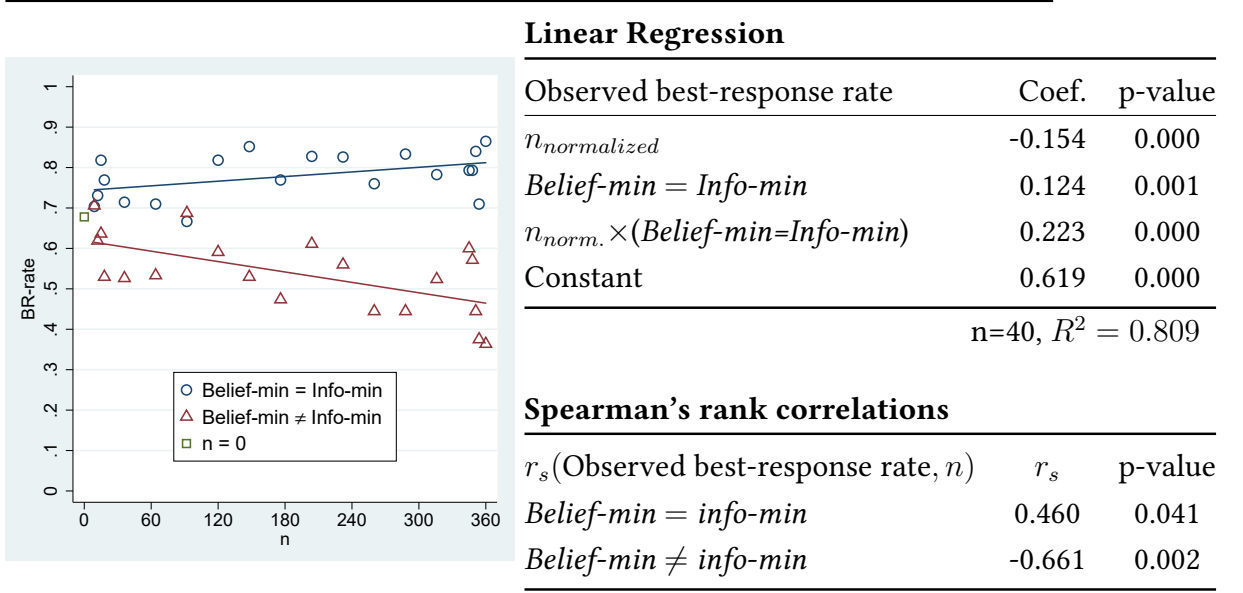


Figure 3: Observed best-response rates for each n , depending on the Belief-min–Info-min relationship. $n_{normalized} = \frac{n}{360}$ and $n > 0$ for the regression analysis and Spearman's rank order correlation. Each blue circle ($Belief-min = Info-min$) on average consists of 27 action–belief-report pairs, and each red triangle ($Belief-min \neq Info-min$) of 18 such pairs.

be affected by the variance in the provided information.¹³ At the same time, participants are surprisingly bad at guessing the earlier participants' behaviour: on average, only some 25.4% of the probability mass is placed on the other participant's true choice, uncorrelated with the period.

Coming to our main results, the left panel of Figure 3 gives a first impression. For the figure, we use all observations where the belief report has a unique best-response. For each value of our treatment variable n we compute the observed best-response rates across all participants, separately for both values of our situation proxy. If prior and information are not clearly incoherent, the best-response rates are increasing in n (HYPOTHESIS 1) and higher on average compared to situations with contradictory information (HYPOTHESIS 2A). Additionally, in the case when the information clearly contradicts the prior, the best-response rate decreases in n (HYPOTHESIS 2B). As a first statistical approximation, we add a linear regression and Spearman's rank correlations in the right panel of Figure

¹³Going from a degenerate distribution with all mass on a single action to a uniform distribution is estimated to reduce the percentage of 'information-following' by almost 20 percentage points, albeit with very low precision ($p = 0.599$). Also note that the effect of the increase in n seems to stem primarily from the highest information levels, as the ratio of observations does not change up until about $n = 316$. It clearly diverges only for $n = 354$ and $n = 360$, where it reaches ratios of about 1:2 and 1:3, respectively.

Best-response to reported belief	Average marginal effects					
	Model 1		Model 2		Model 3	
	Coeff.	p (s.e.)	Coeff.	p (s.e.)	Coeff.	p (s.e.)
Observations = 898, Clusters = 54						
$n_{normalized}$			-0.126 (0.051)	0.015 (0.051)	-0.128 (0.050)	0.010 (0.050)
$Belief-min = Info-min$			0.108 (0.057)	0.057 (0.057)	0.108 (0.055)	0.051 (0.055)
$n_{normalized} \times Belief-min = Info-min$			0.197 (0.096)	0.040 (0.096)	0.202 (0.095)	0.034 (0.095)
‘Strength’ of the reported belief	0.561 (0.309)	0.070 (0.309)			0.576 (0.299)	0.054 (0.299)
Period	0.009 (0.002)	0.000 (0.002)	0.008 (0.002)	0.000 (0.002)	0.008 (0.002)	0.000 (0.002)
Male	0.173 (0.076)	0.023 (0.076)	0.167 (0.076)	0.028 (0.076)	0.132 (0.073)	0.069 (0.073)
Mean Squared Error (Full Sample)	0.1961		0.1869		0.1826	
Mean Squared Error (Out of Sample for even Periods)	0.1960		0.1853		0.1814	

Table 2: Average marginal effects of logit regressions on observed best-responses. Standard errors in parentheses are clustered on the participant level (54 clusters). The interaction is computed using the `inteff` software by Norton, Wang & Ai (2004). See also Ai & Norton (2003). The marginal effect of the interaction is positive for all participants. Additional controls in all models: *age*, *math-grade*, *economics-student* and a self-reported *reliability-of-answers* measure.

3.¹⁴

The results are in line with the predictions of our model. We interpret the different situations created by the interaction of our *Belief-min = Info-min* dummy and our treatment variable n as different levels of belief uncertainty. As predicted in our model, observed best-response rates decrease in belief uncertainty. In the following, we present regressions that account for the dependencies in the data, as well as for certain particularities of individual decisions.

Accounting for error costs and learning

The results in Figure 3 use aggregate best-response rates across all participants and hence ignore individual characteristics and incentives. For a more meaningful analysis, we add controls for two additional influences on observed best-responses. First, we account for feedback-free learning over time by controlling

¹⁴If we revert the base category to *Belief-min = Info-min*, the regression coefficient for $n_{normalized}$ has a p-value of $p = 0.096$.

for the period in which the decision has been made. Second, we account for the cost of making an error. In Section 3, we already pointed to the potential effects of (low) error costs on the observed best-response rate. We account for both factors in the logit regressions whose average marginal effects we report in Table 2.¹⁵ Model 1 tests a model that only includes error costs, while Model 2 only includes belief-uncertainty and no error cost (like in Figure 3). Model 3 tests for both sources of errors jointly. Again, we use all observations with a unique best-response and $n > 0$.

Model 1 regresses individual best-responses on individual characteristics and the ‘strength’ of the belief report ϕ_B^r . By the strength of the belief we mean the utility cost of the cheapest decision error, assuming an expected-utility function (which makes the utility cost linear in the probability of a decision error).¹⁶ The strength of a reported belief is thus the percentage-point difference in beliefs on the options with the minimum and the second-lowest probability mass. If the strength is very low, the participant is almost indifferent between choosing the optimal or the second-best option and, according to a model in the spirit of quantal responses, has a high probability of making such an error. Thus, consistency will be low in these situations *independent* of belief uncertainty. The results of Model 1 suggest that the utility cost of making an error indeed have a large impact on belief-action consistency. High costs of an error strongly increase the probability of an observed best response.

Model 2 replicates our earlier results with respect to belief-uncertainty. The probability of a best-response decreases in belief uncertainty: when the prior is likely to be incongruent with the new information and strong, more information leads to lower best-response rates, as can be seen by the coefficient of $n_{normalized}$. When the prior was not likely to be incongruent, more information tends to increase consistency. To see that, we revert the baseline category and find a coefficient for $n_{normalized}$ of 0.073, with a standard error of 0.063. While the coefficient hence is not significantly different from 0, it is clearly positive, and consistently so throughout all model specification we have run.

Including both sources of error (the strength of belief and belief uncertainty) in the regression shows that the effect of belief-uncertainty is robust also when controlling for the utility cost of an error (Model 3). Finally, feedback-free learning over time leads to more best-responses in later periods in all three models.

To compare all three stochastic-choice specifications, we use out-of-sample predictions. We perform the regressions in Table 2 for all odd periods and predict

¹⁵The results are virtually the same when using the linear-probability model reported in Table C2 in Appendix C.

¹⁶All results are robust to a non-linear specification replacing the belief’s strength by its fourth-order polynomial.

the probability of a best response for each decision in all even periods.¹⁷ The bottom panel of Table 2 shows that the out-of-sample mean squared prediction error decreases from Model 1 to 3. To test the predictive power of the models, we compute the average squared prediction error of each model for every subject individually. The distributions of mean prediction errors differ between Model 1 and 2 (Wilcoxon signed-rank test, $p = 0.043$). This means model 2 outperforms model 1. Further, Model 3 outperforms both Models 1 and 2 (Model 1 vs 3: $p = 0.009$, Model 2 vs 3: $p = 0.083$).

Our results provide evidence that error costs alone cannot explain stochastic choice and belief-action consistency sufficiently in our data. Models 2 and 3, where we add our measures for belief uncertainty clearly outperform the pure error-cost Model 1 both in terms of fit to the data and predictive power. Hence, belief uncertainty plays an important role on top of error costs.¹⁸

Response times as an alternative measure of utility differences

Above, we use the strength of the reported belief as a measure of the utility cost of an error—hence as a measure for the strength of participants preferences. An alternative measure for the strength of preferences are response times. There is ample evidence in the literature that response times are closely linked to preferences: longer response times indicate that a person is close(r) to indifference between two options.¹⁹ In this study, the response time also may serve as an implicit measure of the strength of preference. This measure might be even less noisy than the strength of the reported belief because it does not rely on the participant’s belief report, which, after all, is stochastic according to our model.

We hence rerun our regressions, accounting also for response times. The regressions are reported in Table C4 in Appendix C. We include the normal logarithm of the response time (needed to select and confirm one of the boxes) as an additional explanatory variable in the set of logit regressions reported in Table 2. As expected, the extended models show that quicker response times are associated with higher belief-action consistency. This effect is in line with our

¹⁷The out-of-sample results are robust to predicting the choices of the second half of periods (13-24) by the first half of periods (1-12). However, models 1 and 2 do not differ significantly in that case (Wilcoxon signed-rank test, $p = 0.312$)

¹⁸It would be conceivable that ‘action uncertainty’—the variance of the reported belief—could play an important role for best-response play, too. To control for this possibility, we run another set of control regressions reported in Table C3 in Appendix C. If at all, including ‘action uncertainty’ in this way makes our focal coefficients larger. The same holds true if we instead use the variance in the displayed information or both measures simultaneously (unreported regressions).

¹⁹E.g., Mosteller & Nogee (1951), Moffatt (2005), Chabris *et al.* (2009), Alós-Ferrer *et al.* (2012), Dickhaut *et al.* (2013), or Konovalov & Krajbich (2017). Alós-Ferrer *et al.* (2016) even include the fact as a building block in their economic model to explain preference reversals.

above interpretation, that stronger preferences lead participants to committing fewer errors, which in turn leads to higher belief-action consistency. The effect of response times on consistency is robust to adding the belief strength, our original measure of the utility cost of making an error. Most importantly, though, the effect of belief uncertainty is robust to adding response times as an alternative measure for utility differences. Higher belief uncertainty still leads to less belief-action consistency when we include both measures for sources of stochastic choice—belief strength and response times—either separately or jointly. The effect of belief uncertainty becomes stronger, if at all.

The last period with full information

In the last period, participants saw the full choice distribution. Hence, (belief) uncertainty should be absent in this period. However, we still do not observe 100% best-responses. 11 participants (6.7%) even reported a belief with $Belief_{min} \neq Info_{min}$. We attribute these observations to other error sources, such as Fechner-type errors. It is also conceivable that some participants did not understand that there was no more uncertainty in this period: the data with $n = 360$ lines up perfectly with the rest of the results, as if there was some uncertainty left. All our main results hold (especially all regression results), when excluding the last period with $n = 360$ from the analysis.

6 Testing the mechanism

Having seen that the belief-sampling model explains our data so nicely, we wanted to go one step further and look at the model’s mechanism more closely. In the model, the decision-maker draws a belief from the belief distribution whenever she needs to act on her belief. In particular, when we ask participants for their beliefs more than once, the model predicts beliefs to ‘change’ sometimes—the more often the higher the belief uncertainty.

In this section, we set out to test this implication directly. To do so, we conducted two additional sessions with a total of 58 participants. The design of this ‘mechanism treatment’ was very close to the one in the main experiment. There were two main differences: (i) we exposed participants only to the first six of the frames, and gave them either (9, 9, 9, 354, 354, 354) or (354, 354, 354, 9, 9, 9) observations, respectively, sequences being allocated randomly to participants; and (ii), after participants gave their belief regarding the behavior of their respective opponent, there would be a pause of 7 seconds. After this pause, they read: “Today, we are interested in whether estimates change over time. Therefore, we would like to ask you once again to tell us your (current) estimate of the proba-

bilities with which the participant of the previous experiment randomly allotted to you has chosen the individual boxes of the current arrangement.”

The six rounds were the first part of an experiment that consisted of three independent parts; the content of parts 2 and 3 was not known to the participants at this stage. One of the parts was selected randomly for payment at the end. If part 1 was to be selected, we would pay one out of the six rounds according to their action, one according to their first-stated belief, and one according to their second-stated belief. To prevent hedging, we made sure that the three rounds selected for payment would all be different and that participants knew this.

To test the mechanism of our model, we set up the following hypotheses that come straight out of the model:

HYPOTHESIS M.1: Participants will state different beliefs for the same situation at least some of the time.

Similar to our main hypotheses in Section 4, the degree of belief uncertainty determines the degree of consistency of the two belief reports. High belief uncertainty will lead to a higher probability of stating different beliefs than low belief uncertainty:

HYPOTHESIS M.2: If participants’ first-stated belief is congruent with the information they receive, their second-stated belief will differ less often if they receive a high number of observations compared to when they receive only few observations. If participants’ first-stated belief is incongruent with the information they receive, their second-stated belief will differ more often if they receive a high number of observations.

Finally, having displayed consistent behavior in the action choice and the first belief report is a proxy for a narrow belief distribution. Thus, we posit:

HYPOTHESIS M.3: The belief of participants who choose a best-response to their first-stated belief is less likely to change than the belief of participants who do not best-respond to their first-stated belief.

Note that we deliberately formulated the question for their second belief estimate in an open way (“we are interested in whether estimates change over time.”) to avoid a demand effect. Still, we cannot exclude for sure that our assessment of Hypothesis M.1 is influenced by such an effect. However, there is no reason for

Number of consistent belief statements	0	1	2	3	4	5	6
Number of participants	2	3	11	14	13	10	5

Table 3: Numbers of participants for whom the first-stated and the second-stated beliefs have their minima at the same options.

why this effect should give rise to the pattern spelt out in Hypotheses M.2 and M.3.

To test Hypothesis M.1, we ask for each participant how often the minima of the first-stated and the second-stated beliefs coincide. Table 3 shows that 43 out of 58 participants (74%) state inconsistent beliefs at least twice, and 30 out of 58 participants (52%) state inconsistent beliefs at least half of the time. These numbers clearly support Hypothesis M.1.

To test Hypotheses M.2 and M.3, we run the regressions reported in Table 4. Model 1 shows that if the first-stated belief and the information are incongruent (baseline category), receiving a lot of information ($n = 354$) leads to less second-stated beliefs that are consistent with the first-stated belief; again, this reverses for first-stated beliefs that are congruent with the information (the sum of the coefficients for $\{n = 354\}$ and for $\{n = 354\} \times \text{Belief-min} = \text{Info-min}$ is positive). We therefore find support also for Hypothesis M.2.

Model 2 tests for the difference between cases where the first-stated belief and the chosen action were inconsistent (which proxies belief uncertainty) and cases where they were consistent (little belief uncertainty). Clearly, belief-action consistency predicts whether beliefs change (while the coefficients relating to Hypothesis M.2 are largely unaffected). Hence, the data offer support also for Hypothesis M.3.²⁰

7 Eliciting belief certainty and proxies for priors

As an additional—and more direct—test of our conjecture that belief uncertainty is an important determinant of best-response rates, and that it does so in the way

²⁰An additional test of the model’s mechanism is whether belief quality changes between the first-stated belief and the second-stated belief. The model would predict that this is not the case, as belief reports merely are independent draws from the belief distribution. The evidence is somewhat mixed. While the median difference in probabilities placed on the true choices of participants’ opponents is 0 and the first and third quartile are symmetric, too, the distribution as a whole is slightly right-skewed with a mean of 1.167 percentage points (a plain Wilcoxon signed-rank test yields $p = 0.076$). We are indebted to one of the anonymous reviewers for pointing us to this additional test.

Expected probability of consistent beliefs	Model 1		Model 2	
	<i>Coeff.</i>	<i>p (s.e.)</i>	<i>Coeff.</i>	<i>p (s.e.)</i>
(Intercept)	0.603	< 0.001 (0.061)	0.506	< 0.001 (0.066)
($n = 354$)	-0.175	0.055 (0.091)	-0.165	0.065 (0.089)
<i>Belief-min = Info-min</i>	0.072	0.377 (0.081)	0.058	0.468 (0.080)
($n = 354$) $\times$ (<i>Belief-min = Info-min</i>)	0.257	0.0264 (0.115)	0.236	0.0373 (0.113)
Action is a best-response			0.195	< 0.001 (0.055)
Log Likelihood	-192.115		-187.984	
Num. obs.	283		283	

Table 4: Linear-probability mixed-effects models on whether the first-stated belief will have the minimum (“*Belief-min*”) at the same option as the second-stated belief, with random effects for participants; standard errors in parentheses.

the belief-sampling posits, we add another treatment. The treatment follows the same procedure as the main treatment, with two exceptions. First, at the beginning of each round, we add another belief-elicitation stage before participants receive information about the choice distribution they are playing against. And second, after their report on the posterior belief, we ask them: “Please indicate how certain you are about the probabilities you have just indicated”, with a slider going from 0 for “absolutely uncertain” to 100 for “absolutely certain”.²¹

The analysis will follow the same lines as for the main part, but substitute an “*Info-min = Prior-min*”-dummy for the former “*Belief-min = Info-min*”-dummy. We then use a z-transformation to individually normalise the belief-certainty reports and include them in a second set of regressions. As we failed to run a proper power-analysis or a simulation of the model predictions, we had the following hypotheses:

²¹We thankfully acknowledge the suggestion of the additional measures (and the additional experiment) by an anonymous reviewer.

HYPOTHESIS P.1: If participants' prior belief is congruent with the information they receive, more information leads to a higher observed (posterior-)belief-action consistency. If prior and information are incongruent, the relationship is weaker.

HYPOTHESIS P.2: Participants' reported belief certainty is positively related with their observed belief-action consistency. Including belief certainty in the regression brings the coefficients relating to information-prior consistency and the amount of information (closer) to zero.

Note that we were too quick in following our anonymous reviewer in setting up **Hypothesis P.1**. Simulating the model predictions in analogy to the simulation from Section 4 clearly would have made us expect null-effects (see Figure C2 in Appendix C; unfortunately, while we were suspecting that the relationship would be much noisier than before when using reports of prior beliefs, we only re-ran the simulation after seeing the experimental results).

The key intuition for the model's surprising prediction of a close-to-null effect is not straightforward and comes from a combination of two effects. First of all, note that more often than not, the information will be incongruent with participants' priors (in the experiment we observe information-congruent prior reports in one third of the cases, which compares to one fourth if participants were guessing).

If prior and information are incongruent, there will be three cases as discussed before: the prior is relatively strong, which means that low levels of information hardly impact the posterior's low uncertainty. In this case, we are likely to observe a best-response, but on a different action than what the information would suggest.²² Alternatively, the information brings the posterior close to the point where the prior would be 'overturned', in which case a) best-response rates will be much lower and b) the belief-report still will often have the minimum at a different action than the information; this is what drives the effect we documented for the first two experiments.

In the third possible case, the information 'overturns' the prior, again leading to a low-uncertainty posterior, which in this case will have the minimum at the same action as the information. Note that in two out of three cases (relatively strong priors and 'clearly overturned' priors) we will observe high best-response

²²Judging by the data, this seems to account for roughly one third of the best-responses under prior-incongruent information (starting at two thirds for the lowest level of information and declining from there); under prior-congruent information, such cases are negligible (accounting for 4% out of 63% observed best-responses).

rates; in contrast, the *Info-Min* = *Belief-Min* proxy will exclude most ‘overturned’ priors (and, for higher levels of information, mostly include posteriors at the point where the prior is about to be ‘overturned’).

In terms of the second factor contributing to the explanation, note that participants more often than not will have weak priors (otherwise, the prediction would still be a close-to-null effect, but the mechanism would be somewhat different). This means that the prior-confirming effect of congruent information will play only a minor role, for two reasons: first, a substantial percentage of the information-congruent prior reports come from priors that are actually (when reduced) information-incongruent, and similarly, a substantial percentage of information-incongruent prior reports come from information-congruent priors. Second, when the prior is weak, the information will be the main determinant of the posterior almost irrespective of the prior.

Procedures. Payments were similar to the preceding experiments: participants would obtain the payment of three distinct rounds, one for each paid decision (prior elicitation, game action, and posterior elicitation). In each of these rounds, participants would earn either 6 Euros or nothing, depending on whether they discoordinated (in the game) or whether they were successful in the lottery associated with the binarised-scoring rule (for each belief elicitation). In addition, they received a fixed show-up fee of 10 Euros for online sessions and 12 Euros for in-person sessions.²³ We pre-specified that we wanted to obtain 150 data points or stop after 11 sessions, and exclude any participants who either report a uniform posterior in at least 20 of the 24 rounds, or play a discoordination-game choice that coincides with the action that is assigned the highest probability in their posterior in at least half of all rounds.²⁴ Unfortunately, following the procedure meant that we were able to collect the data of only 110 participants, one of whom chose to discontinue the experiment after round 3, and 30 of whom we had to exclude due to the second exclusion criterion.²⁵ We include the analyses for the full sample in Table C6 in Appendix C. Finally, note that we lost all belief-certainty reports from the first of the 24 rounds due to a programming error.

²³Due to a misunderstanding between us and our research assistants, we had to conduct the first four sessions as online sessions. Note that participants were recruited from our standard participant pool also for these sessions.

²⁴#151416 on asppredicted.org.

²⁵Note that this fraction came as a surprise to us, as hedging was excluded by design: we again stressed in the instructions that participants would be paid for *different* rounds for the game and the posterior-belief report. There is no indication that the frequency of worst-response play would be primarily due to the online sessions, even though the number is somewhat higher (25% of in-person participants vs. 32% online participants).

Best-response to belief rep.	Baseline analysis						With belief certainty					
Obs. in parentheses:	Model 1 (1226)		Model 2 (1226)		Model 3 (1226)		Model 4 (1170)		Model 5 (1170)		Model 6 (1170)	
Clusters = 79	Coeff.	p (s.e.)	Coeff.	p (s.e.)	Coeff.	p (s.e.)	Coeff.	p (s.e.)	Coeff.	p (s.e.)	Coeff.	p (s.e.)
$n_{normalized}$			0.018	0.661 (0.040)	0.037	0.390 (0.043)			-0.001	0.982 (0.045)	0.016	0.736 (0.047)
$Info-Min = Prior-Min$			0.017	0.728 (0.049)	0.019	0.692 (0.047)			0.018	0.721 (0.050)	0.016	0.737 (0.048)
$Info-Min = Prior-Min \times n_{normalized}$			0.028	0.707 (0.075)	0.019	0.803 (0.076)			0.006	0.938 (0.077)	-0.000	0.997 (0.077)
'Strength' of the ...reported belief	1.180	0.002 (0.379)			1.198	0.002 (0.386)	1.258	0.002 (0.399)			1.261	0.002 (0.406)
Normalized belief certainty							0.043	0.003 (0.015)	0.040	0.014 (0.016)	0.041	0.010 (0.016)
Period	0.008	0.001 (0.002)	0.007	0.002 (0.002)	0.007	0.001 (0.002)	0.009	0.000 (0.003)	0.009	0.001 (0.003)	0.009	0.000 (0.003)
Mean squared error	0.2255		0.2288		0.2252		0.2221		0.2260		0.2221	

Table 5: Average marginal effects of logit regressions of observed best-responses. Standard errors in parentheses are clustered on the participant level (79 clusters). The interaction is computed using the `inteff` software by Norton, Wang & Ai (2004). See also Ai & Norton (2003). Additional controls in all models: *age*, *math-grade*, *economics-student*, *male*, and a self-reported *reliability-of-answers* measure.

Results. Before we test our hypotheses, we quickly replicate the main analysis for the new data. The coefficients are roughly similar in magnitude and standard errors (with the coefficients relating to $n_{normalized}$ being somewhat smaller and the one relating to posterior-information congruency somewhat larger). We include the corresponding Table C5 in Appendix C.

Table 5 displays the results of the analysis pertaining to our **Hypotheses P.1** and **P.2**. In line with what we should have expected based on a simulation of the model (C2 in Appendix C), we do not find evidence for **Hypothesis P.1**: the estimated coefficients for $Info-Min = Prior-Min$ and its interaction effect with $n_{normalized}$ are much smaller than those in the main analysis, while the corresponding standard errors are similar. On the other hand, we do find support for **Hypothesis P.2**: the coefficients of participants' normalized belief certainty are consistently estimated to be positive, as posited by the model (and, for what it is worth, they reduce the coefficients relating to **P.1** further to zero).²⁶

8 Relationships to Other Approaches

A recent paper by Drerup, Enke & von Gaudecker (2017) nicely complements our paper. They econometrically relate the degree of belief uncertainty to the

²⁶A Wilcoxon matched-pairs signed-ranks test on participants' normalized belief certainty after trials registered as best-responses *versus* non-best-responses yields $p = 0.124$.

decision to participate in the stock market and show that beliefs are predictive only for those who have little uncertainty in their beliefs. Thus, while in contrast to our setting, they cannot manipulate uncertainty exogenously, their findings can nicely be explained by the belief-sampling model. They hypothesize that non-participants have high uncertainty because they choose not to get informed (because they rely on other choice rules, such as listening to friends and relatives who advise to stay away from stock ownership). While we cannot exclude that their hypothesized channel is at work, we can identify that the relationship runs (also) in the opposite direction. Non-participants may commit more mistakes because their beliefs are based on little information. Plus, if they anticipate (some of) this effect, they may shy away from stock ownership because they are afraid of committing expensive mistakes due to their lack of knowledge.

In a prominent contribution to the literature on belief-action consistency, Costa-Gomes and Weizsäcker (2008) find that, roughly speaking, their participants best-respond to uniform mixing in their actions but best-respond to their opponent's best-response to uniform mixing in their belief reports. While plausible, it is unclear how these findings would translate to the setup we study. First of all, if participants naïvely respond to the game, this may mean either uniform mixing (our participants are facing a pure discoordination game), or a best-response to the information they are shown. Our data correspond to neither of the two. It also is unclear what Costa-Gomes and Weizsäcker's results could mean for the belief reports in our setting: there is no obvious way of thinking more strategically about the situation our participants are facing. Most importantly, it is not obvious how their findings would lead to the treatment effects we find.

Many of the ideas behind our motivation for this study can be found already in the literature on choice under uncertainty.²⁷ The main question in this literature is how people make their decisions when they face uncertainty and there is no clear way of assigning probabilities to the possible states of the world. This literature departs from Savage's (1954) idea that when agents face ambiguity, they simply will form subjective beliefs and act on those subjective beliefs as if the beliefs were proper probabilities. There is a whole array of how the corresponding non-Bayesian subjective probabilities are modelled, and how they are used by the agents.²⁸

The approach that probably is closest to ours in that agents decide based on a single potential belief realization is the multiple-priors approach as axiomatized in Gilboa & Schmeidler (1989). In models following this approach, agents choose

²⁷Even our introductory example is similar to the examples given in this literature, *cf.*, *e.g.*, Gilboa, Postlewaite & Schmeidler (2008).

²⁸For a review, *cf.* Etner, Jeleva & Tallon (2012), who also discuss important economic applications of the models.

among the alternatives using a maximin-utility criterion across all probabilities they consider possible (e.g., Gilboa & Schmeidler, 1989). Other approaches use some form of expectations maximization (over a ‘probability distribution that does not add up’: Schmeidler, 1989; or over some transformation of certainty equivalents: Klibanoff, Marinacci & Mukerji, 2005). In either case, the models are about agents who consistently make a specific type of choice, namely ambiguity-averse choices, for example in Ellsberg-type settings. Our aim complements this literature, as we focus on explaining the variance within people’s choices, and on the likelihood of observing inconsistent choices.

In the huge and important literature on stochastic choice, the random-error model (like in Harless & Camerer, 1994) is a popular model. According to the model, agents make mistakes with a fixed probability. This model seems to be at odds with the low belief-action-consistency rates referred to in Table 1, as the error probability usually is assumed to be in the range of 5-10%. On top, the model would not predict our treatment effects as the error rate is assumed to be constant across decisions.

Another popular type of model are models incorporating Fechner-type errors (e.g., Luce, 1959; McKelvey & Palfrey, 1995; also known as logistic-choice or quantal-response models; in binary settings, the standard drift-diffusion model of Ratcliff, 1978, will also lead to the same choice distribution). In this type of model, agents make mistakes more often, the smaller the utility difference is between the available alternatives. In other words, what is important for choice consistency is the expected cost of making a mistake. We control for error costs in our statistical analyses. The conceptual difference between such models and belief sampling is best seen in an example: consider an agent in a two-player coordination game with options X and Y . The game pays 1 Euro in case both players choose the same action. If the agent is convinced that the other player chooses X with a probability of 51%, the agent always chooses X under belief sampling (because the agent always samples the same (51%, 49%)-belief). In contrast, a logistic-choice model would predict a close-to-uniform choice distribution. We expected error costs to matter, which is also in line with our findings. What our paper shows is that belief uncertainty matters *on top of* error costs.

Two additional types of stochastic-choice models are models that rest on the agent either having random preferences (as, for example, in Becker, DeGroot & Marschak, 1963, or Loomes & Sugden, 1995, where the parameters of the utility function are drawn from a distribution before each choice) or on the agent’s uncertainty about her true utility function (as, for example, in Cerreia-Vioglio *et al.*, 2015). In both cases, however, participants should always produce the same belief report when asked twice under a binarized scoring rule, because both types of model respect first-order stochastic dominance.

Given the mechanism of the belief-sampling model, two additional mod-

els come into mind: Action-sampling (claimed by Selten in Selten & Chmura, 2008) and payoff-sampling equilibrium (Osborne and Rubinstein, 1998). Action-sampling equilibrium represents the long-run outcome when agents best-respond to a restricted sample of their opponents' prior actions. Payoff-sampling equilibrium corresponds to the long-run outcome when agents sample for each of their actions a fixed number of payoffs they achieved using the corresponding action, then choosing the action with the highest payoff. For both concepts it is not clear how they would fit our setup. We present participants with samples of different sizes of what (unrelated) others have done. Given that in many cases, they see a distribution of many previous choices, why would they arbitrarily eliminate some of those choices? For if they don't, both models would predict consistent best-responding to the provided information (which is not what we find).²⁹

Last but not least, the belief-sampling model is related to the vast literature on learning. By introducing a different form of stochasticity, the model also provides a new perspective on learning in unknown situations. When facing a decision for the first time, belief uncertainty is likely to be high. In the model, this leads to high error rates. Hence, there is scope—and need—for learning. As the situation is repeated with feedback, the agent gathers more and more observations. In most situations, gathering more information will decrease the variance of the belief distribution, leading to less errors. Hence, the agent learns how to behave in the situation by identifying the situation better and better, *even when there is no change in the reduced belief*. This sets us apart, for example, from models of belief-based learning like fictitious play or Cournot learning.

With respect to our experiment, the only type of learning that could interfere with the conclusions from our feedback-free environment is feedback-less learning (Weber, 2003). According to this idea, experiment participants learn how to play a game even without feedback. We therefore should see increasing best-response rates over time. We address this potential confound by (individually) randomizing the order in which participants receive the different sample sizes. On top, we control for the period in our analysis, and hence, implicitly also for any form of feedback-less learning on how to play a best-response.

9 Summary and Conclusion

In this study, we had participants play against an unknown process after giving them different amounts of information about this process. We then asked participants for their belief about the process and checked whether they had been playing a best-response to their stated belief. Intriguingly, the raw data do not

²⁹Also, at least the payoff-sampling equilibrium is psychologically unintuitive for quasi-continuous action spaces such as our belief-elicitation task: How would anybody draw a sample for each of the innumerable potential belief reports?

show a positive correlation between the frequency of best-response play and the amount of information provided.

Does this mean that standard theory is correct because the amount of information does not have an influence on participants' best-response rates? Both a belief-sampling model and our data analysis contradict this possibility. As soon as we take into account that some participants will have had a non-uniform prior that is 'incongruent' with the displayed information and proxy for it, we find a clear relationship between information and best-response rates. When the prior is very likely to be at odds with the new information, more information leads to fewer best-responses. Agents become more uncertain about their (likely mistaken) belief, which in the belief-sampling model leads to higher rates of inconsistent behavior. When the prior is likely to be neutral or in line with the new information, more information instead tends to entail higher best-response rates. Agents become more certain about their belief, which translates into more consistent belief-action pairs.

We conducted a second treatment that is meant to test the mechanism behind the model more closely. This 'mechanism treatment' provides evidence for the idea that participants sample their beliefs anew every time they have to act on their belief—much like the person in our initial example who has to decide on whether to accept a bet on her favourite football team winning the next match. In principle, it would be conceivable that inconsistent belief reports merely are an example of stochastic choice (as modelled, *e.g.*, by Cerreia-Vioglio *et al.*, 2015). However, our results put substantial limitations on the type of model that may be alternative explanations of the data. First, given that our belief incentivization is proper for agents who respect first-order stochastic dominance, the model would have to admit violations of first-order stochastic dominance. And second, the model would have to account for the same pattern of inconsistent belief-choices that comes so natural out of the belief-sampling model.

Our third experiment at first sight seems to yield evidence that is more mixed. We do not find the posited relationships with the number of observations if we use a comparison of the information to direct reports of the prior. However, a simulation of the model reveals that the model predicts effects that *are* (very close to) null-effects. Thus, in contrast to our initial idea, the elicitation of priors does not allow for a strong test of the model. In contrast, the aspect of the third experiment that does allow for such a test, the elicitation of belief certainty, does lend further credence to the model, as best-response rates are clearly related to participants' (normalized) belief-certainty reports.

Given the predictive success of the belief-sampling model also on the individual-choice level, the next step might indeed be to generalize the model into a game-theoretic model. The current working paper of Mauersberger (2020) takes a step in this direction. In this context, we fill the empirical gap between the existing

research in the exploration-exploitation literature and the aggregate behaviour at Mauersberger's focus.

References

- Ai, C., & Norton, E. C. (2003). Interaction terms in logit and probit models. *Economics Letters*, 80(1), 123-129.
- Agrawal, S., & Goyal, N. (2012). Analysis of Thompson Sampling for the multi-armed bandit problem. *JMLR Workshop and Conference Proceedings (COLT2012)*, 23:39.1-39.26.
- Alós-Ferrer, C., Granić, D. G., Shi, F., & Wagner, A. K. (2012). Choices and preferences: Evidence from implicit choices and response times. *Journal of Experimental Social Psychology*, 48(6), 1336-1342.
- Alós-Ferrer, C., Granić, D. G., Kern, J., & Wagner, A. K. (2016). Preference reversals: Time and again. *Journal of Risk and Uncertainty*, 52(1), 65-97.
- Becker, G. M., DeGroot, M. H., & Marschak, J. (1963). Stochastic models of choice behavior. *Behavioral Science*, 8(1), 41-55.
- Çelen, B., Geng, S., & Li, H. (2019). Belief error and non-Bayesian social learning: Experimental evidence. *Working paper*.
- Cerreia-Vioglio, S., Dillenberger, D., & Ortoleva, P. (2015). Cautious Expected Utility and the Certainty Effect. *Econometrica*, 83(3), 693-728.
- Chabris, C. F., Morris, C. L., Taubinsky, D., Laibson, D., & Schuldt, J. P. (2009). The Allocation of Time in Decision-Making. *Journal of the European Economic Association*, 7(2/3), 628-637.
- Costa-Gomes, M. A., & Weizsäcker, G. (2008). Stated beliefs and play in normal-form games. *The Review of Economic Studies*, 75(3), 729-762.
- Danz, D. N., Fehr, D., & Kübler, D. (2012). Information and beliefs in a repeated normal-form game. *Experimental Economics*, 15(4), 622-640.
- Dickhaut, J., Smith, V., Xin, B., & Rustichini, A. (2013). Human economic choice as costly information processing. *Journal of Economic Behavior & Organization*, 94, 206-221.
- Drerup, T., Enke, B., & von Gaudecker, H.-M. (2017). The precision of subjective data and the explanatory power of economic models. *Journal of Econometrics*, 200(2), 378-389.
- Ether, J., Jeleva, M., & Tallon, J. M. (2012). Decision theory under ambiguity. *Journal of Economic Surveys*, 26(2), 234-270.
- Fischbacher, U. (2007). z-Tree: Zurich toolbox for ready-made economic experiments. *Experimental Economics*, 10(2), 171-178.
- Folli, D. & Wolff, I. (2022). Biases in Belief Reports. *Journal of Economic Psychology*, 88, 102458.
- Francetich, A., & Kreps, D.M. (2018). Choosing a good toolkit: Bayes-rule based heuristics. *Working paper*.
- Gershman, S.J. (2018). Deconstructing the human algorithms for exploration. *Cognition*, 173: 34-42.
- Gilboa, I., & Schmeidler, D. (1989). Maxmin expected utility with non-unique prior. *Journal of Mathematical Economics*, 18(2), 141-153.
- Gilboa, I., Postlewaite, A. W., & Schmeidler, D. (2008). Probability and uncertainty in economic modeling. *The Journal of Economic Perspectives*, 22(3), 173-188.

- Greiner, B. (2015). Subject pool recruitment procedures: organizing experiments with ORSEE. *Journal of the Economic Science Association*, 1(1), 114-125.
- Groeneveld, R. A., & Meeden, G. (1977). The mode, median, and mean inequality. *The American Statistician*, 31(3), 120-121.
- Harless, D. W., & Camerer, C. F. (1994). The predictive utility of generalized expected utility theories. *Econometrica*, 62(6), 1251-1289.
- Hossain, T., & Okui, R. (2013). The binarized scoring rule. *The Review of Economic Studies*, 80(3), 984-1001.
- Hyndman, K., Ozbay, E. Y., Schotter, A., & Ehrblatt, W. Z. (2012). Convergence: An experimental study of teaching and learning in repeated games. *Journal of the European Economic Association*, 10(3): 573-604.
- Ivanov, A. (2011). Attitudes to ambiguity in one-shot normal-form games: An experimental study. *Games and Economic Behavior*, 71(2): 366-394.
- Klibanoff, P., Marinacci, M., & Mukerji, S. (2005). A smooth model of decision making under ambiguity. *Econometrica*, 73(6): 1849-1892.
- Konovalov, A. & Krajbich, I. (2017). "Revealed Indifference: Using Response Times to Infer Preferences." Working Paper.
- Loomes, G., & Sugden, R. (1995). Incorporating a stochastic element into decision theories. *European Economic Review*, 39(3), 641-648.
- Luce, R.D. (1959). *Individual Choice Behavior: A Theoretical Analysis*. New York: Wiley.
- Manski, C. F., & Neri, C. (2013) First- and second-order subjective expectations in strategic decision-making: Experimental evidence. *Games and Economic Behavior*, 81, 232-254.
- May, B.C., Korda, N., Lee, A., & Leslie, D.S. (2012). Optimistic bayesian sampling in contextual-bandit problems. *Journal of Machine Learning Research*, 13: 2069-2106.
- Mauersberger, F. (2020). Thompson Sampling: Predicting Behavior in Games and Markets. *Working Paper*.
- Mauersberger, F. (2021). Monetary policy rules in a non-rational world: A macroeconomic experiment. *Journal of Economic Theory*, 197: 105203.
- McKelvey, R. D., & Page, T. (1990). Public and private information: An experimental study of information pooling. *Econometrica*, 58, 1321-1339.
- McKelvey, R. D., & Palfrey, T. R. (1995). Quantal Response Equilibria for Normal Form Games. *Games and Economic Behavior*, 10(1), 6-38.
- Moffatt, P. G. (2005). Stochastic choice and the allocation of cognitive effort. *Experimental Economics*, 8(4), 369-388.
- Mosteller, F., & Nogee, P. (1951). An experimental measurement of utility. *Journal of Political Economy*, 59(5), 371-404.
- Norton, E. C., Wang, H., & Ai, C. (2004). Computing interaction effects and standard errors in logit and probit models. *Stata Journal*, 4, 154-167.
- Nyarko, Y., & Schotter, A. (2002). An experimental study of belief learning using elicited beliefs. *Econometrica*, 70(3), 971-1005.
- Osborne, M.J., & Rubinstein, A. (1998). Games with Procedurally Rational Players. *American Economic Review*, 88(4), 834-847.
- Pouget, A., Drugowitsch, J., & Kepecs, A. (2016). Confidence and certainty: distinct probabilistic quantities for different goals. *Nature Neuroscience*, 19(3), 366-374.
- Ratcliff, R. (1978). A theory of memory retrieval. *Psychological Review*, 85(2), 59.

- Rey-Biel, P. (2009) Equilibrium play and best response to (stated) beliefs in normal form games, *Games and Economic Behavior*, 65(2), 572-585.
- Savage, L. J. (1954) *The Foundations of Statistics*. New York: John Wiley and Sons. (Second ed., Dover, 1972).
- Schmeidler, D. (1989). Subjective probability and expected utility without additivity. *Econometrica*, 57(3), 571-587.
- Schulz, E., Konstantinidis, E., & Speekenbrink, M. (2015). Learning and decisions in contextual multi-armed bandit tasks. *Proceedings of the 37th Annual Conference of the Cognitive Science Society*, 2122-2127.
- Selten, R., & Chmura, T. (2008). Stationary Concepts for Experimental 2x2-Games. *American Economic Review*, 98(3), 938-966.
- Weber, Roberto A. (2003). 'Learning' with no feedback in a competitive guessing game. *Games and Economic Behavior*, 44(1): 134-144.
- Wolff, I., & Bauer, D. (2018). Elusive beliefs: Why uncertainty leads to stochastic choice and errors. *TWI Research Paper Series*, No. 111.

Appendix (for online publication only)

A Relating the observable belief-action consistency to the error rate

To obtain our predictions, we need to put some structure on the belief distribution $q(\phi)$. We assume ϕ to be beta-distributed, $q(\phi; \alpha, \beta)$, a very flexible distribution that is able to approximate many different belief distributions.³⁰

PROPOSITION 1: If $q(\phi; \alpha, \beta)$ with $\mu_q \neq 0.5$ is the Beta-distribution with hyperparameters $\alpha, \beta > 1$, the expected observed best-response rate $\widehat{BR}$ decreases in the error rate ε in a symmetric game.

PROOF: $\frac{\partial \widehat{BR}}{\partial \varepsilon} = 4\varepsilon - 2$. Hence, $\widehat{BR}$ decreases in ε if $\varepsilon < 0.5$. The error rate ε is always smaller than 0.5 if the median m_q of the belief distribution $q(\phi; \alpha, \beta)$ is on the same side of the critical value as the mean μ_q (that is, if the median favors the same best response as the mean $BR[m_q] = BR[\mu_q]$) because then, more than 50% of the probability mass are contained in $(1 - \varepsilon)$.

For the symmetric games we consider here, it is hence sufficient to show that either **both** or **neither** the mean and median of $q(\phi)$ are larger than $\phi^{crit} = 0.5$. By the mode-median-mean inequality (Groeneveld & Meeden, 1977), $\mu_q \leq m_q$ if $1 < \beta < \alpha$. However, if $\beta < \alpha$, also $\mu_q = \frac{\alpha}{\alpha + \beta} > 0.5$. Hence, if $1 < \beta < \alpha$, then $0.5 < \mu_q \leq m_q$, and if $1 < \alpha < \beta$, then $m_q \leq \mu_q < 0.5$. ■

Note that **PROPOSITION 1** also holds if either the action or the belief are assumed to be non-stochastic. In these cases, the expected observed best response rate is simply $\widehat{BR}' = (1 - \varepsilon)$ and obviously $\frac{\partial \widehat{BR}'}{\partial \varepsilon} < 0$.

B Simulating the four-option game

To make clear that our predictions for the four-option game do indeed result from our theory, we run a simulation. First, we randomly choose an absolute weight n_q for our prior, with $n_q \sim U[1, 400]$. n_q can be interpreted as the number of

³⁰The beta-distribution is a prominent example of a probability density function with support $(0,1)$ and hence suitable to model a distribution over probabilities. With this distributional assumption, it will be convenient to apply Bayesian-updating, as the beta-distribution is a conjugate prior for the Bernoulli and binomial distributions. Hence, updating a (beta-distributed) prior belief by a number of R - and L -choices in a sample (n i.i.d. Bernoulli variables) will again yield a beta-distributed posterior. See section 3.3. For the four-option setup, we have a three-dimensional space. We thus use the beta-distribution's generalization: the Dirichlet-distribution.

observations in a prior sample. We choose an upper limit of 400 so that there can be priors that outweigh the maximum sample size of 360 used in our experiment. Our prior should be Dirichlet- (α) distributed (the Dirichlet-distribution is the multivariate generalization of the Beta-distribution, and the conjugate prior for the multinomial distribution). So, we randomly draw four probabilities $\pi_i^{(q)}$ for the α s of the prior distribution. We use a Dirichlet-(1, 1, 1, 1) distribution for this random draw. Then, we use the randomly-drawn probabilities together with the drawn n_q , to determine the parameters of the prior Dirichlet distribution: $\alpha_i^{(q)} = n_q \pi_i^{(q)} + 1$.

After randomly defining the prior, we create an “observed sample” of choices. We draw the number of new observations n from a uniform distribution over all levels we use in the experiment but the extreme cases, so that $n \sim U\{9, 12, 15, 18, 36, 64, 92, 120, 148, 176, 204, 232, 260, 288, 316, 345, 348, 351, 354\}$. Then, we randomly determine ‘choice probabilities’ for the random samples. For this purpose, we draw three values $\pi_i^{(aux)}, \pi_i^{(aux)} \sim U[0, 1]$. We then let sampling probabilities be a random perturbation of the following sequence of probabilities: $\pi_1^{(s)} = \pi_1^{(aux)}$, $\pi_2^{(s)} = (1 - \pi_1^{(s)})\pi_2^{(aux)}$, $\pi_3^{(s)} = (1 - \pi_1^{(s)} - \pi_2^{(s)})\pi_3^{(aux)}$, and $\pi_4^{(s)} = (1 - \pi_1^{(s)} - \pi_2^{(s)} - \pi_3^{(s)})$. Using the random perturbation of our probabilities $\pi_i^{(s)}$, we draw a sample of n new ‘observations’. We then apply Bayesian updating to update the prior Dirichlet distribution according to the ‘new observations’, so that $\alpha_i^{(p)} = \alpha_i^{(q)} + n_i$.

So far, we have simulated a prior belief-distribution with absolute weight n_q and an observed sample of n choices. Using Bayes’ rule, we have updated the prior to arrive at a posterior belief distribution. To assess the predicted observed best-response rate for the resulting posterior, we use 10’000 iterations of the following process: from the posterior, we draw a belief ϕ_A^r for the action and a belief ϕ_B^r for the reported belief, with $\phi_A^r, \phi_B^r \sim Dir(\alpha^{(p)})$. If the two beliefs have their minimum on the same option, they are consistent. For each draw of ϕ_B^r , we also record whether it has the same minimum as the distribution of ‘new observations’ $\mathbf{n}$. Then, we record the average consistency for all draws of ϕ_B^r that have the minimum on the ‘anti-mode’ of $\mathbf{n}$. Further, we compute the average consistency for all draws of ϕ_B^r that do not have the minimum on the ‘anti-mode’ of $\mathbf{n}$ (where we define the anti-mode to be the location that occurs least often in the sample). We thus compute best-response rates separately for when the reported belief indicates the same best-response as the observed sample and when it has not.

We iterate the above process 5’000 times. Then, we use a linear regression to relate the level of consistency to the sample size n , a dummy indicating whether the drawn belief ϕ_B^r has its minimum on the anti-mode of the sample $\mathbf{n}$, and the interaction of both terms. We plot the resulting predicted best-response rates

in Figure 2 in Section 4. This prediction has three characteristics: when the reported belief and the sample suggest the same choice, (i) the best-response rate is higher than when they do not; (ii) the predicted best-response rate increases in n ; and (iii) when the reported belief and the sample suggest different choices, the predicted best-response rate decreases in n .

C Figures and Tables

Game	Box 1	Box 2	Box 3	Box 4	χ^2	Sig. on 5%	Sig. on 1%
1	74	106	110	70	14.578	✓	✓
2	110	68	76	106	14.844	✓	✓
3	84	70	86	120	15.022	✓	✓
4	110	100	70	80	11.111	✓	-
5	104	84	101	71	7.933	✓	-
6	76	77	97	110	9.044	✓	-
7	115	63	84	98	16.156	✓	✓
8	83	90	87	100	1.7556	-	-
9	123	74	75	88	17.489	✓	✓
10	104	83	92	81	3.667	-	-
11	97	77	81	105	5.822	-	-
12	101	82	88	89	2.111	-	-
13	86	76	80	118	12.178	✓	✓
14	116	92	72	80	12.267	✓	✓
15	76	104	89	91	4.378	-	-
16	91	66	102	101	9.356	✓	-
17	113	70	90	87	10.422	✓	-
18	85	95	61	119	19.244	✓	✓
19	100	76	71	113	13.178	✓	✓
20	93	87	75	105	5.200	-	-
21	97	84	86	93	1.222	-	-
22	92	71	93	104	6.333	-	-
23	102	75	101	82	6.156	-	-
24	104	67	76	113	16.111	✓	✓
Number of significantly non-uniform distributions:						15	10

Table C1: The 24 historic choice distributions, used to sample the provided information. Corresponding χ^2 -tests with H_0 : choices are uniform across boxes

1 2 3 4	◇ ◇ ◇ ◇
1 x 3 4	y 4 X 3
A B A A	• △ • •
B A B B	△ • ◇ 0
◇ ◇ △ ◇	2 0 i 5
◇ ◇ △ ◇	   
q 8 -)	♥ ♦ ♠ †
a a a B	   
> > < <	B A A A
(o : >	As 2 3 Joker
y • • •	A A B A
△ • 0 △	A A A B

Figure C1: The 24 label sets, used to label the four options of the games. One set for each game.

Best-response to belief	Linear Probability Model		
	Model 1	Model 2	Model 3
$n_{normalized}$		-0.158** (0.068)	-0.160** (0.066)
$Belief-min = Info-min$		0.117* (0.063)	0.117* (0.061)
$n_{normalized} \times (Belief-min = Info-min)$		0.212** (0.102)	0.217** (0.100)
'Strength' of the reported belief	0.511** (0.223)		0.534** (0.222)
Period	0.009*** (0.002)	0.008*** (0.002)	0.008*** (0.002)

Table C2: Linear Probability Model OLS regressions of observed best-responses. Standard errors in parentheses are clustered on the participant level (54 clusters). Asterisks: *** $p < 0.01$, ** $p < 0.05$, * $p < 0.1$. Additional controls in all models: *age*, *math-grade*, *economics-student*, *male*, and a self-reported *reliability-of-answers* measure.

Best-response to reported belief	Average marginal effects					
	Model 1		Model 2		Model 3	
	<i>Coeff.</i>	<i>p</i> (s.e.)	<i>Coeff.</i>	<i>p</i> (s.e.)	<i>Coeff.</i>	<i>p</i> (s.e.)
Observations = 898, Clusters = 54						
$n_{normalized}$			-0.131	0.009 (0.050)	-0.130	0.009 (0.050)
$Belief-min = Info-min$			0.109	0.055 (0.057)	0.108	0.050 (0.055)
$n_{normalized} \times Belief-min = Info-min$			0.206	0.030 (0.095)	0.205	0.030 (0.094)
‘Strength’ of the reported belief	0.551	0.084 (0.319)			0.523	0.088 (0.307)
Normalized variance of the reported belief	-0.029	0.916 (0.271)	-0.453	0.100 (0.275)	-0.173	0.500 (0.256)
Period	0.009	0.000 (0.002)	0.008	0.000 (0.002)	0.008	0.000 (0.002)
Mean Squared Error (Full Sample)	0.1960		0.1852		0.1822	

Table C3: Average marginal effects of logit regressions on observed best-responses including a measure of the reported belief’s variance normalized to the interval $[0, 1]$, where 0 represents a belief that puts 100% probability mass on a single action and 1 stands for a uniform belief. Standard errors in parentheses are clustered on the participant level (54 clusters). The interaction is computed using the `inteff` software by Norton, Wang & Ai (2004). See also Ai & Norton (2003). The marginal effect of the interaction is positive for all participants. Additional controls in all models: *age*, *math-grade*, *economics-student*, *male*, and a self-reported *reliability-of-answers* measure.

Best Response to belief	Average Marginal effects after Logit			
	Model 1'	Model 1''	Model 2'	Model 3'
$\ln(\text{decision time})$	-0.123*** (0.033)	-0.114*** (0.034)	-0.099*** (0.033)	-0.090*** (0.034)
$n_{\text{normalized}}$			-0.124** (0.051)	-0.128** (0.050)
$\text{Belief-min} = \text{Info-min}$			0.095* (0.054)	0.094* (0.053)
$n_{\text{normalized}} \times (\text{Belief-min} = \text{Info-min})$			0.211** (0.103)	0.217** (0.102)
'Strength' of the reported belief		0.490 (0.300)		0.516* (0.299)
Period	0.006*** (0.002)	0.007*** (0.002)	0.006*** (0.002)	0.006*** (0.002)
Mean Squared Error	0.1944	0.1914	0.1834	0.1799

Table C4: Marginal effects of Logit regressions accounting for $\ln(\text{decision time})$. Number of Observations = 898. Standard errors in parentheses are clustered on the participant level (54 clusters). Asterisks: *** $p < 0.01$, ** $p < 0.05$, * $p < 0.1$. Additional controls in all models: *age*, *math-grade*, *economics-student*, *male*, and a self-reported *reliability-of-answers* measure.

Best-response to reported belief	Average marginal effects					
	Model 1		Model 2		Model 3	
Observations = 1226, Clusters = 79	<i>Coeff.</i>	<i>p</i> (s.e.)	<i>Coeff.</i>	<i>p</i> (s.e.)	<i>Coeff.</i>	<i>p</i> (s.e.)
$n_{\text{normalized}}$			-0.091 (0.054)	0.092 (0.054)	-0.080 (0.053)	0.133 (0.053)
$\text{Belief-min} = \text{Info-min}$			0.149 (0.055)	0.007 (0.055)	0.137 (0.054)	0.011 (0.054)
$n_{\text{normalized}} \times \text{Belief-min} = \text{Info-min}$			0.175 (0.073)	0.017 (0.073)	0.180 (0.073)	0.014 (0.073)
'Strength' of the reported belief	1.180	0.002 (0.379)			1.023	0.006 (0.371)
Period	0.008	0.001 (0.002)	0.007	0.002 (0.002)	0.007	0.001 (0.002)
Mean Squared Error (Full Sample)	0.2255		0.2132		0.2103	

Table C5: Average marginal effects of logit regressions on observed best-responses for the data set including priors and belief-certainty reports. Standard errors in parentheses are clustered on the participant level (79 clusters). The interaction is computed using the `inteff` software by Norton, Wang & Ai (2004). See also Ai & Norton (2003). Additional controls in all models: *age*, *math-grade*, *economics-student*, *male*, and a self-reported *reliability-of-answers* measure.

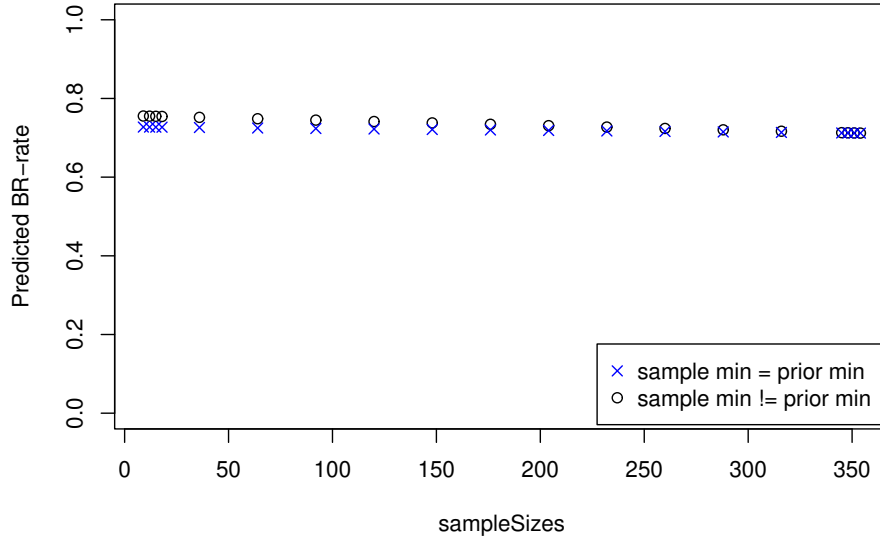


Figure C2: Predicted best-response rates in the four-option game according to our simulation, conditional on *Info-Min = Prior-Min*.

Best-response to belief rep.	Baseline analysis						With belief certainty					
Obs. in parentheses:	Model 1 (1611)		Model 2 (1611)		Model 3 (1611)		Model 4 (1537)		Model 5 (1537)		Model 6 (1537)	
Clusters = 110	Coeff.	p (s.e.)	Coeff.	p (s.e.)	Coeff.	p (s.e.)	Coeff.	p (s.e.)	Coeff.	p (s.e.)	Coeff.	p (s.e.)
$n_{normalized}$			0.007	0.834 (0.035)	0.020	0.567 (0.035)			-0.011	0.775 (0.038)	0.001	0.985 (0.039)
<i>Info-Min = Prior-Min</i>			-0.007	0.873 (0.045)	-0.009	0.842 (0.044)			-0.004	0.935 (0.046)	-0.008	0.857 (0.044)
$Info-Min = Prior-Min \times n_{normalized}$			0.026	0.682 (0.064)	0.025	0.695 (0.063)			0.008	0.899 (0.066)	0.010	0.881 (0.064)
'Strength' of the ...reported belief	0.875	0.016 (0.363)			0.890	0.015 (0.365)	0.949	0.011 (0.375)			0.952	0.012 (0.377)
Normalized belief certainty							0.033	0.008 (0.012)	0.032	0.019 (0.014)	0.033	0.016 (0.014)
Period	0.005	0.011 (0.002)	0.005	0.019 (0.002)	0.005	0.013 (0.002)	0.007	0.003 (0.002)	0.006	0.005 (0.002)	0.007	0.003 (0.002)
Mean squared error	0.2371		0.2395		0.2370		0.2343		0.2371		0.2343	

Table C6: Average marginal effects of logit regressions of observed best-responses for the full sample (before exclusions). Standard errors in parentheses are clustered on the participant level (110 clusters). The interaction is computed using the `inteff` software by Norton, Wang & Ai (2004). See also Ai & Norton (2003). Additional controls in all models: *age*, *math-grade*, *economics-student*, *male*, and a self-reported *reliability-of-answers* measure.

D Experimental Instructions

The instructions are translated from german. Boxes indicate consecutive screens showed to participants.

Today's Experiment

Today's experiment consists of 24 rounds in which you will make two decisions each.

Decision 1 and Decision 2

In the first round, you will see the instructions for both decisions directly before the decision. In later rounds, you can display the instructions again if you need to.

The payment of the experiment

In every decision you can earn points. At the end of the experiment, 2 rounds are randomly drawn and payed. In one of the rounds, we pay the point you earned from decision 1 and in the other round, you earn the points from decision 2. The total amount of points you earned will be converted to EURO with the following exchange rate:

1 Point = 1 Euro

After the experiment is completed, there will be a short questionnaire. For completion of the questionnaire, you additionally receive 5 Euro. You will receive your payment at the end of the experiment in cash and privacy. No other participant will know how much money you earned.

General Instructions

For today's experiment, another experiment plays a central role. This experiment has been conducted earlier, here in the LakeLab. The earlier experiment is described in the following.

The earlier experiment

In the earlier experiment, 360 participants ran through 24 rounds. In every round groups of two randomly matched persons were formed. The group members did not know each others' identity and could not communicate throughout the whole experiment.

One round of the experiment worked in the following way: both participants did see the exact same screen. On the screen, there was an arrangement of four boxes which are marked with symbols. Both of the group members chose one of the boxes. If both group members chose different boxes, both received a price. If both members chose the same box, there was no payoff. All participants learn about which box was chosen by the other participant and which payoff they received in a certain round only at the end of the experiment.

The arrangement of symbols on the boxes differed in every round for every group. The decision of a participant was hence on an unknown arrangement. Below, you can see an example of how such an arrangement could have looked like.

Example: The four boxes are marked from left to right by Diamond, Heart, Spade, Diamond.



In this example, there are two boxes which are marked with the same symbol. However, the boxes on the most left and most right count as different boxes.

Instructions for experiment 1

The number of points you receive in decision 1 depends on your own decision, as well as on a participant of the earlier experiment who will be randomly matched with you. How this works, will be explained in the following.

Decision 1

For decision 1 in every round, you see an arrangement of four boxes which are marked with symbols that was also used in the earlier experiment. The computer then randomly draws one of the participants of the earlier experiment who chose one of the boxes.

In decision 1, you have to choose a box as well.

If you choose another box than your randomly matched partner from the earlier experiment, you receive 7 points. If you choose the same box as your randomly matched partner, you don't receive points.

On the next screen, you receive more information about the earlier experiment.

Additional Information

In every round, before you make decision 1, you receive additional information, how a certain sample of the 360 participants of the earlier experiment decided in the respective arrangement. In every round, a random sample is drawn from all 360 participants of the earlier experiment. For every of the four boxes, you get to know how many participants in the sample chose that box. You can see an example of how this information looks like below:

[Example Screen, see screenshot below]

Please note, that the participant you are matched to in the respective period is not contained in the sample you see. This means, that this participant is always drawn from the remaining participants which are not shown to you.

The size of the respective sample of participants you receive information about will vary from round to round. This means, that you have different amounts of information about the decisions of the participants of the earlier experiment in every round.

Please note, that the participants of the earlier experiment did not have any information how other participants decided. Information like you can see it above, was not displayed to the participants of the earlier experiment.

The information is displayed on the next screen.

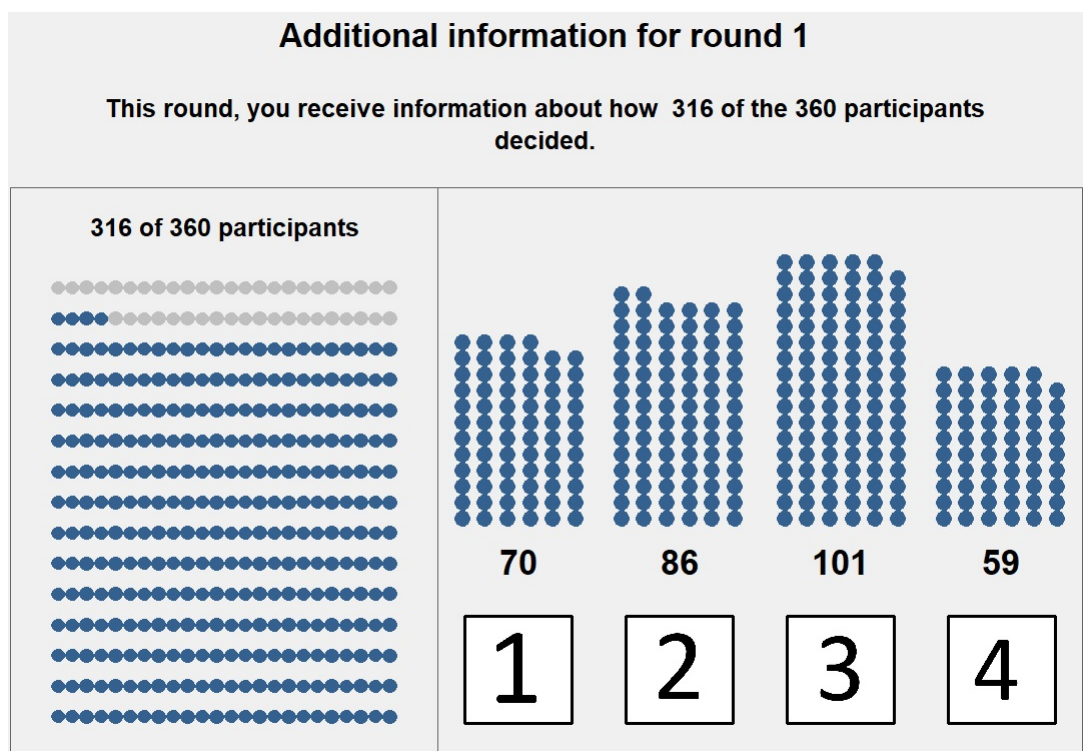


Figure D1: Example screen of the presentation of choice frequencies.

Instructions for decision 2

In decision 2, your payoff also depends on your own decision and on the decision of your matching partner from the earlier experiment. We now explain decision 2 in detail.

Decision 2

Decision 2 refers always to the arrangement from decision 1, which was also used in the earlier experiment. You will hence see the arrangement of boxes from the respective round again. You also can look at the additional information again. Again, the decision of your matching partner from the earlier experiment is relevant for you.

Decision 2 is about your assessment, how your matching partner from the earlier experiment decided. We are interested in your assessment of the following question:

“With what probability did your matching partner chose each of the respective boxes of the current set-up?”

For every box, you can report your assessment with what probability your matching partner chose the respective box. You can enter the percentage numbers in a bar diagram. By clicking into the diagram, you can adjust the height of the bars. You can adjust as many times as you like, until you confirm. Since your assessments are percentage numbers, the bars have to add up to 100%. The sum of your assessment is displayed on the right. You can adjust this value to 100% by clicking. Or you enter the relative sizes of your assessments only roughly and then press the “scale” button. Please note, that because of rounding, the displayed sum may deviate from 100% in some cases.

On the next page, we explain the payoff of decision 2.

The payoff in decision 2

In this decision, you can either earn 0 or 7 points. Your chance of earning 7 points increases with the precision of your assessment. Your assessment is more precise, the more it is in line with the decision behavior of your matching partner. For example, if you reported a high assessment on the actually selected box, your chance increases. If your assessment on the selected box was low, your chance decreases.

You may now look at a detailed explanation of the computation of your payment, which rewards the precision of your assessment.

It is important for you to know, that the chance of receiving a high payoff is maximal in expectation, if you assess the behavior of your matching partner correctly. It is our intention, that you have an incentive to think carefully about the behavior of your matching partner. We want, that you are rewarded if you have assessed the behavior well and made a respective report.

At the end of the experiment, one participant of today's experiment will roll a number between 1 and 100 with dies. If the rolled number is smaller or equal to your chance, you receive 7 points. If the number is larger than your chance, you receive 0 points.

As soon as you reported and confirmed your assessment about the behavior of your matching partner, the round ends. You will then be matched with another participants and the next round begins.

Payment of the assessments

At the end of your assessment, you will receive the 7 points with a certain chance (p) and with $(1 - p)$, you receive 3 points. You can influence your chance p with your assessment in the following way:

As described above, you will report an assessment for each box, on how likely your matching partner is to select that box. One of boxes is the actually selected. At the end, your assessments are compared to the actual decision of your matching partner. Your deviation is computed in percent.

Your chance p is initially set to 1 (hence 100%). However, there will be deductions, if your assessments are wrong. The deductions in percent are first squared and then divided by two.

For example, if you place 50% on a specific box, but [your matching partner selects another box,] your deviation is equal to 50%. Hence, we deduct $0.50 * 0.50 * \frac{1}{2} = 0.125$ (12.5%) from p .

[For the box, which is actually selected by your matching partner, it is bad if your assessment is far away from 100%. Again, your deviation from that is squared, halved and deducted. For example if you only place 60% probability on the actually selected box, we will deduct $0.40 * 0.40 * \frac{1}{2} = 0.08$ (8%) from p .]

With this procedure, we compute your deviations and deductions for all boxes.

At the end, all deductions are summed up and the smaller the sum of squared deviations is, the better was your assessment. For those who are interested, we show the mathematical formula according to which we compute the chance.

$$p = 1 - \frac{1}{2} \left[\sum_i (q_{box_i, estimate} - q_{box_i, true})^2 \right]$$

The value of p of your assessment will be computed and displayed to you at the end of the experiment. The higher p is, the better your assessment was and the higher your chance to receive 7 points (instead of 0) in this part. At the end of the experiment, the computer will roll a random number between 0 and 100 with dies. If this number is smaller or equal to p , you receive 7 points. If the number is larger than p you receive 0 points.

Summary

In order to have a high chance to receive the large payment, it is your aim to achieve as few deductions from p as possible. This works best, if you have an good assessment of the behavior of your matching partner and report that assessment truthfully.



THURGAU INSTITUTE
OF ECONOMICS
at the University of Konstanz

Hafenstrasse 6
CH-8280 Kreuzlingen

T +41 (0)71 677 05 10
F +41 (0)71 677 05 11

info@twi-kreuzlingen.ch
www.twi-kreuzlingen.ch