

Using Advice Without Considering Its Quality: A Laboratory Experiment of Demand for Advice.*

Jacopo Bregolin,[†]Astrid Hopfensitz,[‡]Elena Panova.[§]

Abstract

We experimentally test how the content of advice, namely, its alignment with common priors, influences beliefs about its quality and future demand for it. We reject the theoretical hypothesis that demand for advice can be increased by giving advice in alignment with common priors. We find, furthermore, that such alignment has hardly any impact on the participants' beliefs about quality of advice. Nevertheless, advice influences participants' guesses in an incentivized task, regardless of their beliefs about the quality of advice itself.

Key words: demand for information, belief updating.

JEL codes: D90, C91, D83.

*We thank Stéphane Cezera for his help in organizing the experimental sessions. We are also very grateful to Juan Carrillo and Leeat Yariv for their valuable comments and suggestions. In addition, we thank the participants of the 2025 EEA meetings in Bordeaux as well as the participants of the University of Liverpool Applied Micro internal seminar for their attention and feedback. Elena Panova acknowledges funding from the ANR under grant ANR-17-EURE-0010 (Investissements d'Avenir program).

[†]University of Liverpool, Management School, Department of Economics. E-mail: jacopo.bregolin@liverpool.ac.uk

[‡]emlyon business school, GATE, Lyon, France. E-mail: hopfensitz@em-lyon.com

[§]Toulouse School of Economics, University of Toulouse Capitole. E-mail: e_panova@yahoo.com

1 Introduction

In many situations, private decision makers buy advice from experts. Examples include purchasing financial advice from intermediaries, paying physicians for medical guidance, or buying news from the media. Consumers of professional advice may remain uncertain about its quality. The high volatility of investment returns makes it difficult to evaluate the effectiveness of an investment strategy. Similarly, the lack of counterfactual knowledge makes it hard to assess whether a prescribed treatment was optimal, or which political candidate or policy would have been best. Therefore, experts may bias their advice in their own interest.

An influential theoretical literature on reputational cheap-talk suggests that experts will confirm the common priors regardless of their true opinions to signal their competence.¹ The demand side of this theory rests on two assumptions. First, consumers update their beliefs about an advisor's competence based on the content of the received advice using Bayes' rule. Second, their demand for advice increases with their posterior beliefs regarding the advisor's competence. However, there is substantial evidence of errors in belief updating (for surveys see: [Camerer 1998](#) and [Benjamin 2019](#)). Furthermore, the empirical and experimental literature on demand for information (discussed in more detail below) has documented numerous deviations from theoretical predictions regarding the demand for information (see [Ambuehl](#)

¹The empirical literature has also found that *advisors* bias their advice following demand pressures. For instance, [Camara and Dupuis \(2021\)](#) finds that movie reviewers bias their review towards the common prior, [Alpert et al. \(2023\)](#) show that direct-to-consumer advertising increases doctors' prescriptions of the advertised drugs, and [Mastrorocco et al. \(2025\)](#) find that weather news coverage is biased towards the political leaning of the local market.

and Li 2018, Golman et al. 2022, Reshidi et al. 2021 and references therein).

In this paper, we present results from a laboratory experiment designed to test whether the alignment of given advice with common priors indeed increases the demand for future advice. The setup of our experiment is the following: participants guess the color of a ball “drawn” by their computer from a jar containing ten balls, most of which are green and the remainder are blue. Correct guesses are rewarded. The participants have computerized advisors that generate advice of uncertain quality. Specifically, an advisor is equally likely to be *perfect* or *defective*. A perfect advisor always provides correct advice, while a defective advisor provides advice that is equally likely to be correct or false. Participants make guesses over two periods. In the first period, advice is provided for free. In the second period, participants decide whether or not to purchase advice before guessing. These two periods together constitute a round. A participant’s advisor(s) remain the same throughout a round. Participants are observed across fifteen rounds in total, each time with a new set of advisors.

We consider six different experimental treatments (N=307). In the four baseline treatments (N=208), we vary *the precision of the common prior*, which is determined by the number of green balls in the jar: either 6 out of 10 (*weak*) or 7 out of 10 (*strong*); and the *number of available pieces of advice*: either one or two.² In these treatments, we do not elicit beliefs about the quality of the advice, instead asking about these beliefs only in a post-experimental survey. Additionally (N=99), we conducted one treatment for

²This is inspired by the debate on the impact of competition on the quality of professional advice: in finance (see references in Jin 2024), in healthcare (Sivey and Chen, 2019; Currie et al., 2025), and in the media (Gentzkow and Shapiro, 2008).

each prior (i.e., weak and strong) with a single piece of advice, in which participants were asked to report their beliefs about the advisor’s quality after receiving advice in the first period.

We find that, across all treatments, the demand for advice in period two *does not* depend on whether the first-period advice is confirmatory (green) or contrarian (blue). Two prominent behavioral demand patterns emerge: First, we observe that many participants (between 22% and 34%) either purchase all available advice or abstain from purchasing altogether (we call this in the following an “*all-or-nothing*” demand pattern). Second, the demand of the remaining participants tends to match that observed in the previous round, conditional on their guess in that round being correct, reminiscent of the *hot-hand fallacy* (see [Gilovich et al. 1985](#), [Miller and Sanjurjo 2018](#)). Recall that the quality of advisors across rounds is independent; therefore, following such a “win-stay, lose-shift” rule cannot improve guessing accuracy.

In the same vein, we find that 1) the advice received in period one does not affect participants’ willingness to pay for advice, elicited through the [Becker et al. \(1964\)](#) BDM mechanism, and 2) the rate at which advice is followed in period two, conditional on purchase, does not depend on the content of the advice received in the previous period.

In sum, we observe a lack of elasticity with respect to the impact of the content of the received advice on demand, willingness to pay, and guessing. These patterns align with participants’ posterior beliefs about advisor quality, contingent on the received advice. The beliefs elicited during the experiment, as well as those reported in the post-experimental survey, show little evidence that an advice confirming the common priors is seen as a signal

of quality. About half of the participants do not update their beliefs about the quality of the advisor at all.³

While the above purchasing strategies are very different from Bayesian predictions, we find that guessing patterns are broadly aligned with them. We find clear evidence of participants using both the priors regarding the color distribution and the received advice when guessing. Indeed, in our treatments with single piece of advice, participants tend to follow advice, either received or bought. They are more likely to follow advice recommending green (likely outcome) than blue (unlikely outcome), and this difference is larger when the likelihood of green is higher (strong prior versus weak prior).

These findings suggest that advice induces participants to update their beliefs regarding the state of the world, while largely foregoing inferences regarding the quality of advice itself.

Related literature. Our study directly relates to [Meloso et al. \(2023\)](#), as we also examine reputational cheap-talk theory through laboratory experiments. However, in [Meloso et al. \(2023\)](#), the advisors (referred to as reporters) are human and bias their reports toward common priors to enhance the perceived quality of their private signals. Our study shifts attention to the receivers of advice. We investigate not only their beliefs about the quality of advice depending on its alignment with common priors, but also the

³While persistence of beliefs about the quality of advice could be attributed to *conservatism bias* ([Edwards, 1982](#)), the perfectly rigid beliefs observed in many participants suggest that they may be failing to recognize the alignment of advice content with common priors as a signal of quality, which is rather consistent with *base-rate neglect* ([Kahneman and Tversky, 1973](#)).

demand for and use of advice.⁴

Thereby, our findings contribute to a sizable economic literature studying deviations from the Bayesian paradigm in the demand for information.⁵ Probably most closely related are recent experimental studies investigating the demand for instrumentally valuable information in neutral, non-strategic settings. [Ambuehl and Li \(2018\)](#) elicit the valuation of viewing an informative signal on a binary state of the world before guessing the state. They find overvaluation of low-quality information and undervaluation of high-quality information (compression effect).⁶ [Augenblick et al. \(2025\)](#) show that the compression effect is robust across different environments. Specifically, in addition to abstract and naturalistic experiments, it is observed for sports betting markets and financial markets. Our findings regarding a disconnection of purchasing decisions from beliefs about the quality of advice are consistent with the compression effect.

Our analysis focuses on how the alignment of advice with common priors influences perceptions of its quality and subsequent demand. In this

⁴In some treatments in [Meloso et al. \(2023\)](#), the quality of advice is evaluated by human participants who are motivated to make accurate assessments based on the alignment of the advice with both the common priors and the true state (color). Their assessments are consistent with Bayesian updating, in that posteriors are ordered in a Bayesian way for some beliefs about the advisors' reporting strategy. Note that this finding is not necessarily inconsistent with ours, as there are several differences in experimental design. First, in our experiment, the receivers of advice are certain about the reporting strategy: they know that computers truthfully report their signals. Second, they form their posteriors based solely on the alignment of advice with common priors, while remaining uncertain about the true state. Finally, our belief elicitation is not incentivized, which may lead beliefs to deviate from rational updating ([Gächter and Renner, 2010](#)).

⁵The literature has also found deviations in Bayesian updating in learning from signals. For instance, [Kapons and Kelly \(2025\)](#) provide field evidence of prior-biased inference, while [Aydogan et al. \(2025\)](#) find experimental evidence of confirmatory and conservatory biases.

⁶They also find a preference for information structures that may remove uncertainty.

we are different from recent work in psychology and behavioral economics studying the impact of instrumental, hedonic, and cognitive motivations on the demand for information (Golman et al., 2022; Kelly and Sharot, 2021). Heterogeneity in information demand may stem from the different weights individuals assign to these three motives (Kelly and Sharot, 2021). In our setting, advice has no emotional or moral dimension. While our results may provide some insights into the motivations underlying the purchase of advice (such as the observation that buyers tend to follow the advice at a very high rate, seemingly relying on distinct heuristics and exhibiting substantial heterogeneity), we do not directly study this question.⁷

Also related is work by Schoar and Sun (2024) who show, using a randomized controlled trial, that participants rate financial advice significantly higher when it aligns with their reported priors regarding the best investment strategy (either passive or active). Note that in their setting, there is room for belief confirmation motives, which have been shown to be as important as accuracy concerns (Chopra et al., 2024).

Our neutral setting without scope for motivated reasoning is reminiscent of Charness et al. (2021) and Nunnari and Montanari (2025). However, their focus differs from ours: they study choices over biased sources of instrumentally valuable information. In their designs and unlike our setting, computerized advisors have known quality and are biased either toward or against the common priors. Charness et al. (2021) find a tendency to choose advisors with a specific bias, while Nunnari and Montanari (2025) find a

⁷We briefly examine purchasing motivations using a post-experimental survey. Tables 6 and 7 in Appendix E show that most participants’ purchasing decisions are motivated by cognitive rather than instrumental reasons.

tendency to choose the least biased advisor.

The rest of the paper is organized as follows: Section 2 describes the underlying model and its predictions. Section 3 outlines the experimental design and procedure. Section 4 presents our findings. Section 5 concludes.

2 Theoretical predictions.

2.1 Underlying model.

We begin by describing the optimal demand for advice as modeled in the reputational cheap-talk literature. However, unlike that literature, we focus exclusively on the demand side by assuming that the supply of advice is *non-strategic*. The supply consists of one or two signals of uncertain quality, which we refer to as advice. In the experiment, these signals are mechanically reported by machines (robots).

Our model of demand for advice is as follows. An individual receives reward R for guessing the hidden state of Nature in two successive periods indexed with $t = 1, 2$. The period specific state x_t is drawn from distribution

$$x_t = \begin{cases} G, & \text{with probability } p \\ B, & \text{with probability } 1 - p, \end{cases} \quad (1)$$

where $p \geq \frac{1}{2}$.⁸ The states in different periods are independent.

Before making a guess, the individual can receive one or two period-specific pieces of advice about the prevailing state, indexed by $i = l, r$. When only one piece of advice is available, the index i takes the value l .⁹ Advice i

⁸In our experiment, the state x_t corresponds to the color of the ball drawn from the jar in period t : either Green ($x_t = G$) or Blue ($x_t = B$).

⁹During the experiment, the advice indexed by l (r) is presented by the robot positioned on the left (right) side of the computer screen.

in period t is denoted by a_t^i . In period 1, advice is provided free of charge. In period 2, advice is available at a price of ε per piece.

The period-invariant quality q^i of advice i is equally likely to be *perfect* ($q^i = 1$) or *defective* ($q^i = 0$).¹⁰ The qualities of different pieces of advice are independent. A perfect piece of advice always matches the prevailing state, whereas a defective one is equally likely to match or mismatch the state. To facilitate the updating of beliefs about the quality of advice based on its content when two pieces of advice are available, we assume that the signal structure is *nested*:

$$a_t^i = \begin{cases} x_t, & \text{if } q^i = 1, \text{ or } q^i = 0 \text{ and } z_t = H; \\ \{G, B\} \setminus \{x_t\}, & \text{otherwise,} \end{cases} \quad (2)$$

where z_t denotes the outcome of a period-specific flip of a fair coin: heads ($z_t = H$) or tails ($z_t = T$). Notice that this signal structure guarantees that different pieces of advice of the same quality (either both perfect or both defective) always agree. Therefore, agreement between different pieces of advice provides no information about their quality. Conversely, disagreement implies that one piece is perfect while the other is defective.

2.2 Calibration

We limit the precision of the common priors to the following interval:

$$0.5 < p < 0.75. \quad (3)$$

The lower bound indicates that the common prior is more informative than defective advice. The upper bound provides a necessary and sufficient con-

¹⁰Because in our experiment the advice is non-strategic, as explained in the first paragraph of this section, we use the term “perfect” instead of “smart,” as commonly used in the literature on reputational cheap talk, and “defective” instead of “dumb.”

dition for “contradictory” advice B to be likely correct. For simplicity in our experimental design, we let p be a multiple of $\frac{1}{10}$ and, consequently, focus on two values: 0.6 and 0.7.¹¹

To minimize the impact of risk- or loss aversion on the demand for advice, we set the price of receiving advice substantially lower than the stakes of guessing.¹² Specifically, we assume that the reward for a correct guess is $R = 500$, while the price of receiving advice is only $\varepsilon = 5$. Thus, purchasing advice is worthwhile if and only if it increases the probability of guessing correctly by at least one chance in one hundred.

2.3 Rational beliefs and behavior.

This section summarizes our theoretical predictions.

Rational beliefs and behavior with one advice.

Prediction Set 1. *Suppose only one piece of advice is available.*

(i) Optimal first guess: *It is optimal to follow the advice in period 1 for either $p \in \{0.6, 0.7\}$.*

(ii) Rational posteriors: *If $p = 0.6$, the Bayesian posterior is 0.55 when period 1 advice is G and 0.44 when it is B . If $p = 0.7$, the posterior is 0.58 when period 1 advice is G and 0.38 when it is B .*

(iii) Demand for advice in period 2: *If $p = 0.6$, it is optimal to buy advice in period 2 regardless of whether period 1 advice was G or B . If $p = 0.7$, it is optimal to buy advice in period 2 only if period 1 advice was*

¹¹In our experiment, the period-specific state corresponds to the color of a ball drawn from an urn containing ten balls of two different colors.

¹²This is relevant to the applications discussed in the first paragraph of the Introduction.

G .

(iv) WTP for advice in period 2: The WTP is 86.36 if period 1 advice is G and 61.11 if it is B . If $p = 0.7$, the WTP is 45.83 if period 1 advice is G and 0 if it is B .

(v) Optimal use of advice in period 2 (conditional on purchasing): If $p = 0.6$, it is optimal to follow advice in period 2 conditional on bying regardless of whether period 1 advice was G or B . If $p = 0.7$, it is optimal to follow advice in period 2 only if period 1 advice was G . Otherwise, it is optimal to guess G .

A technical proof of Prediction Set 1 is presented in Appendix A.1. The intuition behind these predictions is as follows.

Suppose first that the advice in period 1 is G , which is aligned with the common prior (hereafter referred to as *confirmatory* advice). In this case, it is trivially optimal to guess G . Suppose instead that the advice is B (hereafter, *contrarian*). It remains optimal to follow the advice, as in our calibration the probability of its correctness exceeds the precision of the common prior. This provides the reason for prediction (i).

Note that defective advice is confirmatory with probability 0.5, whereas perfect advice is confirmatory with probability $p > 0.5$. Therefore, confirmatory advice is most likely perfect, while contrarian advice is most likely defective. The higher the probability p , the more strongly the advice content influences the posteriors, as reported in prediction (ii).

When $p = 0.7$, the difference in posteriors following confirmatory and contrarian advice is large enough that, if the advice is contrarian, it loses its influence on guessing in the second period, becoming thereby useless. When

$p = 0.6$, the effect is weaker, so the advice remains influential and potentially useful no matter whether it was contrarian or confirmatory in period 1.¹³ This explains prediction (v).¹⁴

The best alternative to buying and following advice in period 2 is to guess G , which matches the state with probability p . Therefore, WTP equals the difference between the probability that advice matches the state and the prior probability p , scaled by R , would it be positive, and to 0 otherwise. The figures are reported in prediction (iv).

We find the demand for advice by comparing the WTP for it with its price. The demand is positive in all cases except when $p = 0.7$ and the period 1 advice is contrarian, as stated in prediction (iii).

Rational beliefs and behavior with two pieces of advice.

Prediction Set 2. *Suppose that two pieces of advice, l and r , are available.*

(i) Optimal first guess: *For either $p \in \{0.6, 0.7\}$, the optimal first guess is B if both pieces of advice are B , and G otherwise.*

(ii) Rational posteriors: *When both pieces of advice are G , each is perfect with probability 0.53 if $p = 0.6$, and 0.55 if $p = 0.7$. When both pieces of advice in period 1 are B , each is perfect with probability 0.35 if $p = 0.6$, and 0.24 if $p = 0.7$. When the two pieces of advice differ, advice G is perfect with probability p , while advice B is perfect with probability $1 - p$.*

(iii) Demand for advice in period 2: *If different pieces of advice agree*

¹³The threshold precision of common priors below which advice remains influential, regardless of its period 1 content, is $p < 0.68$.

¹⁴We use backward induction reasoning and therefore establish prediction (v) before predictions (iv) and (iii)

in period 1, it is optimal to buy both of them in period 2. If they disagree, it is optimal to buy only one piece of advice, namely the one which was confirmatory in period 1.

(iv) WTP for advice in period 2: Table 1 below summarizes the WTP as a function of the content of the first advice. “Best” denotes the WTP for the advice most likely to be perfect, “Additional” denotes the WTP for an extra piece of advice in addition to the best one,¹⁵ and “Both” denotes the WTP for receiving both pieces of advice jointly.

Table 1: Optimal willingness to pay for advice.

period 1 advice	$p = 0.6$			$p = 0.7$		
	Best	Additional	Both	Best	Additional	Both
$a_1^i = G, a_1^{-i} = B$	100	0	100	75	0	75
$a_1^i = a_1^{-i} = G$	82.35	12, 46	91.18	38.16	24, 72	56.58
$a_1^i = a_1^{-i} = B$	65.38	12, 43	73.08	2.27	24, 17	15.91

(v) Optimal use of advice in period 2 (conditional on purchasing):¹⁶ If the decision maker has only the advice with the highest perceived quality, it is optimal to follow this advice. If the decision maker has both pieces of advice, it is optimal to guess G if at least one piece of advice is G and B otherwise. These statements hold for either p and for any content of the first advice.

Prediction Set 2 is proved in Appendix A.2. The intuition for these predictions is as follows.

¹⁵In the experiment, we observe only WTP (subjectively) “Best” and WTP “Both”.

¹⁶We do not analyze the optimal use of advice perceived to be of inferior quality.

In period 1, each piece of advice is equally likely to be perfect. Therefore, when the two pieces of advice agree, it is optimal to follow them. When they disagree, it is optimal to rely on the prior. This explains prediction (i).

When different pieces of advice agree in period 1, they are equally likely to be perfect a posteriori (most likely perfect if the advice is confirmatory, and most likely defective if it is contrarian). If they disagree, the decision maker can interpret this in two ways: either the confirmatory advice is perfect and the contrarian advice is defective, or vice versa. The first explanation is more plausible. This is the reason for prediction (ii).

In period 2, it is optimal to rely on the common priors, if there is no further information for guessing. If the advice with the (weakly)¹⁷ highest perceived quality (hereafter, the *best advice*) is available, it is optimal to follow this advice. When both pieces of advice are available, it is optimal, in our calibration, to follow the same guessing rule as in period 1, as stated in prediction (v).¹⁸

Given the above optimal use of advice, we compute the WTP for the best advice as the difference between the probability that it matches the period 2 state and the prior probability p , scaled by R . An additional piece of advice can influence guessing only when the best advice is contrarian, either by correcting the guess in state G or by inducing an error in state B . Therefore, the WTP for this additional piece of advice is equal to the difference between the probability that it corrects guessing and the probability that it creates an error, scaled by R . Finally, the WTP for both pieces of advice is equal to

¹⁷Note that when different pieces of advice agree in period 1, their perceived qualities are the same.

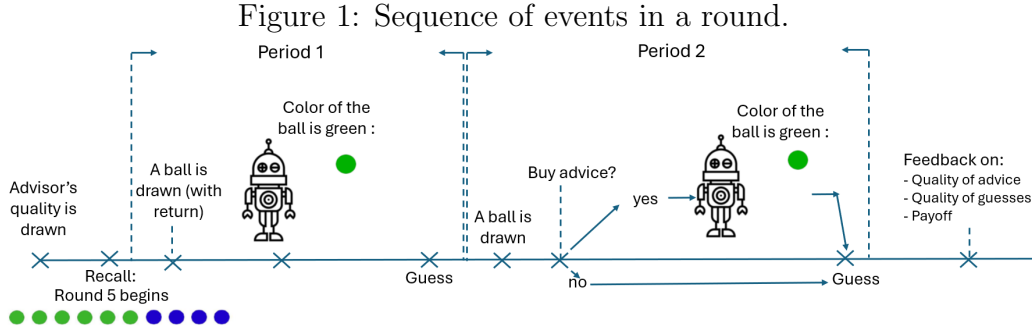
¹⁸Once again, we use backward induction reasoning and therefore establish prediction (v) before predictions (iv) and (iii).

the difference between the probability of making a correct guess using this information and the prior probability p scaled by R . In this way we find the values reported in prediction (iv).

We find the demand for advice by comparing the WTP for it with its price. It is optimal to purchase both pieces of advice if they agree in period 1; otherwise, it is optimal to purchase only the advice that was confirmatory in period 1, as reported in prediction (iii).

3 Experimental design and data.

Overview. Based on the above theoretical framework, we have developed a controlled laboratory task. We will use Figure 1 to describe the details of this task.



Specific features (the number of green balls, the number of robots, the absence of belief elicitation in period 1, and no “bidding” for advice in period 2) indicate that this is a round from 1 to 10 in the weak-prior baseline treatment with one robot.

In the task, the participant (referred to as “she”) was repetitively guessing the color of a ball “drawn” by her computer from a jar containing ten balls, most of which were green and the remainder blue.¹⁹ Each correct guess was

¹⁹Note, once again, that this task has no intrinsic relevance for the participants, which

rewarded with points.

The participants had computerized advisors. An advisor was represented by the drawing of a robot and its advice was visualized by a colored ball (either green or blue), as shown in Figure 1. Each robot was equally likely to be perfect or defective. A perfect robot always provided correct advice, while a defective robot offered advice that was equally likely to be correct or false. All probabilities were objectively specified and explained to participants, and their understanding was checked.

Each trial (hereafter, *round*) consisted of two periods, as illustrated in Figure 1. In period 1, advice was provided free of charge, while in period 2, participants had the option to purchase advice. An advisor’s quality was reassigned at the beginning of each round and remained the same until the end of that round. To help participants remember this, the image of each robot representing an advisor remained the same across the two periods of a given round (as in Figure 1) and changed at the beginning of a new round. Participants received the feedback on their performance and on the quality of their advisors in a round only after submitting their second, and the last, guess in that round (see, once again, Figure 1). Thus, the only information available for their purchasing decision in period 2 was the content of the advice received in period 1.

Overall, there were 15 rounds. In rounds 1 to 10, the price of advice in period two was fixed. In the final five rounds, the participants’ WTP for advice was elicited via a BDM procedure.²⁰

limits the scope for motivated reasoning. Indeed, a growing body of evidence shows that when information relates to personal traits, individuals tend to give more weight to positive news than to negative news (see e.g., [Thaler 2024](#)).

²⁰Recall that BDM stands for Becker, DeGroot, and Marschak (1964). In our experi-

Treatments. We designed six different treatments: four *baseline treatments* and two treatments with *belief elicitation*. In the baseline treatments, we varied the precision of common priors and the number of available pieces of advice. The precision of common priors was manipulated by adjusting the number of green balls in the jar: *weak prior* treatments contained 6 green balls, whereas *strong prior* treatments contained 7 green balls. The number of available pieces of advice was varied by providing either *one robot* or *two robots*. We employed a within-subject design with respect to the number of robots, and the order was counterbalanced.²¹ Common priors were presented on a between-subject level. Participants’ beliefs about the quality of their advisors were elicited only in a post-experimental survey.

In two additional one-robot treatments, which differed in the precision of common priors, we elicited participants’ beliefs about the quality of their advisors immediately after receiving advice in the first period (before guessing). We asked the participants: *Given the message of your advisor above, what is the probability that it is perfect? We remind you that your advisor can either be perfect or defective.* As in the baseline treatments, participants’ earnings were solely based on guessing correctly the color of the ball, taking into account the cost of buying advice. Beliefs about the quality of the advisor were not additionally incentivised.

ment, participants were asked to propose a price for advice. They would receive the advice only if their proposed price was at least as high as a hidden, randomly generated value. This procedure, equivalent to a second-price sealed-bid auction, created an incentive for participants to propose a price equal to their true willingness to pay for advice. Participants were explicitly informed about these incentives, in addition to having the procedure explained and their understanding tested.

²¹About half of the participants began the experiment with the one-robot treatment and then continued with the corresponding two-robot treatment, while the other half proceeded in the opposite order.

Procedure. The experiment was conducted in the experimental economics laboratory of the Toulouse School of Economics. The four baseline treatments were implemented in June 2021 and May 2023, and the belief elicitation treatments in December 2024.²² The experiment was programmed in oTree (Chen et al., 2016) and participants were recruited using a standard recruitment procedure.²³

After signing a consent form (reproduced in Appendix B), each participant was randomly assigned a computer terminal. Participants were aware that they were allowed to leave the experiment at any point, however, only a show-up fee would be paid in this case. Participants in the baseline treatment were informed that there were two parts in the experiment (treatments). They received information about the second part only after having completed the first part. Participants in the belief elicitation treatments participated in only one part (treatment).

In each part, participants first received some preliminary instructions and then read the task instructions (these are reproduced in Appendix C). These instructions were followed by a series of questions to verify their understanding (see Appendix D for the list of questions and the rate of correct responses).²⁴ Participants then started with the task as described above. At

²²We ran a total of 28 sessions: 22 with the baseline treatments and 6 with belief elicitation. On average, sessions had 9 participants in the baseline treatments and 16 in the belief elicitation treatments. Since we run the belief elicitation treatment after investigating the data of the previous treatments, we have pre-registered the experiment (AsPredicted #202830)

²³The experimental protocol was approved by the TSE/IAST ethics committee in May 2021.

²⁴If a participant selected the correct answer to a question, they were notified, and the next question appeared on the screen. Otherwise, the participant was informed that their answer was incorrect; the correct answer and the corresponding part of the instructions were displayed on the screen, and the participant had to click the “Next” button to

the beginning of each round the participants were reminded of the number of green balls in the jar with a visual reminder (see Figure 1, bottom left).

Participants were informed that, at the end of the experiment, one round from each part would be randomly selected for the final payout. Points earned during this round would be added to the show-up fee and then converted into euros at a rate of 100 to 1. Earnings in period 1 were: 500 points if the first guess was correct, and 0 points otherwise. Earnings in period 2 were either 0 or 500 points, depending on whether the second guess was correct, less the price paid for advice if it was purchased. In rounds 1 to 10 the price of advice was fixed at 5 points. In rounds 11 to 15 the participants had to propose a price in a range from 0 to 250, and received advice if and only if this price was higher than a randomly drawn value.²⁵ The show-up fee amounted to 500 points in the belief-elicitation treatments and 500 points in each part of the baseline treatments.

The participants received some post-experimental questions (presented in Appendix E) inquiring into: their choices, their skills in belief updating, their demographic characteristics, their willingness to solicit advice in real-life situations, and their attitude to risk.²⁶ In treatments with belief elicitation, these questions were asked at the end of the session. In the baseline treatments, participants received questions about their choices after each part,²⁷

proceed.

²⁵We drew the random value from a uniform distribution between 0 and 100, but we allowed participants to bid up to 250. The 0-100 range is based on theoretical predictions of behavior. We allowed a larger bid to avoid bunching and influencing choices.

²⁶Questions regarding participants' choices were mandatory, while other questions were optional; nonetheless, participation was nearly universal.

²⁷Recall that these parts corresponded to different treatments with the same number of green balls but a different number of robots.

with all other questions presented at the end of the session.

Participants collected their total earnings at the end of the experiment in an isolated room. Earnings were paid in cash. In the baseline treatments, participants were informed about their earnings for each part immediately after completing it.

Data. Table 2 provides an overview of the participants’ characteristics. Columns 1 to 3 correspond to the baseline treatment sessions, with Column 1 reporting statistics for all participants and Columns 2 and 3 broken down by year. Column 4 corresponds to the belief elicitation treatments.

In total we had 307 participants. The representation of genders was approximately balanced. Most (but not all) participants were students from different fields (53% reported following economics-related fields). Self-reported propensity to solicit advice and to take risk were not skewed in either direction. The participants showed a high performance on control questions, with the median participant answering correctly about 70% of the questions (see Appendix D for details). According to the post-experimental survey, participants did not consider the decisions to be particularly difficult. The experiment lasted, on average, 42 minutes in the baseline treatments (two parts) and 22 minutes in the belief elicitation treatment (one part). Overall earnings averaged 11-12 euros per part.²⁸

²⁸Participants in the baseline treatments earned roughly twice as much as those in the belief elicitation treatments because the baseline used a within-subject design involving two different numbers of robots, whereas the number of robots was held constant in the belief elicitation treatments.

Table 2: Number of observations and participants’ characteristics.

	All line ple	Base- line Sam- ple	June 2021	May 2023	Belief Elicit. December 2024
N. participants	208		122	86	99
N. participants (Prior 0.6)	107		61	46	52
N. participants (Prior 0.7)	101		61	40	47
Female	47.6%		41.8%	55.81%	63.64%
Age	22.71		22.72	22.69	21.08
Native French	82.84%		84.03%	81.18%	82.65%
Color-blind	0.97%		0.83%	1.16%	0.0%
Educ. [years after high school]	2.80		2.90	2.64	2.85
Economics majors	57.86%		53.57%	64.29%	44.74%
Self-report: Advice seeking (1-10)	5.51		5.74	5.19	5.36
Self-report: Risk seeking (1-10)	5.53		5.61	5.42	5.67
Self-report: Difficulty of decisions (1-10)	5.23		5.29	5.15	4.08
% of correct control questions	69.21%		72.16%	65.03%	71.07%
Bayesian Exercise 1	58.32		59.93	56.13	59.05
Bayesian Exercise 2	31.27		31.39	31.10	30.64
Earnings (euros)	23.16		23.36	22.88	11.78

Advice seeking, risk seeking, and difficulty in decision-making represent participants’ self-reported propensity to ask for advice, willingness to take risks in real life, and difficulty in making decisions during the experiment, measured on a scale from 1 to 10. The table reports average values. The Bayesian exercises tested participants’ ability in Bayesian updating and were optional. The correct solutions were 75 for Exercise 1 and 33 for Exercise 2.

4 Results.

This section is divided into three parts. The first presents our findings on the demand for advice in period 2 (hereafter, *demand*). The second reports participants’ beliefs about the quality of their advisors after receiving first period advice (hereafter, *posteriors*). The third describes the guessing patterns.²⁹

²⁹For the baseline treatments, we observe no significant differences between the two sets of sessions conducted in 2021 and 2023. We also find no effect of the order of treatments differentiated by the number of robots on the outcomes (demand for advice, rate of following advice, and posteriors). We, therefore, report our findings for the pooled dataset.

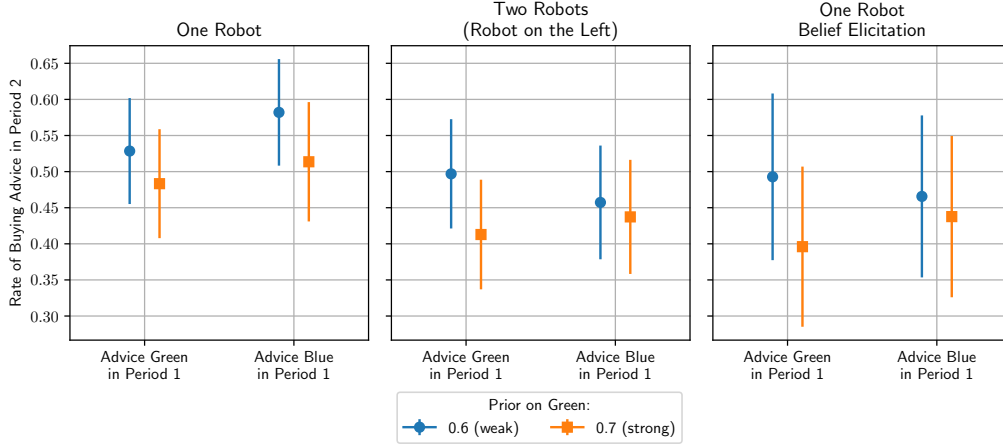
4.1 Inelastic demand and its segments.

Inelastic demand. Recall that, according to theoretical predictions, in the strong-prior treatment with one robot, and in either prior treatment with two robots, the demand for advice should be higher when period 1 advice is confirmatory. However, our data do not show this pattern.

Finding 1 (inelastic demand). *Demand for advice in period 2 does not depend on the content of advice received in period 1 in any treatment.*

Indeed, Figure 2 shows that the average rate of advice purchases at a fixed price does not depend on the content of period 1 advice: there is no statistically significant difference between purchases following confirmatory versus contradictory advice in either the weak- or strong-prior treatments.³⁰

Figure 2: The average rate of advice purchases across rounds 1 to 10.



For treatments with two robots, we present the demand for advice generated by the robot on the left of the screen. Reports 95% confidence intervals.

³⁰Unless specified otherwise, throughout the paper, figures report averages at the participant levels and 95% confidence intervals, computed as $\pm 1.96\hat{\sigma}$ where $\hat{\sigma}$ represents standard errors computed at the participant level.

We find that this result is quite robust as illustrated by Figures in Appendix F. It continues to hold when: restricting the sample to rounds 6 to 10, in order to control for learning (see Figure 11); when considering demand for advice generated by the robot on the right of the screen (instead of the left) in treatments with two robots (see Figure 13); when restricting the sample to participants who understood that the advisor does not change across the two periods of the same round without additional explanations (see Figure 12); and when controlling for demographics and year-fixed effects (see Tables 8 and 9).³¹

In the same line, Figure 3 shows that the WTP for advice among participants who wish to purchase it does not depend on the content of advice received in the previous period. This result passes all the robustness checks discussed in the previous paragraph.³²

Remark 1 *The average values presented in Figures 2 and 3 differ substantially from the theoretical predictions*

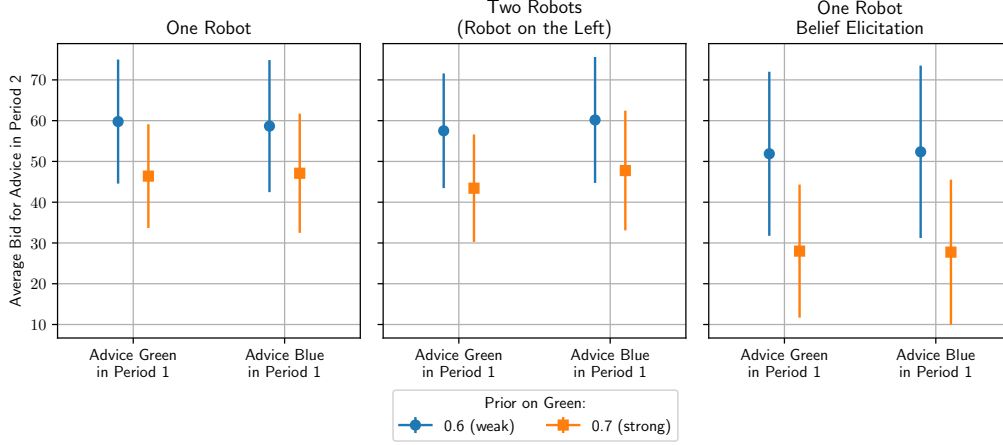
In particular, demand at a given price in the weak-prior treatment with one robot in Figure 2, left, is almost twice as low as predicted, despite the price of advice being fixed at only 1% of the reward for a correct guess. The average rate of purchasing advice and average WTP are similar across treatments.³³

³¹Actually, in the weak-prior treatment with one robot, adding participant-fixed effects yields a small but significant effect *opposite* to the theoretical predictions.

³²For brevity, we invite interested readers to request the supporting evidence.

³³Recall that in period 2 of rounds 11 to 15, we first ask participants whether or not they would like to purchase advice and only then elicit their WTP. This may explain why the average rate of purchasing at a price of 5 points, depicted in Figure 2, is only about one-half, while the average elicited WTP, depicted in Figure 3, is much higher than 5 points.

Figure 3: The average WTP for advice among participants who wish to purchase advice.



In rounds 11 to 15, we first ask participants whether or not they would like to purchase advice in period 2 and only then elicit their WTP. Reports 95% confidence intervals.

While Finding 1 and Remark 1 highlight deviations from theoretical predictions, the following demand pattern aligns well with rational behavior.

Remark 2 (*precision of common priors and demand*). *Demand in the weak prior treatments is higher than in the strong prior treatments*

This difference may be explained by the fact that more informative priors offer better guidance for independent guessing, which aligns well with participants' tendency to guess "green" when no advice is provided (see Set of Findings 3 below).

Remark 3 (*belief elicitation and demand*). *Belief elicitation about quality of advice is associated with a lower demand for advice*

Considering the average rate of purchasing advice (regardless of the content of period 1 advice) in rounds 1 to 10, we find 55% in the baseline weak prior

treatment with one robot, compared to only 47% in the weak-prior treatment with belief elicitation. In the strong prior treatments, the corresponding figures are 49% and 40%. Hence, participants’ rate of purchasing advice is about 8 percentage points lower when they are asked about its quality. One possible explanation is that this question makes the uncertainty of advice quality more salient.

Behavioral demand patterns. A closer examination of demand patterns reveals two behavioral phenomena. The first is participants’ tendency to treat advice as an *all-or-nothing* commodity.

Remark 4 (“all-or-nothing” demand). *Two prominent demand rules emerge: buy any available advice and buy no advice.*

Indeed, Table 3 shows that between 22% to 34% of participants (depending on treatment), either always bought any available advice (columns *All*) or never bought any advice (columns *Nothing*) across all rounds.³⁴

Table 3: Percentage of participants always and never buying any advice, by treatment.

Prior on Green	One Robot (baseline)			Two Robots			Belief Elicitation		
	Total	All	Nothing	Total	All	Nothing	Total	All	Nothing
0.6 (weak)	107	22 (20.56%)	14 (13.08%)	107	10 (9.35%)	20 (18.69%)	52	7 (13.46%)	11 (21.15%)
0.7 (strong)	101	10 (9.9%)	18 (17.82%)	101	0 (0.0%)	23 (22.77%)	47	4 (8.51%)	10 (21.28%)

Number and share of participants in a given treatment who always or never bought advice in rounds 1 to 15. *Total* corresponds to the sample size, *All* to the subsample who always buys, and *Nothing* to the subsample who never buys.

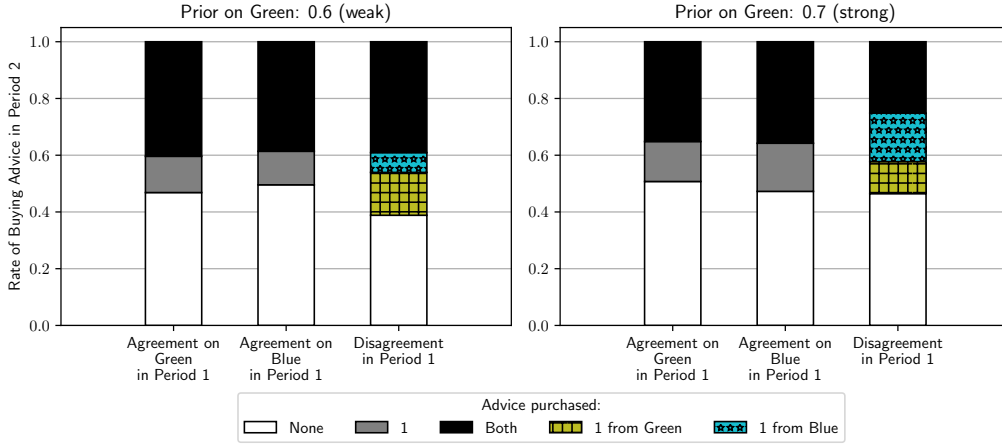
Comparison of the rows in Table 3 shows that a higher precision of common priors is associated with a lower share of participants fully adopting

³⁴Figures 14 and 15 in Appendix F depict the distribution of participants’ rates of purchasing advice in the baseline treatments with one robot and in the belief elicitation treatments.

the rule “All” and a higher share of participants fully adopting the rule “Nothing,” consistent with demand being relatively high in the weak prior treatments, as highlighted in Remark 2.

Furthermore, in the baseline treatments with two robots, participants who do not fully adopt one of the above two demand rules tend to purchase either both pieces of advice or none in each round, as shown in Figure 4. Notice that this holds even when one piece of advice in period 1 is confirmatory while the other is contradictory, revealing their different qualities. Recall that, theoretically, this should induce participants to select one piece of advice, namely the one that was confirmatory in period 1.

Figure 4: All-or-nothing demand pattern in treatments with two robots.



Participants can purchase 0, 1 or 2 pieces of advice. The bars stack the rates at which they made each of these strategies. *1 from Green(Blue)* refers to buying one piece of advice from the robot that advised Green(Blue).

Focusing on participants who do not fully adopt either demand rule identified in Remark 4, we observe an alternative demand pattern.

Remark 5 (*hot-hand fallacy*). *The participants who do not fully adopt*

either rule All or Nothing, tend to maintain the same demand as in the previous round, as long as it is associated with successful guessing

Figure 5 illustrates this tendency by showing the rates at which participants maintained their previous-round demand (excluding those who always bought all advice or never bought any), depending on whether their previous-round guess was correct (“won”) or incorrect (“lost”).³⁵ Participants were significantly more likely to maintain their demand after a correct guess, whereas after an incorrect guess they were about equally likely to stick with or change their demand. Note that such “win-stay, lose-shift” demand rule cannot improve guessing accuracy because advisor quality is independent across rounds. This may reflect a hot-hand fallacy, as documented in the literature, pioneered by Gilovich et al. (1985).

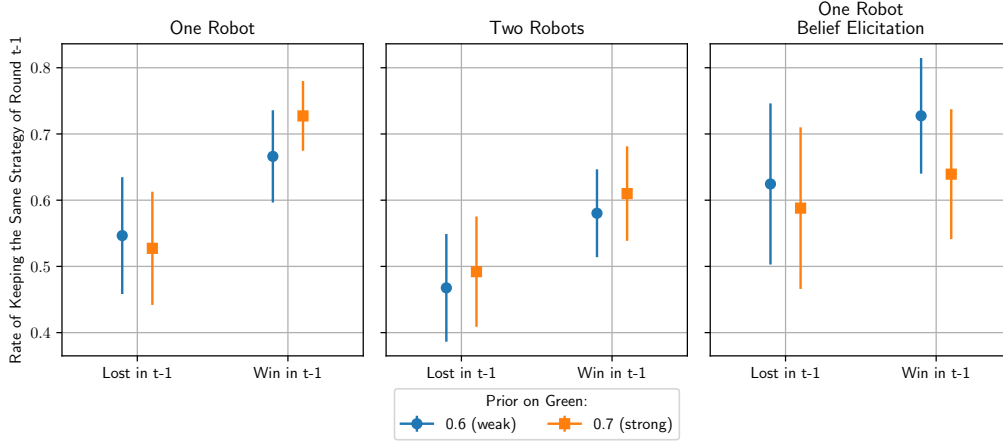
4.2 Persistent beliefs about advice quality.

Finding 1 suggests that at least one ingredient of reputational cheap-talk theory is missing: either participants do not revise their beliefs about the quality of advice based on its content, or their subsequent demand for advice is influenced by factors other than the perceived quality of that advice.³⁶ We use data from treatments with beliefs elicitation to examine whether participants revise their beliefs in response to the advice received and whether their

³⁵In treatments with one robot, we distinguish between two demand choices: buying one piece of advice or buying none. In treatments with two robots, we distinguish among three choices: buying both pieces of advice, buying one, or buying none. We observe similar patterns when considering a finer categorization of four demand choices, separating buying the advice shown by the robot on the left from buying the advice shown by the robot on the right instead of grouping them as “buying one piece of advice”.

³⁶Note that Remarks 4 and 5 support the latter, while leaving open the possibility that the former also holds.

Figure 5: Hot hand demand.



Reports 95% confidence intervals.

subsequent demand increases with these updated beliefs.

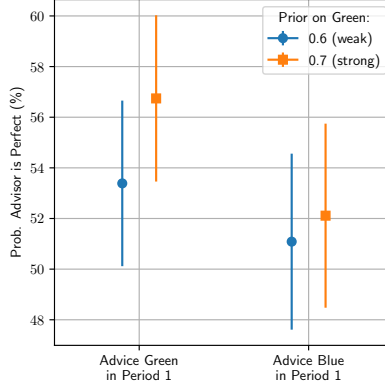
Finding 2 (beliefs about quality of advice). *In the strong prior treatment, average posteriors about the quality of advice are 4.8 percentage points higher following confirmatory advice than following contrarian advice. In the weak prior treatment, there is no statistically significant difference.*

Indeed, the average posteriors in treatments with belief elicitation during the experiment are generally higher when period-one advice is confirmatory (see Figure 6), however, this difference is small relative to theoretical predictions and is significant only in the strong prior treatment.³⁷

While 28.8% of participants in the weak prior treatment and 25% in the strong prior treatment assign, on average, a higher posterior probability to advice being perfect following confirmatory rather than contradictory advice, consistent with Bayesian updating, 11.54% in the weak prior treatment and

³⁷Figure 16 in Appendix F plots the distribution of average beliefs, by advice received.

Figure 6: Average posteriors in rounds 1-10 of belief elicitation treatments.



Reports 95% confidence intervals. Rounds 1-10.

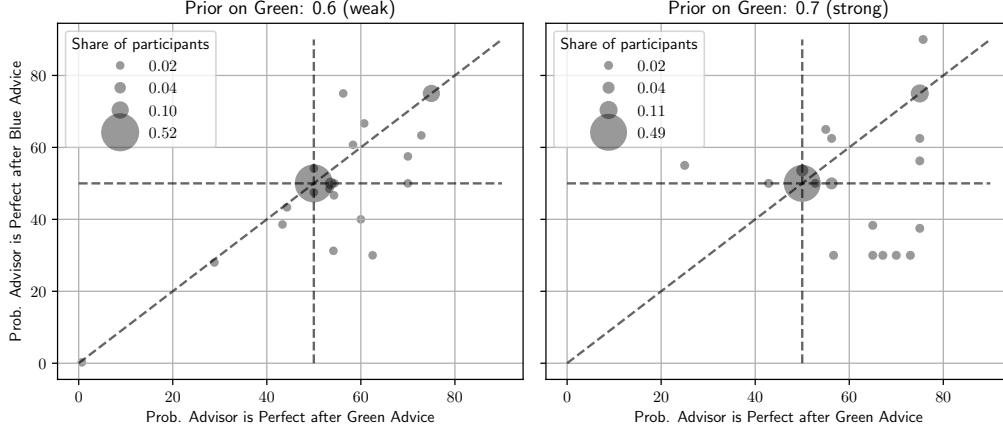
13.46% in the strong prior treatment update their beliefs in the opposite direction, and about half of participants (51.92% in the weak prior treatment and 42.31% in the strong prior treatment) do not update their beliefs about advice quality based on its content *at all*.³⁸ This is illustrated in Figure 7, which plots participants’ average posterior beliefs following confirmatory advice (horizontal axis) against posterior beliefs following contrarian advice (vertical axis). Each circle represents a particular combination of posterior values, with the circle’s radius proportional to the percentage of participants exhibiting that posterior ratio. In either treatment, we observe a large concentration at 50%, resulting in a prominent “bubble.”³⁹

Finding 3 (posteriors and demand). *Higher posteriors on the quality of advice are not associated with a higher demand for advice.*

³⁸We observe a similar pattern in the post-experimental survey from the baseline treatments, in which participants were asked to report their beliefs about the quality of an advisor who recommended guessing “green”: 58.6% of participants in the weak prior treatment and 35.1% in the strong prior treatment answered 50%.

³⁹Figure D.6 in Appendix D depicts, for each treatment, the distribution of average beliefs by advice received.

Figure 7: Participants' posteriors: confirmatory (horizontal) vs. contrarian (vertical) advice.



Each circle represents a particular posterior ratio, with the radius proportional to the percentage of participants exhibiting that ratio.

Table 4 reports OLS regression estimates at the round level, where the dependent variable is a dummy taking value 1 if the participant purchased advice, and the explanatory variable is his or her posterior on the quality of advice. The estimates are not statistically different from zero, even with fixed effects, which absorb variation from participants who do not update their beliefs after receiving the first advice. Thus, a higher perceived advice quality is *not* associated with a higher likelihood of purchasing advice.

4.3 Use of advice.

While participants are rather reluctant to update their beliefs about advice quality based on its (mis)alignment with common priors, and their demand for advice appears driven by considerations other than perceived quality, they still rely on advice in their guesses in a manner broadly consistent with

Table 4: Demand as a function of posteriors.

	(1)	(2)	(3)	(4)
	Weak prior	Weak prior	Strong prior	Strong prior
posteriors	-0.00219 (0.00246)	0.000756 (0.00118)	-0.000510 (0.00252)	-0.000303 (0.00195)
Constant	0.590*** (0.134)	0.435*** (0.0618)	0.432** (0.151)	0.421*** (0.107)
Observations	520	520	470	470
ParticipantFE	No	Yes	No	Yes

Standard errors in parentheses

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Bayesian updating.

Set of findings 4 (influential advice). *(i) Participants tend to follow the advice in treatments with one robot, and are even more likely to follow agreeing advice in treatments with two robots.*

(ii) These tendencies are amplified when the advice is confirmatory, especially under strong priors.

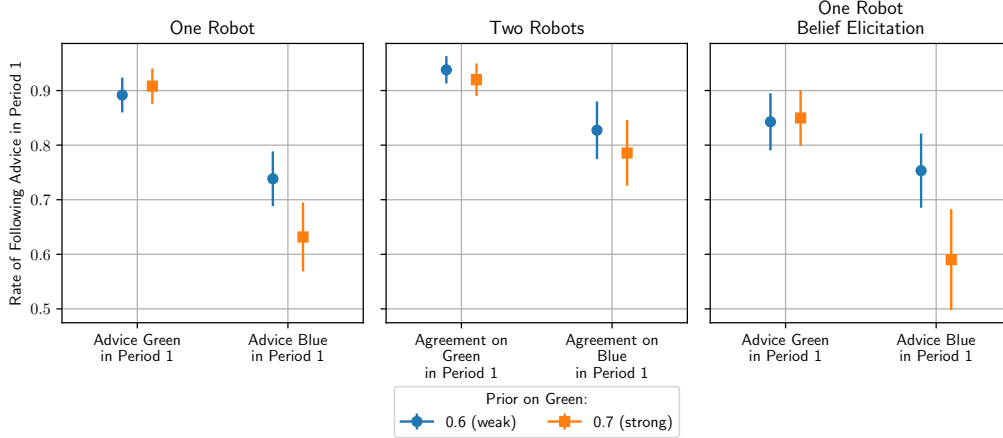
(iii) When the two robots provide conflicting advice, or when participants guess independently, most rely on the common priors.

Figure 8 illustrates patterns (i) and (ii) focusing on guessing in period 1.

Similar patterns arise for guessing in period 2, with a generally stronger impact of advice on guessing (conditional on purchasing). Using the data from baseline treatments with one robot, we find that participants are 6.21pp more likely to follow confirmatory advice while 2.11pp less likely to follow contrarian advice in period 2 than in period 1. These differences are positive and respectively 3.83% and 2.91% in strong prior treatments.⁴⁰ Yet, the

⁴⁰These results contrast with empirical evidence where individuals tend to overweight

Figure 8: Guessing patterns.



Average rate of guessing the color that was advised in period 1, rounds 1 to 15. Reports 95% confidence intervals.

rate of following purchased advice is not universal, suggesting that some participants buy advice for non-instrumental reasons.

Remark 6 (*influence regardless of contents of first advice*). *The rate of following advice in period 2 (conditional on buying) does not depend on the contents of advice received in period 1.*

Combined with Finding 1, Remark 6 suggests that the alignment of advice with priors in period 1 has no impact on decisions in period 2.

As to pattern (iii), when faced with conflicting pieces of advice in the baseline treatments with two robots, or when not purchasing advice in period 2, about *three quarters* of participants guess “green”, relying thereby on common priors.⁴¹ Specifically, the share of “green” guesses following con-

their priors and forgo expert advice (Bouacida et al., 2025) and learning from others (Weizsäcker, 2010). However, Malmendier and Shanthikumar (2007) find, instead, that small investors follow upward-biased analyst recommendations literally.

⁴¹This is consistent with previous evidence where, facing conflicting advice, participants neglected the furthest ones from their prior (Yaniv and Milyavsky, 2007).

flicting advice in period 1 is: 74% under weak priors and 75% under strong priors; and in period 2: 63% under weak priors and 75% under strong priors. When advice is not purchased in period 2, participants guess “green” between 61% and 76% of the time, depending on the treatment.

The above guessing patterns are consistent with participants combining common priors and advice to update their beliefs about the correct guess, broadly in line with Bayesian updating.

5 Conclusion.

We conducted a laboratory experiment to test how the content of advice in an incentivised guessing task influences beliefs about its quality and future demand. Although participants could infer advice quality from its alignment with common priors, such alignment had no significant effect on either their beliefs about quality of advice or subsequent demand for it. At the same time, guessing patterns suggest that participants combined priors and advice when forming decisions.

We hope that these results will motivate further research on the demand for professional advice in real markets. Field experiments, in particular, could provide valuable insights into the external validity of the mechanisms examined in this paper. An important open question concerns which indicators or cues people use to evaluate the quality of advice. In our experiment, alignment with common priors fulfilled this role; however, in real world settings, other signals may play a more prominent part. The extent to which individuals’ demand for advice depends on such evaluations likely varies across markets. For example, media consumers may rely on news reports to guide

private decisions (such as voting) without critically assessing their credibility, whereas consumers of financial advice are likely to be more attentive to quality. A better understanding of these and related questions could help assess the influence of professional advice on private decision making.

References

- Alpert, A., D. Lakdawalla, and N. Sood (2023, May). Prescription drug advertising and drug utilization: The role of Medicare Part D. *Journal of Public Economics* 221, 104860. [2](#)
- Ambuehl, S. and S. Li (2018). Belief updating and the demand for information. *Games and Economic Behavior* 109, 21–39. [2](#), [6](#)
- Augenblick, N., E. Lazarus, and M. Thaler (2025, February). Overinference from Weak Signals and Underinference from Strong Signals. *The Quarterly Journal of Economics* 140(1), 335–401. [6](#)
- Aydogan, I., A. Baillon, E. Kemel, and C. Li (2025). How much do we learn? Measuring symmetric and asymmetric deviations from Bayesian updating through choices. *Quantitative Economics* 16(1), 329–365. [6](#)
- Becker, G. M., M. H. Degroot, and J. Marschak (1964). Measuring utility by a single-response sequential method. *Behavioral Science* 9(3), 226–232. [4](#)
- Benjamin, D. J. (2019). Errors in probabilistic reasoning and judgment biases. *Handbook of Behavioral Economics: Applications and Foundations* 1 2, 69–186. [2](#)

Bouacida, E., R. Foucart, and M. Jalloul (2025, August). When expert advice fails to reduce the productivity gap: Experimental evidence from chess players. *Journal of Economic Behavior & Organization* 236, 107124. 32

Camara, F. and N. Dupuis (2021). Structural Estimation of Expert Strategic Bias: The Case of Movie Critics. *Working Paper*. 2

Camerer, C. (1998). Bounded rationality in individual decision making. *Experimental economics* 1, 163–183. 2

Charness, G., R. Oprea, and S. Yuksel (2021, June). How do people choose between biased information sources? evidence from a laboratory experiment. *Journal of the European Economic Association* 19(3), 1656–1691. 7

Chopra, F., I. Haaland, and C. Roth (2024). The demand for news: Accuracy concerns versus belief confirmation motives. *The Economic Journal* 134(661), 1806–1834. 7

Currie, J., A. Li, and M. Schnell (2025). The effects of competition on physician prescribing. *NBER Working Paper*. 3

Edwards, W. (1982). *Conservatism in human information processing*, pp. 359–369. Cambridge University Press. 5

Gentzkow, M. and J. M. Shapiro (2008, June). Competition and Truth in the Market for News. *Journal of Economic Perspectives* 22(2), 133–154. 3

- Gilovich, T., R. Vallone, and A. Tversky (1985, July). The hot hand in basketball: On the misperception of random sequences. *Cognitive Psychology* 17(3), 295–314. 4, 27
- Golman, R., G. Loewenstein, A. Molnar, and S. Saccardo (2022, September). The Demand for, and Avoidance of, Information. *Management Science* 68(9), 6454–6476. 3, 7
- Gächter, S. and E. Renner (2010). The effects of (incentivized) belief elicitation in public goods experiments. *Experimental Economics* 13(3), 364–377. 6
- Jin, C. (2024). Does competition improve information quality? evidence from the security analyst market. *Toulouse School of Economics Working Paper No. 1553*. 3
- Kahneman, D. and A. Tversky (1973). On the psychology of prediction. *Psychological Review* 80(4), 237–251. 5
- Kapons, M. and P. Kelly (2025, May). Biased Inference due to Prior Beliefs: Evidence from the Field. *Journal of Political Economy Microeconomics*. 6
- Kelly, C. A. and T. Sharot (2021, December). Individual differences in information-seeking. *Nature Communications* 12(1), 7062. 7
- Malmendier, U. and D. Shanthikumar (2007, August). Are small investors naive about incentives? *Journal of Financial Economics* 85(2), 457–489. 32

- Mastorocco, N., A. Ornaghi, M. Pograxha, and S. Wolton (2025). Man bites dog: Editorial choices and biases in the reporting of weather events. *Working Paper*. 2
- Meloso, D., S. Nunnari, and M. Ottaviani (2023, September). Looking into crystal balls: A laboratory experiment on reputational cheap talk. *Management Science* 69(9), 4973–5693. 5, 6
- Miller, J. B. and A. Sanjurjo (2018). Surprised by the Hot Hand Fallacy? A Truth in the Law of Small Numbers. *Econometrica* 86(6), 2019–2047. 4
- Nunnari, S. and G. Montanari (2025). Audi alteram partem: An experiment on selective exposure to information. *Journal of the Economic Science Association*, 1–12. 7
- Reshidi, P., A. Lizzeri, L. Yariv, J. H. Chan, and W. Suen (2021, December). Individual and Collective Information Acquisition: An Experimental Study. Working Paper 29557. 3
- Schoar, A. and Y. Sun (2024). Financial advice and investor beliefs: Experimental evidence on active vs. passive strategies. *NBER Working Paper* (33001). 7
- Sivey, P. and Y. Chen (2019, 06). Competition and quality in healthcare. 3
- Thaler, M. (2024, May). The Fake News Effect: Experimentally Identifying Motivated Reasoning Using Trust in News. *American Economic Journal: Microeconomics* 16(2), 1–38. 16

Weizsäcker, G. (2010, December). Do We Follow Others When We Should? A Simple Test of Rational Expectations. *American Economic Review* 100(5), 2340–2360. [32](#)

Yaniv, I. and M. Milyavsky (2007, May). Using advice from multiple sources to revise and improve judgments. *Organizational Behavior and Human Decision Processes* 103(1), 104–120. [32](#)

A Proof of theoretical predictions.

A.1 Proof of Prediction Set 1.

Proof of prediction (i). By set of equations [\(2\)](#) and Bayes' rule,

$$\Pr(x_1 = G \mid a_1^l = G) = \frac{3p}{2p+1}, \quad \Pr(x_1 = B \mid a_1^l = G) = \frac{1-p}{2p+1}, \quad (4)$$

$$\Pr(x_1 = G \mid a_1^l = B) = \frac{p}{3(1-p)+p}, \quad \Pr(x_1 = B \mid a_1^l = B) = \frac{3(1-p)}{3(1-p)+p}. \quad (5)$$

By set of equations [\(4\)](#),

$$\Pr(x_1 = G \mid a_1^l = G) > \Pr(x_1 = B \mid a_1^l = G). \quad (6)$$

By set of equations [\(5\)](#) and the upper bound in [\(3\)](#),

$$\Pr(x_1 = B \mid a_1^l = B) > \Pr(x_1 = G \mid a_1^l = B) \text{ if and only if } p < 0.75,$$

Therefore, the condition holds for either p in set $\{0.6, 0.7\}$.

Proof of prediction (ii). By set of equations [\(2\)](#) and Bayes' rule,

$$\Pr(q^l = 1 \mid a_1^l = G) = \frac{2p}{2p+1}, \quad (7)$$

$$\Pr(q^l = 1 \mid a_1^l = B) = \frac{2(1-p)}{2(1-p)+1}. \quad (8)$$

By substituting $p \in \{0.6, 0.7\}$ into equations (7) and (8) we obtain the posterior values reported in part (ii) of Prediction Set 1.

Proof of predictions (iv) and (v). In period 2, the advice is correct with probability $\Pr(q^l = 1 \mid a_1^l) + \frac{1}{2} \Pr(q^l = 0 \mid a_1^l)$. The best alternative to purchasing and following advice is to guess G , which is correct with probability p . Therefore, the decision maker's WTP for advice is given by:

$$([\Pr(q^l = 1 \mid a_1^l) + \frac{1}{2} \Pr(q^l = 0 \mid a_1^l)] - p) R, \text{ where } R = 500. \quad (9)$$

Applying straightforward calculus with equations (9), (7) and (8) yields the values reported in part (iv) of Prediction Set 1.

Proof of prediction (iii). Comparing the WTP for advice with its price of 5 points, we find the demand described in part (iv) of Prediction Set 1. Note that demand following advice G is positive if and only if $p < 0.79$, while that following advice B is positive if and only if $p < 0.68$.

A.2 Proof of Prediction Set 2.

Proof of prediction (i). It is optimal to guess G if $\Pr(x_1 = G \mid a_1^l, a_1^r) > \frac{1}{2}$ and B otherwise. By Bayes' rule,

$$\Pr(x_1 = G \mid a_1^l = a_1^r = G) = \frac{5p}{4p+1} > \frac{1}{2},$$

$$\Pr(x_1 = G \mid a_1^l = G, a_1^r = B) = \Pr(x_1 = G \mid a_1^l = B, a_1^r = G) = p > \frac{1}{2}, \quad (10)$$

$$\Pr(x_1 = B \mid a_1^l = B, a_1^r = B) = \frac{5(1-p)}{5(1-p)+p} > \frac{1}{2} \text{ if and only if } p < 0.8(3), \quad (11)$$

hence, for either $p \in \{0.6, 0.7\}$.

Proof of prediction (ii). Suppose first that $a_1^i = G$ and $a_1^{-i} = B$, where $i \in \{l, r\}$ and $-i = \{l, r\} \setminus \{i\}$ (here and throughout). By Bayes' rule, the decision maker infers that the coin has landed tails, that is, $z_1 = T$, and that the qualities of the two pieces of advice differ:

$$\Pr(q^l = q^r = 0 \mid a_1^l \neq a_1^r) = \Pr(q^l = q^r = 1 \mid a_1^l \neq a_1^r) = 0. \quad (12)$$

With probability p , advice i is perfect and advice $-i$ is defective:

$$\Pr(q^i = 1 \mid a_1^i = G, a_1^{-i} = B) = \Pr(q^i = 1, q^{-i} = 0 \mid a_1^i = G, a_1^{-i} = B) = p, \quad (13)$$

$$\Pr(q^{-i} = 1 \mid a_1^i = G, a_1^{-i} = B) = \Pr(q^i = 0, q^{-i} = 1 \mid a_1^i = G, a_1^{-i} = B) = 1 - p. \quad (14)$$

Suppose now that $a_1^l = a_1^r = G$. By Bayes' rule,

$$\Pr(q^l = q^r = 1 \mid a_1^l = a_1^r = G) = \frac{2p}{1+4p}, \quad (15)$$

$$\Pr(q^l = q^r = 0 \mid a_1^l = a_1^r = G) = \frac{1}{1+4p}, \quad (16)$$

$$\Pr(q^i = 1, q^{-i} = 0 \mid a_1^l = a_1^r = G) = \frac{p}{1+4p}. \quad (17)$$

By equations (15) and (17),

$$\Pr(q^i = 1 \mid a_1^l = a_1^r = G) = \frac{3p}{1+4p}. \quad (18)$$

Suppose, finally, that $a_1^l = a_1^r = B$. By Bayes rule,

$$\Pr(q^l = q^r = 1 \mid a_1^l = a_1^r = B) = \frac{2(1-p)}{5-4p}, \quad (19)$$

$$\Pr(q^l = q^r = 0 \mid a_1^l = a_1^r = B) = \frac{1}{5-4p}, \quad (20)$$

$$\Pr(q^i = 1, q^{-i} = 0 \mid a_1^l = a_1^r = B) = \frac{1-p}{5-4p}. \quad (21)$$

By equations (19) and (21),

$$\Pr(q^i = 1 \mid a_1^l = a_1^r = B) = \frac{3(1-p)}{5-4p}, \text{ where, recall, } i \in \{l, r\}. \quad (22)$$

Proof of prediction (v). We describe the best period 2 guess depending on the decision maker's information set⁴²

$$\Omega \in \{\emptyset, \{a_2^{i^*}\}, \{a_2^l, a_2^r\}\}, \text{ where}$$

$$i^* = \arg \max_{i=l,r} \{\Pr(q^i = 1 \mid a_1^l, a_1^r)\}.$$

Trivially, if $\Omega = \{\emptyset\}$, the optimal guess is G . Suppose that $\Omega = \{a_2^{i^*}\}$. Then, it is optimal to guess $a_2^{i^*}$. Indeed, by Bayes' rule and set of equations (13), (14), (15) to (18) and (19) to (22),

$$\begin{aligned} \Pr(x_2 = G \mid a_2^{i^*} = G, a_1^l, a_1^r) &\geq \Pr(x_2 = G \mid a_2^{i^*} = G, a_1^l = a_1^r = B) = \\ &= \frac{(1 + \Pr(q^{i^*}=1 \mid a_1^l = a_1^r = B))p}{(1 + \Pr(q^{i^*}=1 \mid a_1^l = a_1^r = B))p + (1 - \Pr(q^{i^*}=1 \mid a_1^l = a_1^r = B))(1-p)} = \frac{(8-7p)p}{(8-7p)p + (2-p)(1-p)} > \frac{1}{2} \end{aligned}$$

for any $p < 1.625$.

$$\begin{aligned} \Pr(x_2 = B \mid a_2^{i^*} = B, a_1^l, a_1^r) &\geq \Pr(x_2 = B \mid a_2^{i^*} = B, a_1^l = a_1^r = B) = \\ &= \frac{(1 + \Pr(q^{i^*}=1 \mid a_1^l = a_1^r = B))(1-p)}{(1 + \Pr(q^{i^*}=1 \mid a_1^l = a_1^r = B))(1-p) + (1 - \Pr(q^{i^*}=1 \mid a_1^l = a_1^r = B))p} = \frac{(8-7p)(1-p)}{(8-7p)(1-p) + (2-p)p} > \frac{1}{2} \end{aligned}$$

for any $p < 0.703$. Note that both these inequalities hold for either $p \in \{0.6, 0.7\}$.

Finally, consider information set $\Omega = \{a_2^l, a_2^r\}$. We show that the optimal guess is B if $a_2^l = a_2^r = B$ and G otherwise, for both $p \in \{0.6, 0.7\}$ and for

⁴²Trivially, the advice j is superior information to the advice $-j = \{l, r\} \setminus \{j\}$. We therefore do not include a_2^{-j} into set Ω .

all possible realizations of the first advice. Suppose that $a_1^l = a_1^r = B$. By Bayes' rule and equations (19) to (21),

$$\Pr(x_2 = B \mid a_t^i = B, t = 1, 2) = \frac{(7-6p)(1-p)}{(7-6p)(1-p)+p} > \frac{1}{2} \text{ for any } p < 0.725,$$

$$\Pr(x_2 = G \mid a_2^i = G, a_1^i = B) > \Pr(x_2 = G \mid a_2^i = G, a_2^{-i} = B, a_1^i = B) = p > \frac{1}{2}.$$

Hence, the optimal guess is B if $a_2^l = a_2^r = B$, and G otherwise. Suppose now that $a_1^l = a_1^r = G$. By Bayes' rule and equations (15) to (17),

$$\Pr(x_2 = B \mid a_2^i = B, a_1^i = G) = \frac{(6p+1)(1-p)}{(6p+1)(1-p)+p} > \frac{1}{2} \text{ for any } p < 0.86,$$

$$\Pr(x_2 = G \mid a_2^i = G, a_1^i = G) > \Pr(x_2 = G \mid a_2^i = G, a_2^{-i} = B, a_1^i = G) = p > \frac{1}{2}.$$

Once again, the optimal guess is B if $a_2^l = a_2^r = B$, and G otherwise. Suppose, finally, that $a_1^i = B, a_1^{-i} = G$. By Bayes' rule and equations (10) and (11),

$$\Pr(x_2 = B \mid a_2^{i*} = B, a_2^{-i*} = B, a_1^i \neq a_1^{-i}) = 1 > \frac{1}{2},$$

$$\begin{aligned} & \Pr(x_2 = G \mid a_2^i = G, a_1^i \neq a_1^{-i}, i = i, r) > \\ & \Pr(x_2 = G \mid a_2^{i*} = G, a_2^{-i*} = B, a_1^i \neq a_1^{-i}, i = i, r) > \\ & \Pr(x_2 = G \mid a_2^{-i*} = G, a_2^{i*} = B, a_1^i \neq a_1^{-i}, i = i, r) = p > \frac{1}{2}. \end{aligned}$$

Once again, the optimal guess is B if $a_2^l = a_2^r = B$, and G otherwise.

Proof of prediction (iv). The expected efficiency of the above optimal guess is: p if $\Omega = \{\emptyset\}$;

$$\frac{1 + \Pr(q^{i*} = 1 \mid a_1^l, a_1^r)}{2} \quad (23)$$

if $\Omega = \{a_2^{i*}\}$ and

$$\begin{aligned} & p \left(\Pr(q^l + q^r \geq 1 \mid a_1^l, a_1^r) + \frac{1}{2} \Pr(q^l = q^r = 0 \mid a_1^l, a_1^r) \right) + \\ & (1-p) \left(\Pr(q^l = q^r = 1 \mid a_1^l, a_1^r) + \frac{1}{2} \Pr(q^l = q^r = 0 \mid a_1^l, a_1^r) \right) \end{aligned} \quad (24)$$

if $\Omega = \{a_2^l, a_2^r\}$. Comparing the expression (23) to p , we find that the decision maker's WTP for advice a_2^{i*} is equal to

$$\left(\frac{1 + \Pr(q^{i*} = 1 | a_1^l, a_1^r)}{2} - p \right) R, \text{ where } R = 500. \quad (25)$$

Comparing the expressions (23) and (24), we find that the decision maker's WTP for advice a_2^{-i*} (in addition to advice a_2^{i*}) is equal to

$$\frac{R}{2} (p \Pr(q^{-i*} = 1 | a_1^l, a_1^r) (1 - \Pr(q^{i*} = 1 | a_1^l, a_1^r)) - (1 - p) \Pr(q^{i*} = 1 | a_1^l, a_1^r) (1 - \Pr(q^{-i*} = 1 | a_1^l, a_1^r))) . \quad (26)$$

Finally, comparing expression (24) with p , we find that the decision maker's WTP for two pieces of advice is equal to the difference between (a) the probability that the above optimal guess based on this advice - namely, choosing G if at least one advice is G and B otherwise - matches the state, and (b) the prior probability p , scaled by R :

$$\begin{aligned} & (p (1 - \frac{1}{2} \Pr(q^l = q^r = 0 | a_1^l, a_1^r)) + (1 - p) (\Pr(q^l = q^r = 1 | a_1^l, a_1^r) + \\ & \quad + \frac{1}{2} (1 - \Pr(q^l = q^r = 1 | a_1^l, a_1^r))) - p) R. \end{aligned} \quad (27)$$

Straightforward calculus using expressions (25) to (27) and equations (10) to (22), yields figures in Table 1.

Proof of prediction (iii). Comparing the WTP values in Table 1 with the price of advice (5 points per piece), we obtain the demand described in point (v) of Prediction Set 2.

B Consent form.

The experiment was conducted in French; the English translations are presented here and throughout the text.

Consent Form and Information on Data Privacy

The Institute for Advanced Studies of Toulouse and the Toulouse School of Economics (1 Esplanade de l'Université, 31080 Toulouse, Cedex 06, France) are conducting a laboratory experiment today at the Experimental Economics Laboratory of the Toulouse School of Economics. You have been invited because you are registered in our recruitment system. The procedures of the experiment will be explained to you before the experiment begins.

During the experiment, you will be asked to complete certain tasks, and your responses will be recorded in our computer system. The information recorded during the experiment will not allow any conclusions to be drawn about the participation or behavior of individual persons. The analysis of the data and the presentation of the results of this experiment will be carried out exclusively in anonymous form. Anonymous data will be archived and may be made available to other researchers for research purposes.

There will be no link between the data generated during the experiment and the data in the recruitment system. Receipts completed during payment will always be kept separately. Your participation today is entirely voluntary. If you do not participate, there will be no disadvantage to you. However, note that in this case your earnings from the experiment will be adjusted appropriately. You may withdraw from the experiment at any time.

For your participation and the use of your data, we ask for your consent. Consent may be revoked at any time, for example by email to: [*email contact*]

This consent is the legal basis for any use of data.

You can find a copy of this information at the experimental laboratory.

I have read the information on data privacy and agree to participate in

the experiment and to the use of the related data as described above:

☐ Yes ☐ No

Date

Signature

Without consent, you cannot participate in the experiment today. Please inform the person in charge of the experiment in this case. Thank you.

C Instructions.

C.1 Instructions for weak prior treatment with one robot.⁴³

Thank you for participating in our experiment!

In this experiment, you can earn money based on your decisions. We therefore ask you to read the instructions carefully. Your earnings during the experiment will be calculated in points. Your points will be exchanged for euros at the end of the experiment at the following exchange rate:

$$100 \text{ points} = 1 \text{ euro}$$

In today's experiment, you will be asked to guess the color of a ball drawn from a pot several times. You will earn points based on your answers.

You are not allowed to communicate with other participants during the experiment. We also ask you to turn off your mobile phones now. If you

⁴³Instructions for the strong prior treatment with one robot are the same, with "6" (green balls) replaced by "7" and "4" (blue balls) replaced by "3."

have any questions, please raise your hand and someone will come to your desk to answer them.

What you will do

You will play a session of 15 rounds.

One of these rounds will be selected to determine your compensation at the end of the session.

Each round consists of two periods.

In each period, a ball is drawn from (and returned to) a pot containing 10 balls, of which 6 are green and 4 are blue, as will be shown on your computer screen. You will guess the color of this ball. You earn 500 points if your answer is correct.

The First Period

Before giving your answer in the first period, you will automatically receive advice from your computerized advisor. In other words, your advisor will show either a blue circle (suggesting you respond “blue”) or a green circle (suggesting you respond “green”).

Your advisor is either perfect or defective, with equal probability. The quality of your advisor is determined anew at the beginning of each round and remains the same during both periods of that round. You will not be informed about the quality of your advisor in advance.

If the advisor is perfect, their advice is correct (i.e., it shows a green circle if the selected ball is green, and a blue circle if the selected ball is blue).

If the advisor is defective, the content of their advice depends on the outcome of a coin flip.

After receiving the advice, you must guess whether the selected ball is blue or green. You will learn whether your guess was correct at the end of the round. You earn 500 points if your guess is correct. You will also learn the quality of your advisor at the end of each round.

The Second Period

Your advisor is the same as in the first period. You do not automatically receive advice. You can request advice by clicking the “Advice” button. Advice is costly (the cost is detailed below). Alternatively, you can answer without advice. Once your answer is submitted, you will learn whether it is correct or incorrect. You earn 500 points if your answer is correct. Additionally, at the end of each round, you will also learn whether your advisor was perfect or defective.

Cost of Advice in the Second Period

Rounds 1 to 10: If you request advice in period 2, you pay 5 points.

Rounds 11 to 15: If you request advice in period 2, you must propose an amount between 0 and 250 (in points) that you are willing to pay for the advice. The experiment is designed such that it is optimal for you to propose the highest amount you are willing to pay.

Indeed, a threshold price X for the advice will be drawn randomly. If the amount you propose to pay exceeds this threshold, you will receive the requested advice and pay the threshold price X (not the amount you proposed). Otherwise, you will have to guess without advice.

C.2 Instructions in weak prior treatment with two robots.⁴⁴

Thank you for participating in our experiment!

In this experiment, you can earn money based on your decisions. We therefore ask you to read the instructions carefully. Your earnings during the experiment will be calculated in points. Your points will be exchanged for euros at the end of the experiment at the following exchange rate:

$$100 \text{ points} = 1 \text{ euro}$$

In today's experiment, you will be asked several times to guess the color of a ball drawn from a pot. You will earn points based on your answers.

You are not allowed to communicate with other participants during the experiment. We also ask you to turn off your mobile phones now. If you have any questions, please raise your hand, and someone will come to your desk to answer them.

What you will do

You will play a session of 15 rounds.

One of these rounds will be selected to determine your compensation at the end of the session.

Each round consists of two periods.

In each period, a ball is drawn from (and returned to) a pot containing 10 balls, of which 6 are green and 4 are blue, as will be shown on your computer

⁴⁴Instructions for the strong prior treatment with two robots are the same, with “6” (green balls) replaced by “7” and “4” (blue balls) replaced by “3.”

screen. You will guess the color of this ball. You earn 500 points if your answer is correct.

First Period

Before guessing in the first period, you will automatically receive advice from your computerized advisors. In other words, each advisor will show you a circle, either blue, suggesting you guess “blue,” or green, suggesting you guess “green.”

Each advisor is either perfect or defective, with equal probability.

The quality of your advisors is independent; you may therefore have two perfect advisors, two defective advisors, or one perfect and one defective advisor.

The quality of your advisors is determined at the beginning of each round and remains the same during both periods of that round. You will not be informed of the quality of your advisors in advance.

If both advisors are perfect, their advice is correct (i.e., if the selected ball is blue, both show a blue circle).

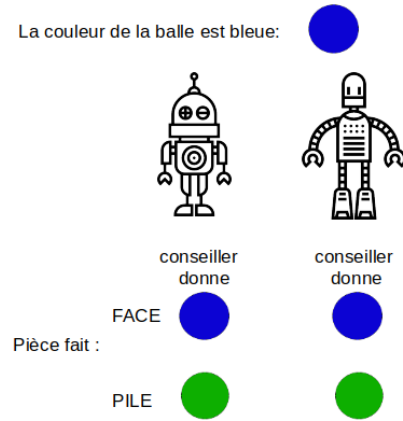
If both advisors are defective, the content of their advice depends on the outcome of a coin flip:

If the coin lands on heads, both give correct advice (i.e., show a blue circle if the selected ball is blue).

If the coin lands on tails, both give incorrect advice (i.e., show a green circle if the selected ball is blue).

The figure 9 illustrates the advice given when the actual color of the ball is blue.

Figure 9: Illustration of what advice is provided with defective advisors.



If one advisor is perfect and the other is defective, the content of their advice depends on a coin flip:

Heads: both give correct advice.

Tails: the perfect advisor gives correct advice, while the defective advisor gives incorrect advice.

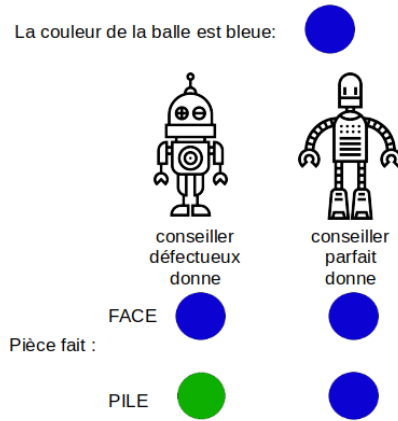
The figure 10 illustrates the advice given when the actual color of the ball is blue.

After receiving the advice, you must guess whether the ball is blue or green. You earn 500 points if your answer is correct. You will learn whether your guess is correct at the end of each round (i.e., after the end of period 2). You will also learn the quality of your advisors at the end of each round.

Second Period

Rounds 1 to 10: Your computerized advisors are the same as in the first period. You do not automatically receive their advice. You can request advice by clicking the button “Advice from left advisor”, “Advice from right

Figure 10: Illustration of what advice is provided with asymmetric quality.



advisor”, or “Advice from both advisors.” Requesting advice is costly (the cost is detailed below). Alternatively, you can guess without advice. Once you submit your answer, you will learn whether it is correct. You earn 500 points if your answer is correct. Additionally, at the end of each round, you will learn whether your advisors were perfect or defective.

Rounds 11 to 15: The second period of rounds 11 to 15 is the same as in rounds 1 to 10, except that you must propose a price you are willing to pay for advice if you request it. Your computerized advisors are the same as in the first period. You do not automatically receive their advice. You can request advice by clicking “Advice from left advisor”, “Advice from right advisor”, or “Advice from both advisors.” This request is costly (details below). You can also guess without advice. You earn 500 points if your answer is correct.

Cost of Advice in the Second Period

Rounds 1 to 10: You pay 5 points per advisor if you request advice.

Rounds 11 to 15: If you request advice, you must propose an amount

you are willing to pay:

If you request only one advisor, propose an amount between 0 and 250 points.

If you request both advisors, propose an amount between 0 and 250 points.

The experiment is designed so that it is optimal for you to propose the highest amount you are willing to pay.

A threshold price X will be drawn randomly.

If the amount you propose is greater than X , you will receive the requested advice and pay the threshold price X .

Otherwise, you will have to guess without any advice.

D Comprehension Check Questions.

D.1 Baseline treatments with one robot and belief elicitation treatments.

List of questions with correct answers (CA) verifying understanding in baseline treatments with one robot and belief elicitation treatments, translated from French:

Q1. Without any advice, what is the probability that the selected ball is green? CA: 0.6 (0.7).

Q2. Suppose you have accessed our computer program and you learn that the color of the ball drawn from the urn in period 1 of a given round is green. What is the probability that a green ball will be drawn from the urn in period 2 of the same round? CA: 0.6 (0.7).

Q3. Without any additional information what is the probability that your advisor is defective? CA: 0.5.

Q4. You have learned that in a previous round your advisor was perfect. What is the probability that your advisor is perfect in this round? CA: 0.5.

Q5. You have learned that your advisor in period 1 of a given round is perfect. What is the probability that your advisor is perfect in period 2 of the same round? CA: 1.

Q6. Suppose the color of the selected ball is green. Suppose also that you have learned that your computerized advisor is defective. What is the probability it advises “green”? CA: 0.5.

Q7. Without any additional information, what is the probability that your advisor gives you correct advice? CA: 0.75.

Q8. [True or False]: defective advisor gives correct advice with probability $\frac{1}{2}$ regardless of the color of the ball drawn from the jar? CA: TRUE

Q9. Suppose that after receiving advice in period 1 of a given round, you make a guess. How many points do you earn if this guess is correct? CA: 500.

Q10. [True or False]: Suppose that you request advice in period 2 of some round between 1 and 10 and you make a guess. If your guess is correct, you earn 495 points. CA: TRUE.

Q11. [True or False]: Suppose that you request advice in period 2 of some round between 1 and 10 and you make a guess. If your guess is false, you lose 5 points. CA: TRUE.

Table 5 reports the percentage of participants who answered each question correctly.

Table 5: Percentage of participants responding correctly.

Question number	Baseline prior 0.6	Baseline prior 0.7	Belief Elicit. prior 0.6	Belief Elicit. prior 0.7
Q1.	77.57	84.16	86.54	82.98
Q2.	75.70	86.14	71.15	80.85
Q3.	92.52	90.1	92.31	100
Q4.	66.36	72.25	55.77	63.83
Q5.	47.66	47.52	51.92	51.06
Q6.	71.96	62.38	67.31	68.02
Q7.	36.45	26.73	25	21.28
Q8.	84.11	83.17	84.62	85.11
Q9.	86.14	90.10	80.85	95.74
Q10.	85.05	84.16	78.85	80.85
Q11.	80.37	73.27	71.15	68.09

D.2 Treatments with two robots.

List of questions with correct answers (CA) verifying understanding in baseline treatments with two robots:

Q1. Without any advice, what is the probability that the selected ball will be blue? CA: 0.4 (0.3)

Q2. Suppose you manage to break into our computer program and learn that the color of the ball drawn from the pot in period 1 of round 1 is blue. What is the probability that a blue ball will be drawn from the pot in period 2 of round 1? CA: 0.4 (0.3)

Q3. What is the probability that both of your advisors are faulty in period 1 of any given round? 0.25

Q4. What is the probability that only one of your advisors will be defective in period 1 of any given round? CA: 0.5

Q5. At the end of the previous round, it turned out that one of your

advisors gave correct advice in both periods, while your other advisor gave incorrect advice in one period. What is the probability that both of your advisors will be perfect in this round? 0.25

Q6. Suppose you manage to break into our computer program and learn that one of your advisors is perfect in period 1 of any round. What is the probability that this advisor is perfect in period 2 of the same round? CA: 1

Q7. Suppose, once again, that you manage to break into our computer program and learn that one of your advisors in period 1 of a round is perfect. What is the probability that one of your advisors will be perfect in period 1 of the next round? CA: 0.5

Q8. Suppose the ball drawn from the urn is green. Suppose, furthermore, that you manage to break into our computer program and learn that both of your computerized advisors are perfect. What is the probability that they will both advise "green"? CA: 1

Q9. Suppose you manage to break into our computer program and discover that only one of your computerized advisors is perfect. What is the probability that your advisors will disagree? CA: 0.5

Q10. Suppose, once again, that you access our computer program and learn that both of your computerized advisors are faulty. What is the probability that your advisors will disagree? CA: 0

Q11. How many pieces of advice can you buy during period 2 of any given round? CA: All answers are correct [i.e. 0, 1, or 2]

Q12. Suppose that after automatically receiving advice during period 1 of a round, you guess that the selected ball is blue. Let's assume this guess is correct. It allows you to win: CA: 500 points

Q13. [True or False]: Suppose you ask for the two available pieces of advice in period 2 of a round between 1 and 10. If you guess correctly, you win 490 points. CA: TRUE

Q14. [True or False]: Suppose you ask for the two available pieces of advice in period 2 of a round between 1 and 10. If you guess correctly, you win 500 points. CA: FALSE

Q15. [True or False]: Suppose you ask for the two available pieces of advice in period 2 of a round between 1 and 10. If you guess incorrectly, you lose 10 points. CA: True

E Post experimental survey.

E.1 Explanatory Questions (mandatory) in treatments with one robot.

Q1. Please indicate how difficult it was to make decisions in this experiment using a 10-point scale, from 1 very easy to 10 very difficult

Q2. Which decisions were the most difficult?

Q3. Did you use a decision-making rule?

Q4. What probability did you assign to your advisor of being perfect after receiving their "green" advice in period 1? (as a percentage, a round number between 0 and 100)

Q5. We have noticed that in the first 15 rounds, it happened that you *[purchasing pattern]*. Why? *[See details in Table 6.]*⁴⁵

Q6. How much did you offer to pay for advice in period 2 of the last round?' *options provided: ["The amount you were really willing to pay",*

⁴⁵The number 15 should have been 10 because of a mistake in the program.

“More than you were willing to pay”, “Less than you were willing to pay”]

Table 6: Reasons behind advice purchasing strategies in baseline treatments with one robot.

Purchasing pattern	Justification	Prior 0.6 (weak)	Prior 0.7 (strong)
All	Other reasons	0.20	0.17
	You like to gather all the information relevant to your choices	0.15	0.15
	You have always been optimistic about the quality of your advisor	0.17	0.02
	You found it easy to rely on your advisor	0.12	0.02
Sometimes Buy	Other reasons	0.02	0.04
	You like to gather all the information relevant to your choices	0.13	0.11
	You wanted to be more confident in your decision	0.10	0.12
	You were curious	0.15	0.14
Sometimes No Buy	You were optimistic about the quality of your advisor	0.11	0.09
	Other reasons	0.02	0.07
	You like to make independent decisions	0.13	0.12
	You tried your luck	0.18	0.20
Nothing	You wanted to simplify your decision-making	0.07	0.05
	You were pessimistic about the quality of your advisor	0.12	0.06
	Other reasons	0.03	0.16
	You like to make independent decisions	0.24	0.08
	You tried your luck	0.03	0.08
	You wanted to simplify your decision-making	0.08	0.14
	You were pessimistic about the quality of your advisor	0.03	0.14

Share of participants who, conditional on a purchasing pattern, gave one of the suggested explanations for their behavior. Purchasing patterns are relative to round 1 to 10 and include always bought advice (*All*), bought advice sometimes, and never bought advice (*Nothing*). For the second strategy, we asked for justification on why they purchased at least once, and why they did not purchase at least once. Remark: because of a design bug, participants were told that their purchasing patterns referred to rounds 1 to 15 instead of 1 to 10.

E.2 Explanatory Questions (mandatory) in treatments with two robots.

Q1. What probability did you assign to your two advisors of being perfect after receiving ”green” advice from each of them in period 1? (as a percentage, a round number between 0 and 100)

Q2. What probability did you assign to your two advisors of being perfect after receiving "blue" advice from one of them and "green" advice from the other in period 1? (as a percentage, round number between 0 and 100)

Q3. What probability did you assign to the advisor on the left of your screen of being perfect after receiving "blue" advice from them and "green" advice from the advisor on the right of your screen? (as a percentage, a round number between 0 and 100)

Q4. What probability did you assign to the advisor on the left of your screen of being perfect after receiving "green" advice from them and "blue" advice from the advisor on the right of your screen? (as a percentage, a round number between 0 and 100)

Q5. Did you know in advance how you were going to use the advice when you asked for it in period 2? If so, tell us how.

Q6. We have noticed that in the first 15 rounds, it happened that you *[purchasing pattern]*. Why? *[See details in Table 7.]*⁴⁶

Q7. How much were you willing to pay for advice during the second period of the final rounds? *Options provided: ["The amount I was actually willing to pay"; "Less than the amount I was actually willing to pay", "More than the amount I was actually willing to pay".]*

E.3 Questions on Demographic and Other Characteristics (optional).

E.3.1 Demographics.

Could you please share with us the following information:

⁴⁶The number 15 should have been 10 because of a mistake in the program.

Table 7: Reasons behind advice purchasing strategies in baseline treatments with two robots.

Purchased (at least once)	Justification	0.6 (weak)	0.7 (strong)
No advice	Other reasons	0.06	0.11
	You like to make independent decisions	0.13	0.07
	You tried your luck	0.15	0.15
	You wanted to simplify your decision-making	0.07	0.08
	You were pessimistic about the quality of your advisor	0.08	0.10
From both robots	Other reasons	0.03	0.02
	You like to gather all the information relevant to your choices	0.25	0.20
	You wanted to be more confident in your decision	0.10	0.13
	You were curious	0.10	0.05
From one robot only	You were optimistic about the quality of your advisor	0.07	0.04
	This was sufficient information for your decision	0.24	0.25
	Other reasons	0.08	0.07
	You didn't want to face a situation where your advisors disagreed	0.19	0.17

Share of participants who, conditional on a purchasing pattern, gave one of the suggested explanations for their behavior. Purchasing patterns are relative to round 1 to 10. Remark: because of a design bug, participants were told that their purchasing patterns referred to rounds 1 to 15 instead of 1 to 10.

Q1. Your gender.

Q2. Your age.

Q3. Your native language.

Q4. What year are you in, and what is your degree?

Q5. Are you colorblind? Options: Yes, No

E.3.2 Self-assessed characteristics.

Q1. How often do you seek advice in real-life situations? Use a 10-point scale:

1 = never, 10 = always.

Q2. How willing are you to take risks in general? Use a 10-point scale: 1 = completely unwilling, 10 = completely willing.

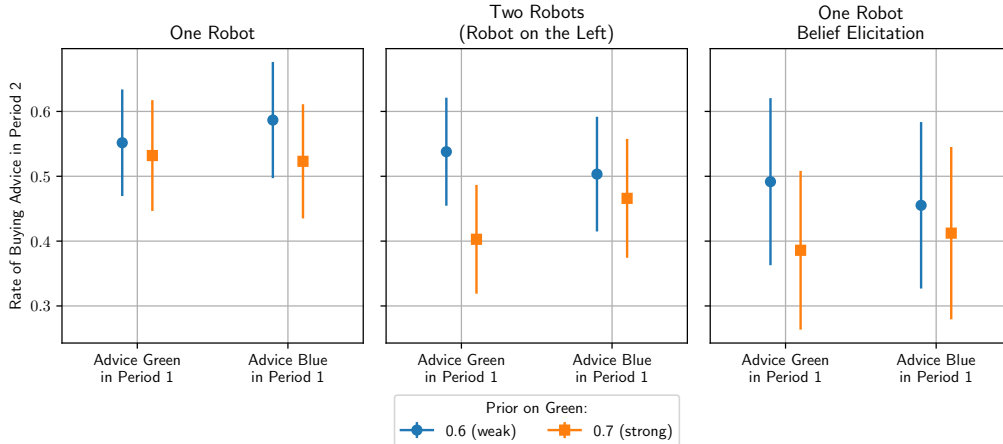
E.3.3 Skills in Bayesian Updating.

Q1. “There are two pots, A and B, each containing 4 balls. Pot A has 3 black balls and 1 white ball. Pot B has 3 white balls and 1 black ball. One ball is drawn from one of the pots. It turns out to be black. What is the probability that it was drawn from pot A?” (Answer in percent, round to a whole number between 0 and 100).

Q2. “There is a room with 100 people, half men and half women. All men are economists. Half of the women are economists. What is the probability that a randomly selected economist from this room is a woman?” (Answer in percent, round to a whole number between 0 and 100).

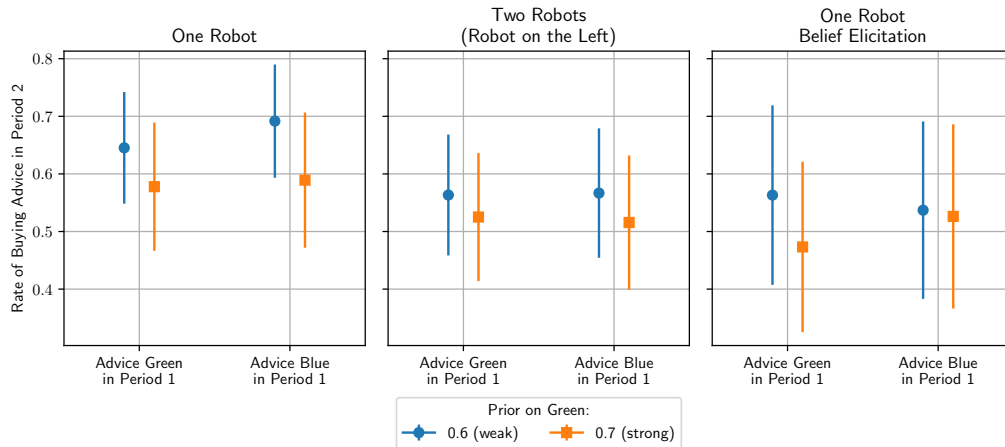
F Supplementary figures.

Figure 11: The average rate of advice purchases across rounds 6 to 10, by treatment.



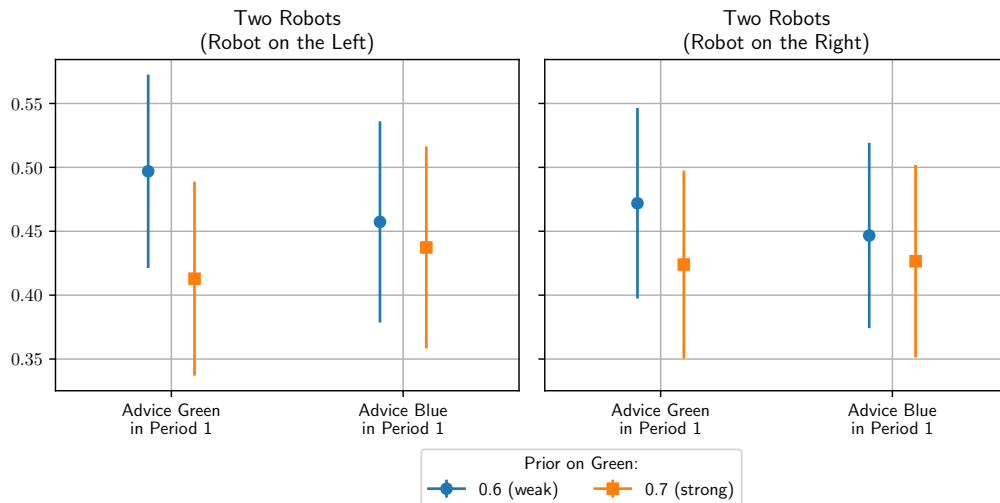
For treatments with two robots, we present the demand for advice generated by the robot on the left of the screen.

Figure 12: The average rate of advice purchases by participants who answered question Q5 correctly (in treatments with one robot) across rounds 1 to 10, by treatment.



Recall the question (in English).

Figure 13: The average rate of advice purchases across rounds 1 to 10 in treatments with two robots.



Left: advice generated by the robot on the left of the screen. Right: advice generated by the robot on the right

Table 8: Demand for advice in treatments with one robot.

	(1)	(2)	(3)	(4)	(5)	(6)
	Weak prior	Weak prior	Weak prior	Strong prior	Strong prior	Strong prior
Adv1SaidGreen	-0.0511* (0.0240)	-0.0341 (0.0224)	-0.0384 (0.0234)	-0.0208 (0.0279)	-0.00169 (0.0281)	-0.0185 (0.0286)
Experiment2023		-0.0376 (0.0666)	0.0106 (0.0857)		0.0243 (0.0692)	0.0985 (0.101)
D5AdviceSeeking		0.0446** (0.0151)	0.0417* (0.0186)		0.0335* (0.0135)	0.0280 (0.0171)
D7RiskSeeking		-0.00680 (0.0167)	0.00994 (0.0216)		0.00467 (0.0159)	0.0166 (0.0205)
Female		0 (.)	0 (.)		0 (.)	0 (.)
Male		0.176** (0.0656)	0.212* (0.0868)		0.124 (0.0702)	0.0946 (0.0966)
bac_plus			0.0466 (0.0403)			0.0699* (0.0282)
has_econ			0.00417 (0.0997)			0.0267 (0.0830)
BayesianExercise1			-0.0000876 (0.00175)			0.000919 (0.00244)
BayesianExercise2			-0.00846* (0.00366)			-0.00330 (0.00446)
Constant	0.580*** (0.0134)	0.263 (0.147)	0.124 (0.208)	0.499*** (0.0165)	0.197 (0.127)	-0.000306 (0.206)
Observations	1070	1575	1065	1010	1515	945
ParticipantFE	Yes			Yes		

Standard errors in parentheses

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Round-level OLS regression where the dependent variable is a dummy equal to one if the participant bought advice in the given round. Includes demographics or participants' fixed effects as control variables.

Table 9: Demand for advice by the robot located on the left of the screen in treatments with two robots.

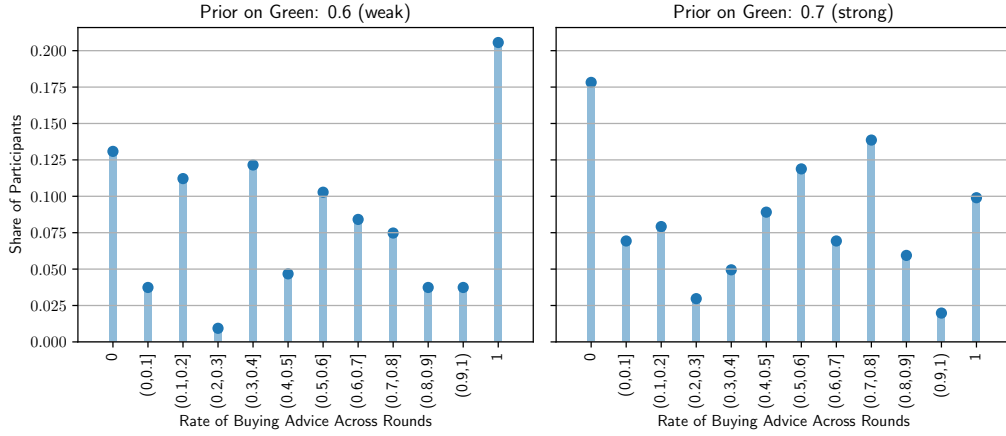
	(1)	(2)	(3)	(4)	(5)	(6)
	Weak prior	Weak prior	Weak prior	Strong prior	Strong prior	Strong prior
Adv1SaidGreen	0.0398 (0.0268)	0.0297 (0.0242)	0.0540 (0.0299)	-0.0152 (0.0283)	-0.0241 (0.0263)	-0.0258 (0.0289)
Experiment2023		-0.115 (0.0633)	-0.150 (0.0769)		0.0166 (0.0731)	0.0616 (0.102)
D5AdviceSeeking		0.0415** (0.0149)	0.0339* (0.0162)		0.0318* (0.0141)	0.0183 (0.0178)
D7RiskSeeking		-0.00448 (0.0183)	0.0141 (0.0197)		0.00700 (0.0156)	0.0145 (0.0212)
Female		0 (.)	0 (.)		0 (.)	0 (.)
Male		0.288*** (0.0623)	0.323*** (0.0811)		0.100 (0.0741)	0.00567 (0.0980)
bac_plus			0.0238 (0.0375)			0.0591 (0.0306)
has_econ			0.0354 (0.0910)			0.0265 (0.0900)
BayesianExercise1			-0.00192 (0.00176)			0.00155 (0.00212)
BayesianExercise2			0.000590 (0.00304)			-0.00536 (0.00466)
Constant	0.450*** (0.0146)	0.149 (0.145)	0.0270 (0.185)	0.425*** (0.0175)	0.150 (0.115)	0.0999 (0.213)
Observations	1070	1575	1065	1010	1515	945
ParticipantFE	Yes			Yes		

Standard errors in parentheses

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

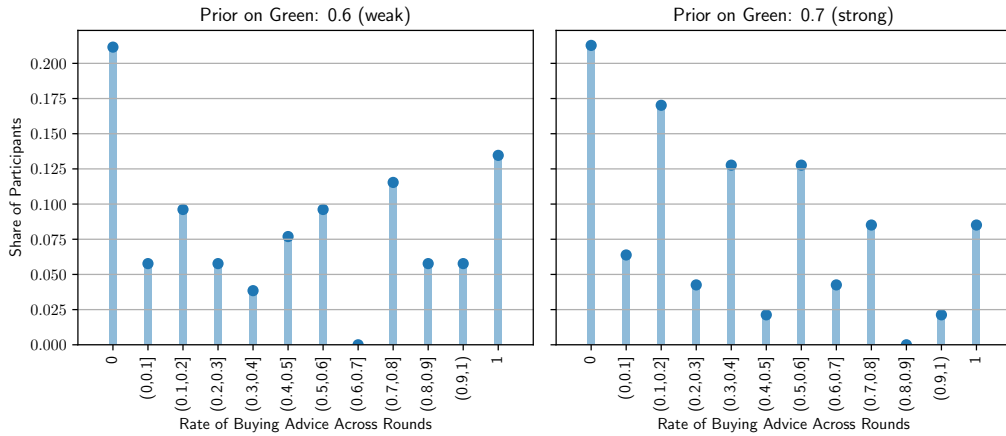
Round-level OLS regression where the dependent variable is a dummy equal to one if the participant bought advice in the given round. Includes demographics or participants' fixed effects as control variables.

Figure 14: Distribution of participants' rates of purchasing advice in baseline treatments with one robot across rounds 1 to 15.



Each bar shows the share of participants whose average rate of buying advice falls within the corresponding bin.

Figure 15: Distribution of participants' rates of purchasing advice in treatments with belief elicitation across rounds 1 to 15.



Each bar shows the share of participants whose average rate of buying advice falls within the corresponding bin.

Figure 16: Distribution of participants' posteriors on advisor's quality.

