



UNIVERSITY OF ALBERTA
FACULTY OF ARTS
Department of Economics

Working Paper No. 2024-10

**On Prior Confidence
and Belief Updating**

Kenneth Chan
UC Santa Barbara

Gary Charness
UC Santa Barbara

Chetan Dave
University of Alberta

J. Lucas Reddinger
Purdue University

December 2024

Copyright to papers in this working paper series rests with the authors and their assignees. Papers may be downloaded for personal use. Downloading of papers for any other activity may not be done without the written consent of the authors.

Short excerpts of these working papers may be quoted without explicit permission provided that full credit is given to the source.

The Department of Economics, the Institute for Public Economics, and the University of Alberta accept no responsibility for the accuracy or point of view represented in this work in progress.

On Prior Confidence and Belief Updating*

Kenneth Chan¹, Gary Charness², Chetan Dave³, J. Lucas Reddinger⁴

December 13, 2024

Abstract We investigate the difference between confidence in a belief distribution versus confidence over multiple priors using a lab experiment. Theory predicts that the average Bayesian posterior is affected by the former but is unaffected by the latter. We manipulate confidence over multiple priors by varying the time subjects view a black-and-white grid, of which the relative composition represents the prior. We find that when subjects view the grid for a longer duration, they have more confidence, under-update more, placing more weight on priors and less weight on signals when updating. Confidence within a belief distribution is varied by changing the prior beliefs; subjects are insensitive to this notion of confidence. Overall we find that confidence over multiple priors matters when it should not and confidence in prior beliefs does not matter when it should.

JEL Codes C11, C91, D83

Keywords Bayesian updating, belief updating, confidence, lab experiment

*We thank Erik Eyster, Liang Yucheng, Daniel Martin, Ryan Oprea, and Sevgi Yuksel for their helpful comments. The University of Alberta provided generous funding for this project.

¹Department of Economics, UC Santa Barbara; kenneth_chan@ucsb.edu.

²Department of Economics, UC Santa Barbara.

³Department of Economics, University of Alberta; cdave@ualberta.ca.

⁴Department of Economics, Purdue University; reddinger@ucsb.edu.

1 Introduction

Bayes' rule is the standard for updating prior beliefs upon receiving new information in classical economic models. Owing to its appeal, Bayes' rule is the objectively correct belief updating process and has microeconomic foundations in decision-making under uncertainty (Ortoleva 2024). However, the behavioral and experimental economics literature provides, perhaps unsurprisingly, substantial evidence that individuals do not behave in accordance with the Bayesian paradigm (Benjamin 2019). Examining how people incorporate new information into their existing beliefs, relative to Bayes' law, is critical to the formalization of new belief updating models to explain individual information synthesis and subsequent decision-making.¹

Our objective here is to consider how an individual's *confidence* affects the belief-updating process and the resultant updated belief. We consider two notions of confidence. First, we consider confidence in the prior belief distribution as being related to the dispersion of the belief distribution; specifically, for some distributions, greater confidence implies less dispersion in the belief distribution.² This notion of confidence offers an intuitive result in which a Bayesian agent updates more when less confident in a prior belief.³ The most well-known example is the updating of a Gaussian prior belief distribution with a signal drawn from a Gaussian distribution. When updating, an agent places more weight on the prior belief when there is less variance in the prior distribution.⁴ Our analysis examines a Bernoulli distribution, for which the greatest updating occurs given a 50% prior belief (maximal uncertainty regarding a binary state).

Second, we consider confidence in prior beliefs as represented by uncertainty over multiple prior beliefs. In this setting, an agent has a set of prior beliefs considered plausible and

¹For example, in finance and macroeconomics, Gennaioli and Shleifer's (2010) formalization of Kahneman and Tversky's (1972b) representativeness heuristic allows Bordalo, Gennaioli, and Shleifer (2018) to develop the notion of diagnostic expectations and explain credit cycles. That broad notion of diagnostic expectations allows Bianchi, Ilut, and Saijo (2024) to better reconcile macroeconomic models with data than with use of the rational expectations standard.

²Variance (or standard deviation) is a useful measure of dispersion for a model of Bayesian updating.

³Specifically, the difference between the mean updated belief and the mean prior belief is increasing in the variance of the prior belief distribution.

⁴Other well-known cases include (i) a beta prior belief distribution with signals drawn from a binomial distribution, and (ii) a Dirichlet prior distribution and a signal drawn from a multinomial distribution.

assigns probability weights to these priors.⁵ If one is only interested in average updated beliefs, a counter-intuitive result obtains: uncertainty over multiple priors does not affect how a Bayesian agent updates average beliefs as long as the average prior belief remains unchanged. To update beliefs, a Bayesian agent updates each individual prior and the beliefs over the set of priors. But if one is only interested in the average Bayesian posterior belief, a Bayesian agent simply updates beliefs using the average prior belief. Our focus is on this second notion of uncertainty over multiple priors.

Our experimental treatment induces exogenous variation in confidence over multiple priors. We then elicit confidence over multiple priors with our novel incentive-compatible mechanism, a process that can be employed to measure confidence in any stated probability by a subject in an experimental design. The experimental design thus allows us to explore the relationship between confidence across multiple priors and the belief updating process.

Our belief updating task is a re-framed version of the seminal taxi-cab problem discussed in Kahneman and Tversky (1972a). Instead of being informed of a numerical prior probability, a subject views a 10×10 grid consisting of black and white squares; the proportion of white squares represents the actual probability of success (prior to any signal). In one treatment the grid is flashed for 0.25 seconds, giving the subject only a rough sense of the proportion.⁶ This feature induces multiple priors in our subjects as they are unsure of the actual proportion of white squares. In the other treatment, the viewing time is 30 seconds, giving subjects sufficient time to count the number of squares exactly. After seeing the grid, we inform subjects that a square from the grid is randomly selected with uniform probability and we provide each subject with a signal on the color of their realized square (with a fixed, symmetric accuracy of 60% or 80%).⁷ We then elicit each subject’s updated belief regarding the color of the selected square before and after seeing the signal. We incentivize the elicitation of beliefs by paying subjects a fixed amount if their stated value is within 3 percentage points of the actual prior or the true Bayesian posterior.

As stated above, our main goal is to study how one’s *confidence* in prior beliefs af-

⁵This set of weights can be interpreted as second-order prior beliefs.

⁶Esponda, Oprea, and Yuksel (2023) use a 0.25 second viewing time to provide a noisy *signal* for updating, while we use it to induce a noisy *prior*.

⁷A *symmetric accuracy* rate r implies $P(\text{black signal} \mid \text{black square}) = P(\text{white signal} \mid \text{white square}) = r$.

fects one’s belief-updating process. To measure confidence, we develop a novel incentive-compatible elicitation method as follows. Following the elicitation of a subject’s point-estimate π of the probability of an event, we offer the subject two options: a fixed payment that is obtained when π is within 3 percentage points of the true parameter (a *subjective gamble*), and a simple gamble that obtains the same fixed payment with probability q (an *objective gamble*). Each subject ultimately provides a switching point q^* , above which the subject prefers the objective gamble and below which the subject prefers the subjective gamble. This switching point provides a measure of the subject’s confidence in their prior.⁸

To analyze our data, we construct two different measures to quantify the degree of over- and under-updating relative to the Bayesian posterior. We then examine the relationship between one’s confidence in one’s prior belief and any over- or under-updating relative to the Bayesian benchmark. We find that higher confidence is associated with more conservative updating. Our estimates of the Grether (1980) model further validate this result. We find that a subject places less weight on their prior and more weight on the signals when they view the grid for 0.25 seconds. We also find that our subjects are insensitive to uncertainty within the prior distribution. Contrary to the Bayesian prediction, we find that our subjects respond to confidence as represented by multiple priors but do not respond to confidence as represented by (less) dispersion in a belief distribution.

We contribute to the existing literature in three main ways. First, we establish the relationship between confidence over multiple priors and belief updating. While results for prior mixture models are well-established, these results have not been experimentally tested in the context of individual belief updating. Our results show that confidence over multiple priors matters even when it should not. Second, we propose and implement a simple incentive-compatible mechanism to elicit confidence in one’s belief. Finally, we present empirical evidence that subjects are surprisingly overconfident in their ability to perform Bayesian updating.

The paper is structured as follows: section 2 reviews relevant literature, section 3 describes our design in detail, section 4 presents our results, and section 5 concludes.

⁸We elicit confidence on both the prior belief *and* the updated belief. Inference on the subject’s confidence depends on how beliefs are incentivized; we discuss this in greater detail in section 3.2.

2 Related Literature

The survey in Benjamin (2019) summarizes well the belief updating literature, including studies of mistakes in probabilistic reasoning. Laboratory experiments have documented several non-Bayesian updating patterns. Some of the more prominent biases include base-rate neglect (Kahneman and Tversky 1972b; Esponda, Vespa, and Yuksel 2024), conservatism bias (Phillips and Edwards 1966; Augenblick, Lazarus, and Thaler 2024), correlation neglect (Eyster and Weizsacker 2016; Enke and Zimmermann 2019), confirmation bias (Nickerson 1998; Charness and Dave 2017), and motivated beliefs (Eil and Rao 2011; Coutts 2019; Thaler 2021; Charness, Oprea, and Yuksel 2021; Möbius et al. 2022). Agranov and Reshidi (2023) also show that in a sequential updating problem with multiple signals, the order in which the signals are presented to the subjects affects the final belief profile.⁹ Beyond the lab, a growing literature documents non-Bayesian belief updating in finance (Bordalo, Gennaioli, and Shleifer 2018; Augenblick and Lazarus 2023), job search (Conlon et al. 2018; Brown and Chan 2024), and sports-betting markets (Augenblick and Rabin 2021; Augenblick, Lazarus, and Thaler 2024).

To accommodate these belief-updating biases, behavioral economists have constructed numerous non-Bayesian updating models (Grether 1980; Rabin and Schrag 1999; Epstein, Noor, and Sandroni 2010; Cripps 2018; Woodford 2020; Levy, Barreda, and Razin 2022). Among these, the most popular model for empirical analyses is that of Grether (1980), which accommodates several updating biases and can be estimated with linear regression. We apply this model to our environment in section 3.6 and present the results in section 4.3.

When there is a set of prior beliefs, the ambiguity literature has proposed a few non-Bayesian updating models. One is termed *Full Bayesian Updating* (Pires 2002; Gilboa and Marinacci 2016) in which an agent ignores the second-order belief over the priors and simply updates each prior beliefs with Bayes' rule. Another is *Maximum Likelihood Updating* (Gilboa and Schmeidler 1993), in which an agent selects the most likely prior after observing the signal. In contrast to these extreme models of belief updating, Kovach (2024) introduces *Partial Bayesian updating*, a generalized version of *Full Bayesian Updating* and *Maximum*

⁹This is a violation of the divisibility property in belief updating (Cripps 2018).

Likelihood Updating, in which an agent considers a subset of priors deemed plausible, and updates these priors with Bayes’ rule. Klibanoff and Hanany (2007) introduces a model of updating that retains dynamic consistency in preferences, and Ortoleva (2012) introduces the *Hypothesis Testing model* in which an agent may switch between prior beliefs upon receiving a surprising signal.

While deviations from Bayes’ rule comprise numerous systematic patterns, we are primarily interested in examining over- and under-updating relative to the Bayesian benchmark. Benjamin (2019) notes that empirical biases in belief updating generally take two forms: base-rate neglect, a phenomenon where people underweight the prior belief, and conservatism, a tendency to under-react to new information. Base-rate neglect results in over-updating relative to the Bayesian benchmark, while conservatism results in under-updating relative to the Bayesian benchmark. The dominating bias then gives the direction of the overall effect.

Our work is closely related to two recent studies that examine the causes of over- and under-updating. Augenblick, Lazarus, and Thaler (2024) show that individuals tend to under-infer from signals when test reliability is high (greater than 60% accuracy) and over-infer when reliability is low (below 60% accuracy). Ba, Bohren, and Imas (2022) show that when the number of possible outcomes or states increases, subjects tend to over-update when there are more than two states, while subjects under-update when there are two states. In contrast to these studies, we focus on prior confidence as an explanation for over- or under-updating. We also diverge from these papers with our use of Kahneman and Tversky’s (1972a) famous taxi-cab problem, while the aforementioned studies both use some variant of the balls-and-urn (or book-bag-and-poker-chip) design (Phillips and Edwards 1966) commonly used in belief updating studies (Benjamin 2019).¹⁰ We note that while the experiments of these two papers were conducted online, our study was conducted in-person with a university subject pool.¹¹

Finally, the aforementioned papers employ a noisy cognition model (Woodford 2020; Enke and Graeber 2023) to explain over- and under-updating. In particular, Enke and Graeber

¹⁰Kahneman and Tversky’s (1972a) task is also recently used by Esponda, Vespa, and Yuksel (2024) and Agranov and Reshidi (2023).

¹¹Charness et al. (2023) discuss the merits of in-person laboratory experiments.

(2023) argue that when humans are cognitively uncertain they tend toward a cognitive default. Our results cannot be explained by the standard model of noisy cognition alone.

3 Experimental Design and Hypotheses

Our design investigates whether confidence in one’s prior influences how one updates beliefs in tasks with two components: a belief-updating exercise and a confidence-in-stated-belief elicitation. In each task, subjects provide four probability reports in the following order: (1) their prior belief, (2) their confidence in their prior belief, (3) their updated belief after receiving a signal, and (4) their confidence in their updated belief. In this section we first describe the belief-updating exercise, describe the elicitation method, and conclude with our hypotheses.

3.1 Belief Updating

Our belief updating task is a version of Kahneman and Tversky’s (1972a) taxi-cab problem. To make the task less abstract, we follow the framing of Esponda, Vespa, and Yuksel (2024) in which subjects are asked to play the role of a manager who is evaluating whether a project is a success or failure. Subjects are told that each square in a 10×10 grid represents a project, with white representing a success and black a failure (see figure 1). We induce prior beliefs by showing a grid realization to the subject. The proportion of successes to failures varies across tasks (grids): the number of successful projects is one of 0, 20, 40, 50, 70, or 90. We included a task with 0 successful projects, so correct responses are straightforward, to evaluate a subject’s comprehension of the environment.

We use a within-subject design to study how an individual updates beliefs given various levels of confidence. To achieve this, we have two treatment conditions that subjects complete in order. Subjects first complete tasks in a *Low Confidence* treatment that flashes each grid for 0.25 seconds. Subjects then complete tasks in a *High Confidence* treatment that displays each grid for 30 seconds (subjects are allowed to proceed after five seconds).¹²

¹²We chose 30 seconds because it allows a subject to tally a grid twice. In our experiment, most subjects proceeded without using all 30 seconds to view the grid.

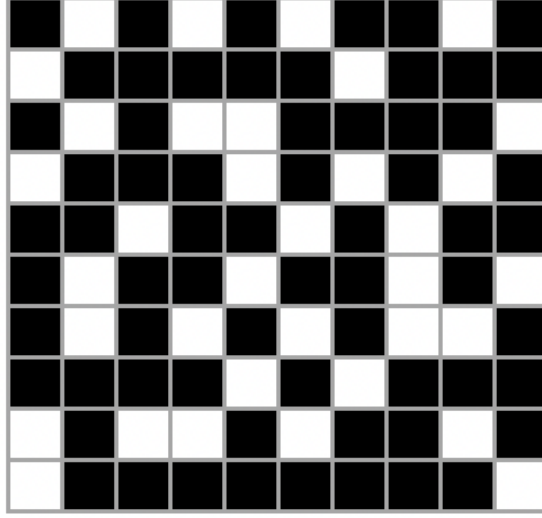


Figure 1: An example grid displayed in an experimental task

The difference in the duration of the display of the grid induces variation in one’s confidence in one’s prior, hence the treatment names.¹³ We postulate that confidence is lower in the Low Confidence treatment relative to the High Confidence treatment, the latter giving subjects sufficient time to tally squares. All subjects complete the Low Confidence treatment before the High Confidence treatment to minimize the formation of beliefs regarding our experimental parameters.

After we show a subject a realized grid, we tell the subject that a project (*i.e.*, a square from the grid) has been randomly selected with equal probability. The subject is then asked to state the probability that the selected project is a success. We refer to this reported probability as the subject’s *prior* belief. We pay \$3 if this stated probability is within 3 percentage points of the actual probability and zero otherwise (Dufwenberg and Gneezy 2000; Charness and Dufwenberg 2006). We choose this simple incentive-compatible mechanism to simplify the confidence-in-beliefs elicitation method which we discuss in section 3.2. Moreover, studies have shown that complex belief-elicitation methods like the binarized scoring rule (Hossain and Okui 2013) can bias subjects’ reported beliefs (Danz, Vesterlund, and Wilson 2022), despite being incentive compatible in theory.

¹³Another possible way to induce different degrees of confidence in the prior belief is to inform the subject that the prior is either 75% or 25% (Liang 2024). We induce the multiple prior in this manner because it avoids describing the prior in a compound manner, which is known to pose challenges for subjects in the context of risk (Halevy 2007).

Next, to permit belief updating, we tell subjects that they will be shown a computer test result to help with their evaluation of the selected project. Subjects are assigned randomly to a treatment with either a 60% or an 80% test reliability.¹⁴ The test reliability indicates the proportion of times when the computer test correctly predicts whether the project is a success or a failure with symmetric false positive and false negative rates. For example, when the reliability is 80%, if the selected project is a success (failure) the test result will be positive (negative) with 80% chance and negative (positive) with 20% chance.

After informing a subject of the result of the test (positive or negative), we ask the subject to report the probability of the selected project being a success, which we refer to as the subject’s *update*. We pay each subject \$3 for reporting a probability within three percentage points of the Bayesian posterior. When the prior is non-degenerate, we elicit the subject’s updated belief for both positive and negative test results. We have one task where the prior is degenerate at 0; in this case, we only elicit the update when the test result is positive. Each subject thus completes eleven tasks in each of the two treatments.

3.2 Confidence Elicitation

We elicit confidence after each belief elicitation—that is, after a subject states their prior belief and also following the updated belief.¹⁵ To elicit a confidence-in-stated-belief, we offer the subject a choice between either (i) a subjective gamble, which obtains \$3 if their stated belief is within 3 percentage points of the true value, and (ii) an objective gamble that obtains \$3 with probability q and nothing otherwise.

We seek to identify the point q^* at which a subject is indifferent between the subjective and objective gambles by varying q . We then interpret q^* as the confidence level for an expected utility maximizer. For example, if the subject is 70% confident that their stated prior belief is within 3 percentage points of the actual prior (or, following the signal, the Bayesian posterior), the subject expects to be paid 70% of the time. Thus the subject should only accept lotteries that pay \$3 at least 70% of the time. By reporting 70% in the confidence elicitation, a subject ensures that they will only receive lotteries that obtain \$3 at least 70%

¹⁴We held signal reliability constant for each subject to minimize potential confusion.

¹⁵This captures confidence over multiple priors.

of the time.

We elicit the switching point q^* using a BDM mechanism (Becker, DeGroot, and Marschak 1964) following Healy’s (2020) instructions; we explain the BDM with a multiple price list (Holt and Laury 2002) to improve comprehension of the mechanism. Because subject comprehension of the BDM mechanism may be a concern (see Cason and Plott 2014), we include a simple question with no successful projects where the theoretically optimal response is obvious. In this question, the grid that is shown contains all black squares, and subjects should have maximal confidence in their belief. Our data on the degenerate prior case demonstrates that our subjects have an excellent understanding of the mechanism and the tasks.

Our confidence elicitation is incentive-compatible for standard expected utility maximizers but also for some non-expected utility maximizers. Let $U(m, q)$ be the utility of a simple gamble that obtains monetary payoff m with probability q .¹⁶ For our confidence elicitation method to be incentive-compatible, we require the three following assumptions. First, the agent only has preferences over monetary payoffs and the probabilities of winning. For example, we assume that the agent does not have a preference to earn money for the sake of merit-based self-signaling (*e.g.*, Bénabou and Tirole 2006), wherein the agent will report higher confidence. Second, we assume utility is strictly increasing in q . This assumption allows us to interpret q^* as the confidence, and the strict condition ensures that the agent is not indifferent between two different values of q^* . Our third and most substantial assumption is that the agent is ambiguity neutral (Ellsberg 1961). The subjective gamble over one’s stated confidence in the belief elicitation task relies on subjective probabilities, whereas the alternative gamble is based on an objective probability of winning \$3. When an agent is ambiguity averse (seeking), the agent will report a lower (higher) q^* than the agent’s true confidence.

In summary, our confidence elicitation method consists of two stages that can be used for belief elicitation in any setting with a verifiable truth. In the first stage, each subject states their belief (*i.e.*, a probability report) or makes a guess that is incentivized. In the second stage, each subject is offered a choice of sticking to their reported probability or a payoff-

¹⁶In the case of expected utility theory, $U(m, q) = q \cdot u(m) + (1 - q) \cdot u(0)$ for some function u that expresses the utility of each respective payoff.

equivalent simple gamble that obtains q of the time. Instead of using alternatives such as a quadratic or a binary scoring rule (see Brier 1950; Hossain and Okui 2013), we choose our particular incentive-compatible mechanism for its intuitive appeal and ease of comprehension. We simply propose paying a subject if their guess is within a reasonably small interval of a correct value because a subject can readily recognize that truthful revelation of their confidence is in their best interest (Danz, Vesterlund, and Wilson 2022).

3.3 Session Summary

After we presented the instructions to subjects, they responded to three comprehension check questions and finally three practice tasks (with Low Confidence treatment) to familiarize themselves with the interface and the task. Later, when subjects begin the High Confidence treatment, they complete one practice question before the paid tasks.

For the main experiment, each subject completes twenty-two tasks; the eleven in each treatment are presented in a random order. Each of these tasks involves four probability reports: a prior belief, an updated belief, and a confidence elicitation for each of these prior and updated beliefs. Subjects simply type an integer between 0 and 100 in a text box to represent a percentage. For each probability report, we randomly select one of the twenty-two tasks for payment. Because these four draws are independent, the probability reports that we pay generally correspond to different tasks.¹⁷ Following the updating tasks, subjects complete a demographic questionnaire.

3.4 Theory: Confidence and Bayesian Updating

We provide a model to formalize our notion of confidence, reflecting uncertainty over multiple priors that relates to our experiment; appendix A offers a more general model. Consider a binary state space $\Omega := \{S, F\}$ that corresponds to the randomly-selected project being either a success or a failure. Consider an agent who is uncertain about their prior belief about the proportion of successful projects, $\pi_0 := P(S)$. Let us assume that the agent considers a finite

¹⁷We did this to avoid hedging behavior within each task as it is possible for subjects to hedge in reporting the prior and their updated beliefs. For example, if subjects believe that there are two possible values for the prior they can report one of the priors and report the updated belief derived from the other prior belief.

set $\Pi_0 \subset [0, 1]$ of N possible priors $\pi_{0,i}$. The agent has a belief distribution over Π_0 , assigning subjective probability mass $k_{0,i}$ to prior $\pi_{0,i}$ for all $i \in \{1, \dots, N\}$. Ultimately the agent's prior belief π_0 about the project being a success is a weighted average of prior beliefs

$$\pi_0 := \sum_{i=1}^N k_{0,i} \pi_{0,i}.$$

The agent receives an informative signal $\sigma \in \{\sigma_P, \sigma_N\}$. Define the conditional probability of observing signal s conditioned on the project being a success or a failure as $P(\sigma|S)$ and $P(\sigma|F)$. The posterior belief is also a weighted average of posteriors. To obtain this posterior, the agent must update each prior $\pi_{0,i}$ to a posterior $\pi_{1,i}$ *as well as* update each respective prior weight $k_{0,i}$ to a posterior weight $k_{1,i}$. Intuitively, the signal provides information about both the project type and the relative likelihood of the priors. However, for a Bayesian agent, updating with the average prior belief is equivalent to the aforementioned procedure.

Proposition 1. *The average Bayesian posterior after observing signal σ is*

$$\pi_1^{\text{Bayes}}(\sigma) := P(S|\sigma) = \sum_{i=1}^N k_{1,i}^{\text{Bayes}}(\sigma) \pi_{1,i}^{\text{Bayes}}(\sigma),$$

which is the same as updating with the average prior belief.

Proof. Let us update each prior $\pi_{0,i}$ and its respective weight $k_{0,i}$ individually using Bayes' rule to obtain

$$\pi_{1,i}^{\text{Bayes}}(\sigma) := P_i(S|\sigma) = \frac{\pi_{0,i}P(\sigma|S)}{P_i(\sigma)} \quad \text{and} \quad k_{1,i}^{\text{Bayes}}(\sigma) := \frac{k_{0,i}P_i(\sigma)}{P(\sigma)},$$

where $P(\sigma) = \pi_0 P(\sigma|S) + (1 - \pi_0) P(\sigma|F)$ is the probability of observing signal σ given the mixture prior belief distribution or the average prior belief π_0 , while $P_i(\sigma) = \pi_{0,i} P(\sigma|S) + (1 - \pi_{0,i}) P(\sigma|F)$ is the probability of observing signal σ given prior $\pi_{0,i}$. We now show that updating the beliefs of each prior and the mixture distribution over these is equivalent to simply updating the average prior belief.

$$\pi_1^{\text{Bayes}}(\sigma) = \frac{\pi_0 P(\sigma|S)}{P(\sigma)}$$

$$\begin{aligned}
&= \frac{\sum_{i=1}^N k_{0,i} \pi_{0,i} P(\sigma|S)}{P(\sigma)} \\
&= \frac{\sum_{i=1}^N k_{0,i} P_i(\sigma) \pi_{0,i} P(\sigma|S) / P_i(\sigma)}{P(\sigma)} \\
&= \frac{\sum_{i=1}^N k_{0,i} P_i(\sigma) \pi_{1,i}^{\text{Bayes}}}{P(\sigma)} \\
&= \sum_{i=1}^N k_{1,i}^{\text{Bayes}}(\sigma) \pi_{1,i}^{\text{Bayes}}(\sigma).
\end{aligned}$$

□

This result is slightly counter-intuitive because the updating process is comprised of two countervailing effects. First is the convexity or concavity of the Bayesian updating function of each individual prior. In our task, when a positive (negative) signal is observed, the function is concave (convex). If the weights over the prior beliefs are held fixed, less updating occurs compared to updating given an average prior belief. Second, after a signal is observed, the Bayesian agent will place more weight on those prior beliefs which have been revealed to be more likely, in turn resulting in a greater degree of updating. These two countervailing effects cancel each other out exactly, providing the result in proposition 1.

3.5 A Measure of Over-updating

Our study's primary outcome of interest is over- or under-updating. We follow Ba, Bohren, and Imas's (2022) approach and define two measures of excess updating relative to the Bayesian benchmark.¹⁸ Let π_0 be a subject's reported prior beliefs, π_1 be the subject's reported updated beliefs, and π_1^{Bayes} be the subject's corresponding Bayesian posterior belief given the prior they previously reported. Note that a subject may not have an accurate perception of the actual prior, especially in the Low Confidence treatment. Our first measure

¹⁸Ba, Bohren, and Imas (2022) refers to these as *over-* and *under-reaction*. We prefer *over-* and *under-update* and reserve *reaction* for inference or consequent action. Over-updating can arise from under-weighting of the prior or over-reacting to a signal; these two updating biases are modeled differently under the Grether (1980) model, each having a different implication.

captures the magnitude of the over-update in percentage points, with

$$\text{over-update} := \begin{cases} \pi_1 - \pi_1^{\text{Bayes}} & \text{if signal is positive, and} \\ \pi_1^{\text{Bayes}} - \pi_1 & \text{if signal is negative.} \end{cases} \quad (1)$$

When a positive signal is observed, we expect the subject to update upwards, that is, the agent over-updates if the subject's updated belief π_1 is greater than the subject's Bayesian belief π_1^{Bayes} , resulting in $\text{over-update} > 0$. Similarly, in the event of a negative signal, if the subject's updated belief π_1 is less than the subject's Bayesian belief π_1^{Bayes} , this again implies $\text{over-update} > 0$. Thus over-update is positive when the subject's belief moves too far in the correct direction and negative when it moves too little in the correct direction. This measure is also negative if a subject updates in the wrong direction.

We next construct a metric to capture the magnitude of an over-update relative to the magnitude of a Bayesian update. Notice that an over-update of 10 percentage points has different implications if the magnitude of the Bayesian update is 5 percentage points compared to 40 percentage points. The former represents a relatively large updating error, while the latter represents a relatively small error. To account for this we normalize the over-update variable by the magnitude of the Bayesian update:

$$\text{over-update-ratio} := \frac{\text{over-update}}{\left| \pi_0 - \pi_1^{\text{Bayes}} \right|}. \quad (2)$$

An update in the wrong direction results in a strictly negative over-update measure and an over-update-ratio strictly less than one. We address this empirically with a robustness check that excludes observations with updates in the wrong direction (about 27% of our observations). We find similar results, so we simply present results from the full sample.

3.6 Estimating the Grether (1980) Model

The Grether (1980) model is a generalized version of Bayes' rule that accommodates several updating biases. In our setting,

$$\pi_1(\sigma) = \frac{P(\sigma|S)^\alpha \pi_0^\beta}{P(\sigma|S)^\alpha \pi_0^\beta + P(\sigma|F)^\alpha (1 - \pi_0)^\beta},$$

where $\alpha \geq 0$ is the weight that the agent places on the signal and $\beta \geq 0$ is the weight on the prior. Thus $\alpha < 1$ represents an under-reaction to a signal and $\alpha > 1$ an over-reaction. Meanwhile $\beta < 1$ represents base-rate neglect and $\beta > 1$ over-weighting the prior.

We can rewrite this model using an odds ratio as

$$\frac{\pi_1(\sigma)}{1 - \pi_1(\sigma)} = \left(\frac{P(\sigma|S)}{P(\sigma|F)} \right)^\alpha \left(\frac{\pi_0}{1 - \pi_0} \right)^\beta.$$

This formulation is widely used in the analysis of experimental data because in logarithms one obtains a model whose parameters can be estimated with a linear regression:

$$\ln \frac{\pi_1(\sigma)}{1 - \pi_1(\sigma)} = \alpha \ln \frac{P(\sigma|S)}{P(\sigma|F)} + \beta \ln \frac{\pi_0}{1 - \pi_0}.$$

Thus, the model permits estimation of the weights that subjects place on prior information and on information conveyed by the signal.

In our experiment, we are interested in the weights in the High and Low treatments. To estimate the difference we specify the following regression:

$$\ln \frac{\pi_1(\sigma)}{1 - \pi_1(\sigma)} = (\gamma_0 + \gamma_2 \mathbb{1}_H) \ln \frac{P(\sigma|S)}{P(\sigma|F)} + (\gamma_1 + \gamma_3 \mathbb{1}_H) \ln \frac{\pi_0}{1 - \pi_0} + \varepsilon \quad (3)$$

where $\mathbb{1}_H$ is an indicator variable for the High Confidence treatment. We use this specification to estimate α_T and β_T for each treatment $T \in \{L, H\}$:

$$\alpha_L = \gamma_0, \quad \beta_L = \gamma_1, \quad \alpha_H = \gamma_0 + \gamma_2, \quad \text{and} \quad \beta_H = \gamma_1 + \gamma_3. \quad (4)$$

While previous studies provide subjects with explicit numerical priors, our subjects must report their prior beliefs after seeing the grid, thus measurement error may be a concern. To

address this, we use the actual prior as an instrument for the prior log-odds ratio.¹⁹

3.7 Hypotheses

Our primary research question asks how confidence over multiple priors affects how one updates their beliefs. While standard economic theory predicts that confidence over multiple priors should not matter, we hypothesize otherwise.

Hypothesis 1. *The over-update and over-update-ratio measures are larger for the Low Confidence treatment relative to the High Confidence treatment.*

Hypothesis 2. *Greater reported confidence (over multiple priors) is associated with lower over-update and over-update-ratio measures.*

Hypothesis 3. *Greater weight is placed on prior beliefs in the High Confidence treatment relative to the Low Confidence treatment. Specifically, in our Grether (1980) regression equations (3) and (4), $\beta_H > \beta_L$.*

4 Results

We recruited 118 subjects for sessions held in April 2024 at the University of California, Santa Barbara.²⁰ We paid each subject a \$7 show-up fee. The session duration was one hour, with \$17.70 in average earnings. Of the 118 subjects, 70 (59.3%) identified as female and 21 years was the average age of all subjects. At the beginning of each session, we presented detailed instructions using slides, subjects were encouraged to ask questions as the instructions were being read, and they retained hard-copy instructions for reference.

4.1 Experimental Design Validation

We first verify that the subjects understood the experiment and that our method of displaying the grid for a short duration achieved its intended manipulation.

¹⁹The instrumental variable regression was also used by Möbius et al. (2022). Their subjects faced cognitive tasks of varying difficulty, regarding which subjects then reported beliefs about their performance. The authors use the task difficulty as an instrumental variable for their log ratio of prior beliefs.

²⁰The UCSB Human Subjects Committee exempted our protocol 69-23-0349. All subjects gave informed consent. We implemented the experiment in Qualtrics and recruited subjects with ORSEE (Greiner 2015).

4.1.1 Comprehension of Tasks

Ensuring subject comprehension of our environment is important given our novel confidence-in-stated beliefs elicitation and general concerns about the comprehension of the BDM (Cason and Plott 2014). To test subjects’ understanding of our confidence elicitation procedure, we include a task with a degenerate prior (*i.e.*, with zero successful projects) as represented by a grid with only black squares. One can readily discern that all squares are black, even when the grid is only displayed for a quarter of a second, as in the Low Confidence treatment. In this task, 115 subjects (97.5%) reported the prior correctly, suggesting that they understood the task of reporting the prior.²¹ Next, if a subject understands the confidence elicitation method, they should report confidence of 100% when the grid contains all black squares, assuming the subject is indeed confident. Fourteen subjects (11.9%) failed to provide such reports.²²

In total, each subject completed four comprehension checks, reporting a prior and confidence for both Low and High treatments. Thirteen subjects made one error in total, four subjects made two, while the remaining 101 subjects (85.6%) made no errors. Our subjects thus generally demonstrate a good understanding of the confidence elicitation task. Appendix C.1 offers qualitatively similar results for the subsample of subjects who passed these comprehension checks. We proceed to analyze and present results using the full sample.

Upon completion of the tasks, we ask subjects to state their agreement (on a Likert scale) with the following: “The more confident that I am in my initial belief about the proportion of successful projects, the less I should respond to the outcome of the computer test result.”²³ Our subjects largely agreed with our overarching hypothesis: 97 of 118 (82.2%) subjects agree that when they have higher confidence, they should update less (see figure 2).

²¹That is, three out of 118 total subjects did not report 3% or less when asked about the probability of a successful project selected from the grid. Two subjects erred only in the Low Confidence treatment, while one subject erred in both treatments.

²²None of these fourteen subjects provided an incorrect prior belief. Four erred only in the Low Confidence treatment, seven only in High Confidence, and three subjects erred in both treatments. Among the eleven subjects who erred in only one treatment, seven were 99% confident and ten were at least 95% confident.

²³This is in a similar spirit to DellaVigna and Pope (2018), who solicited predictions from academic experts on their experiment.

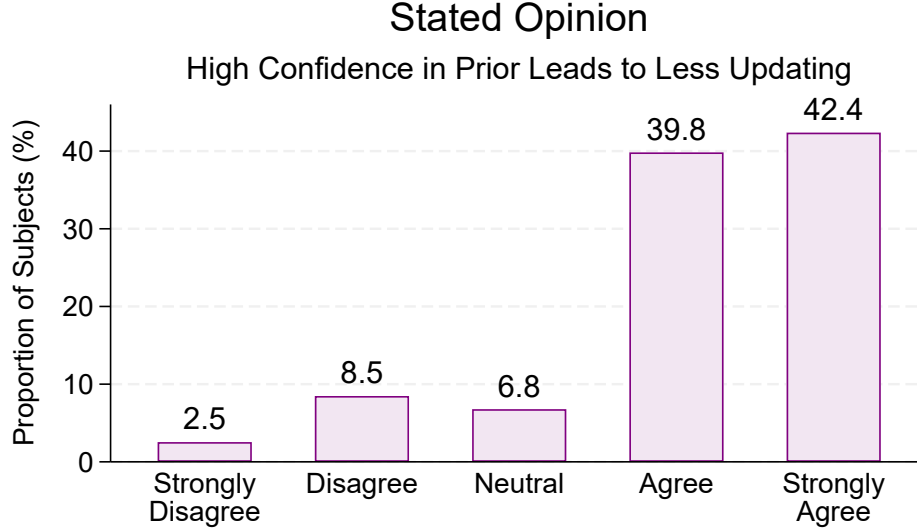


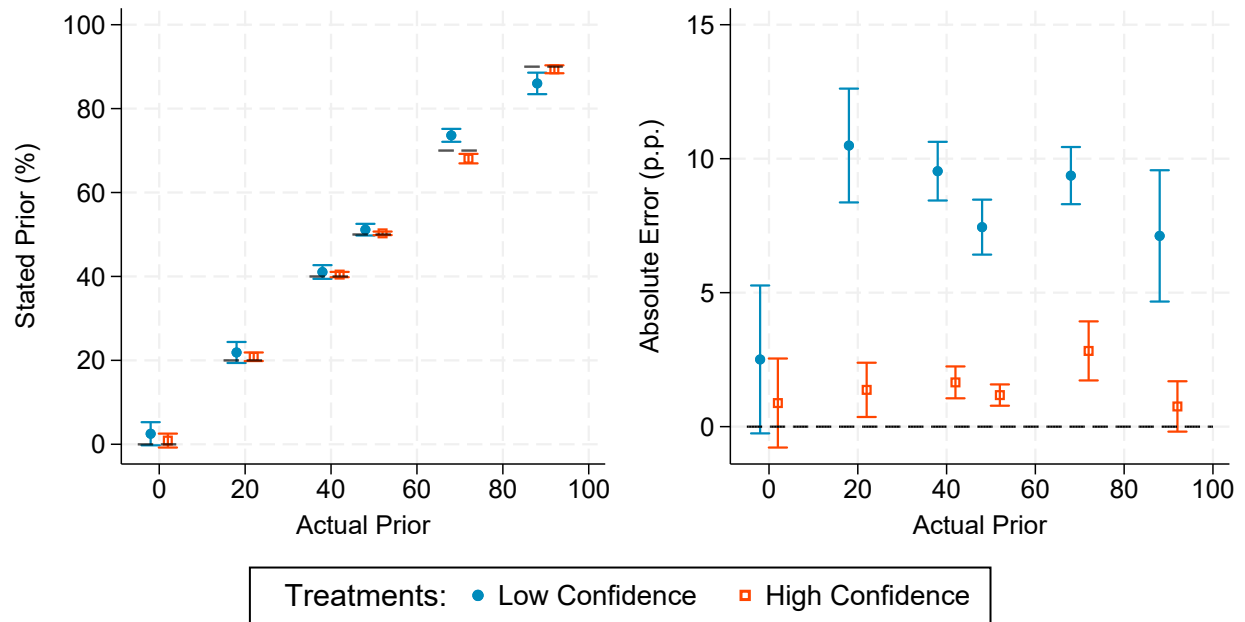
Figure 2: Self-reported relationship between confidence in prior and updating

4.1.2 Accuracy of Priors

Our design induces a prior belief by showing subjects a grid for a limited amount of time. As a result, subjects may have incorrect prior beliefs. We now present results showing the accuracy of subjects' prior beliefs across the two treatments.

As shown in the left panel of figure 3, subjects on average are surprisingly accurate in stating the actual prior even in the Low Confidence treatment. In the right panel, we plot the absolute deviation in subjects' stated priors from the actual prior. As expected, in the Low Confidence treatment, we see that subjects have highly inaccurate prior beliefs. In the High Confidence treatment, errors may result from the miscounting of squares or from the incentives not requiring perfect precision in the stated prior belief. Regardless these errors should be smaller in the High Confidence treatment, which is exactly what the data show. Next, displaying the grid for only 0.25 seconds in the Low treatment gives subjects on average a noisier perception of the actual prior relative to the High treatment; we see wider confidence intervals for the Low Confidence treatment in the data, as shown in the right panel of figure 3.

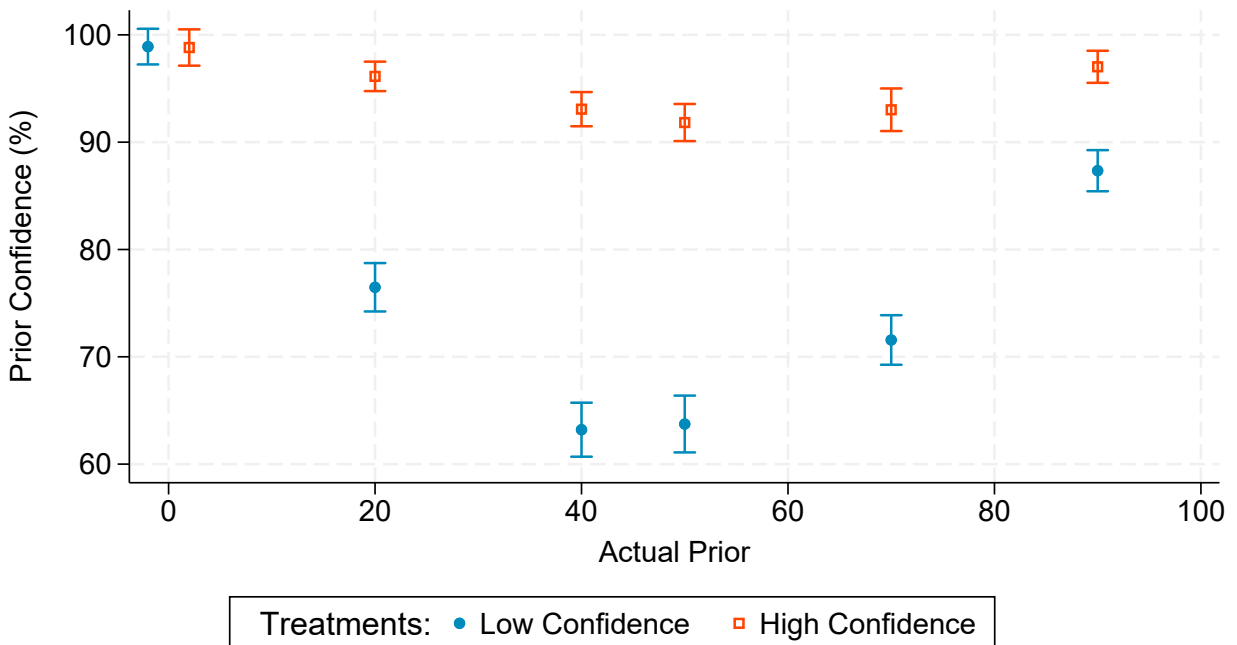
Accuracy of Stated Prior



Notes: Mean value by treatment with 95% confidence intervals as bars. Jitter added to the x-axis. Absolute error is $|\text{stated_prior} - \text{actual_prior}|$.

Figure 3: Stated prior accuracy by treatment

Prior Confidence



Notes: Mean value by treatment with 95% confidence intervals as bars. Jitter added at x=0.

Figure 4: Confidence over Multiple Priors by treatment

4.1.3 Confidence over Multiple Priors Across Treatments

Our design varies the display time of grids to induce different confidence levels in a subject’s prior. We find that subjects’ confidence in their perceived (and hence stated) priors vary in an expected direction by treatment, shown in figure 4, which demonstrates that subjects are less confident in the Low Confidence treatment compared to the High Confidence treatment.

Result 0. *Subjects in the Low Confidence treatment expressed lower confidence in their prior compared to the High Confidence treatment.*

4.2 Over-updating

We plot both of our over-updating measures (over-update and over-update-ratio) in figure 5.²⁴ The first result is that subjects under-update relative to the Bayesian benchmark. This is consistent with the results from Ba, Bohren, and Imas (2022), who found under-updating relative to the Bayesian benchmark in updating tasks with only two states.²⁵

Result 1. *Subjects under-update relative to the Bayesian benchmark in both treatments.*

Figure 5 presents how our over-update and over-update-ratio measures vary by treatment, actual prior, and signal accuracy. Overall the data indicate that the average over-update and over-update-ratio measures are larger in the Low Confidence treatment for every actual prior value. The middle panel for the 80% signal accuracy shows a similar pattern, but with a more pronounced treatment effect.²⁶

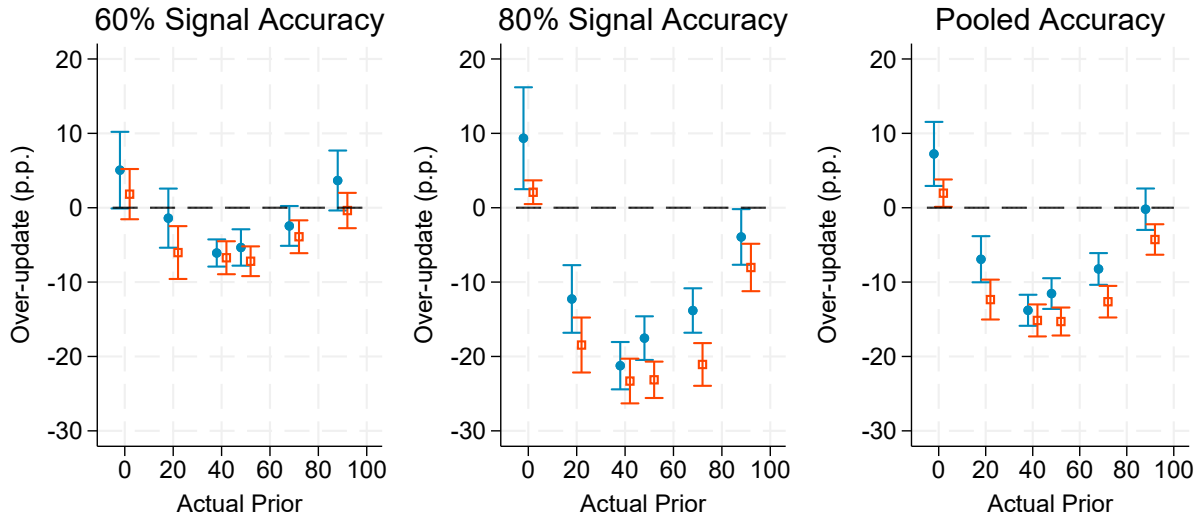
Table 1 presents linear regressions for over-update and over-update-ratio, including specifications both with and without subject fixed-effects. The independent variables include an indicator for the High Confidence treatment and a constant term. On average we see

²⁴We drop the task in which the prior is degenerate because it involves no updating and our over-update-ratio is not well-defined.

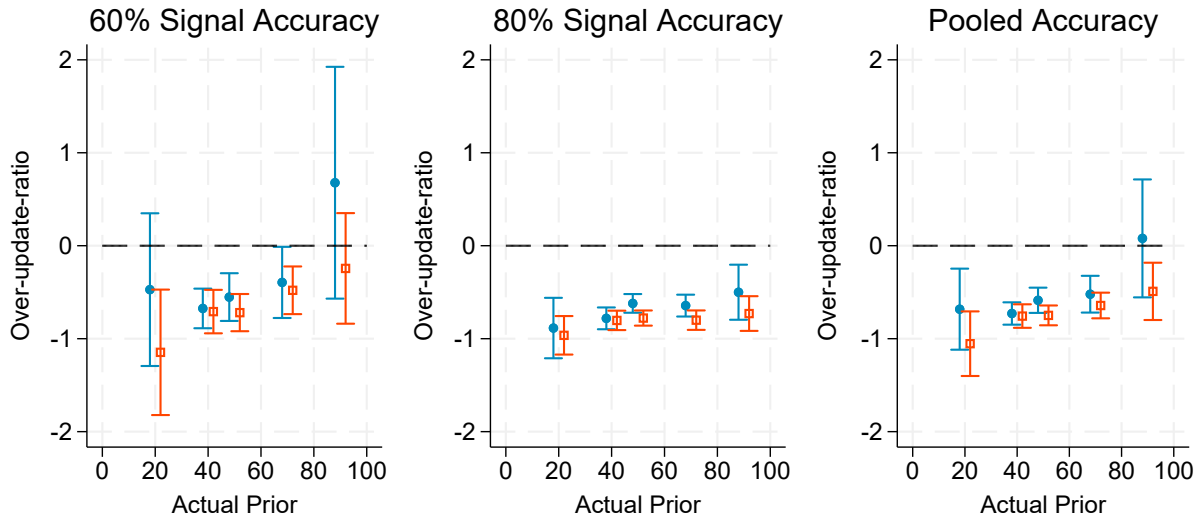
²⁵Ba, Bohren, and Imas (2022) experiment was a “balls and urn” framing, while ours is the Kahneman and Tversky (1972a) belief updating problem. This validates their finding across a different framing of belief-updating tasks.

²⁶We also plot the CDF of the subject average for the over-update and over-update-ratio by treatment, presented in figure 12 in appendix B. We find that the distribution of the over-update in the Low Confidence treatment first-order stochastically dominates the distribution of the over-update in the High Confidence treatment. Regarding the over-update-ratio, the Low Confidence treatment is more likely to have larger values of this measure.

Over-update by Signal Accuracy



Over-update-ratio by Signal Accuracy



● Low Confidence treatment ■ High Confidence treatment — — Bayesian benchmark

Notes: Jitter added to x-axis. Observations dropped with actual prior of zero in panel of ratio measure.

Figure 5: Mean over-update and over-update-ratio by signal accuracy and treatment

Table 1: Ordinary least squares regressions of over-update measures

	Dependent variable			
	over-update		over-update-ratio	
	(1)	(2)	(3)	(4)
High Confidence treatment	−3.804 (0.585)	−3.804 (0.585)	−0.251 (0.085)	−0.252 (0.085)
Constant	−8.139 (0.876)	−8.139 (0.292)	−0.488 (0.088)	−0.488 (0.042)
Subject fixed-effects	No	Yes	No	Yes
R^2	0.010	0.013	0.003	0.003
Subjects	118	118	118	118
Observations	2360	2360	2359	2359

Notes: Standard errors (in parentheses) clustered on subject. Observations with undefined log-ratio dropped.

that High Confidence treatment decreases the over-update measure by 3.8 percentage points ($p < 0.001$) and decreases the over-update-ratio by 0.25 units ($p = 0.004$).

We also perform a non-parametric Wilcoxon signed-rank test to determine whether the over-update and over-update-ratio differ between the High and Low Confidence treatments. We find that the distribution of each measure differs between treatments ($p < 0.0001$ for each).

Result 2. *In the High Confidence treatment subjects exhibit more under-updating relative to the Low Confidence treatment.*

This result is consistent with hypothesis 1. Interestingly, given a noisier prior belief, subjects better approximate Bayesian updating. Finally, given that we find more under-updating in the High Confidence treatment, we ask if this effect is driven by confidence over multiple priors; we investigate this question next.

4.2.1 An Instrumental Variables Regression

We are primarily interested in how confidence over multiple priors affects the degree of over-updating. However, endogeneity is a concern between our over-updating measures and confidence measures. Figure 4 shows that prior confidence varies with actual priors. Accurate

Table 2: OLS and instrumental variable regressions of over-update measures

	Dependent variable			
	over-update		over-update-ratio	
	(1) OLS	(2) FE2SLS [†]	(3) OLS	(4) FE2SLS [†]
Prior confidence, q^*	0.020 (0.021)	-0.175 (0.030)	0.001 (0.002)	-0.012 (0.004)
Constant	-11.670 (1.723)	4.537 (2.470)	-0.717 (0.174)	0.349 (0.340)
First-stage F -stat		320.36		320.36
Subjects	118	118	118	118
Observations	2360	2360	2359	2359

Notes: Standard errors (in parentheses) clustered on subject. Subject fixed-effects included. Observations with undefined log-ratio dropped. Prior confidence q^* is measured in percentage points (between 0 and 100). [†]Fixed-effect two-stage least squares (FE2SLS) regressions use an indicator of High Confidence treatment as the instrument.

perception of a prior is more difficult for priors closer to 50%, which naturally results in lower confidence. Actual priors are also correlated with the over-update and over-update-ratio measures because updated beliefs are a function of priors. We resolve this endogeneity problem by using the treatment as an instrument in a generalized two-stage least squares regression with subject fixed-effects, as presented in table 2. The F -statistic for the first-stage regression is 320, indicating that our treatment is a strong instrument for our confidence measure (Stock and Yogo 2005).

The results are consistent with hypothesis 2.²⁷ Columns 2 and 4 show a statistically significant negative coefficient with the instrumental variable. We find that a percentage-point increase in prior confidence reduces the over-update measure by 0.175 percentage points ($p < 0.001$) and the over-update-ratio by 0.012 units ($p = 0.005$).

Result 3. *Greater under-updating in the High Confidence treatment is due to greater confidence over multiple priors.*

²⁷Columns 1 and 3 of table 2 report positive and non-significant coefficient estimates for prior confidence (the probability report) using an ordinary least squares regression.

4.3 Grether Model Estimation Results

We now present estimates for the Grether (1980) model, using the IV regression specification in equations (3) and (4) to study how our subjects respond to our treatments. We employ the actual prior as the instrument for the log-prior-ratio to account for measurement error in subjects' reported prior beliefs.²⁸ The F -statistic for the first stages are all sufficiently large (> 600), again validating that the actual prior is a strong instrument for the log-prior-ratio. Table 3 presents results for signal accuracy of 60% (column 1), 80% (column 2), both these pooled (column 3), and pooled with subject fixed-effects (column 4).²⁹

In the pooled regressions we find that the weight on prior beliefs, β , is larger by about 0.113 units in the High Confidence treatment compared to the Low treatment (Column 3, $p = 0.027$), indicating that subjects place more weight on the prior in the High Confidence treatment. The signal in our experiment did not change across the High and Low Confidence treatments, thus it is reasonable to expect the weights that the subjects place on the signals would remain unchanged across the treatments. However, this is not the case. We estimate that the weight on the signals, α , is smaller by about 0.111 units in the High Confidence treatment compared to the Low treatment (Column 3, $p < 0.001$). Our subjects respond to our treatments by placing more weight on the signals.

We estimate the Grether (1980) parameters by signal accuracy in columns 1 and 2 of table 3. We see that the overall direction of the results is qualitatively similar. The difference in the estimated Grether (1980) parameters is greater when the signal accuracy is lower. A possible explanation is that given a weaker signal, subjects increasingly rely on the signals more to form their updated beliefs when they have less confidence in their prior beliefs.

We also find that the weight placed on the signal (α) is larger when the signal is weaker. Although we do not find over-reaction to the 60%-accurate signal, our result is similar to those of Augenblick, Lazarus, and Thaler (2024), who find that subjects tend to under-infer from stronger signals. Our design differs from the balls-and-urns framing of Augenblick, Lazarus, and Thaler (2024); we thus validate these findings across different styles of belief-updating tasks.

²⁸A subject may misreport their latent prior belief yet use the latent belief for the updated belief.

²⁹Table 7 in appendix C.2 presents the standard Grether regression using ordinary least squares.

Table 3: Grether model TSLS regressions of log updated belief ratio

	Signal accuracy			
	60%	80%	Pooled	
	(1)	(2)	(3)	(4)
Reduced-form regression:				
α_L	0.634 (0.113)	0.326 (0.042)	0.349 (0.040)	0.354 (0.041)
β_L	0.751 (0.062)	0.774 (0.058)	0.763 (0.043)	0.767 (0.042)
$(\alpha_H - \alpha_L)$	-0.247 (0.108)	-0.101 (0.032)	-0.111 (0.030)	-0.119 (0.030)
$(\beta_H - \beta_L)$	0.147 (0.078)	0.080 (0.067)	0.113 (0.051)	0.108 (0.051)
Linear combinations:				
α_H	0.388 (0.104)	0.224 (0.041)	0.237 (0.039)	0.235 (0.038)
β_H	0.898 (0.039)	0.854 (0.053)	0.876 (0.033)	0.875 (0.033)
p -value of F -test:				
$\alpha_L = \beta_L = 1$	< 0.001	< 0.001	< 0.001	< 0.001
$\alpha_H = \beta_H = 1$	< 0.001	< 0.001	< 0.001	< 0.001
$\alpha_H = \alpha_L$	0.023	0.001	< 0.001	< 0.001
$\beta_H = \beta_L$	0.059	0.231	0.027	0.035
Subject fixed-effects	No	No	No	Yes
First-stage F -stat	2318.53	614.7	1941.18	1929.85
R^2	0.662	0.567	0.606	0.608
Subjects	58	60	118	118
Observations	1025	1068	2093	2093

Notes: Standard errors (in parentheses) clustered on subject. Observations with undefined log-ratio dropped. The null hypothesis of Bayesian updating requires $\alpha = 1$ and $\beta = 1$. Fixed-effect two-stage least squares (FE2SLS) regressions use an indicator of High Confidence treatment as the instrument.

While the regression results present estimates of the Grether (1980) parameters at the aggregate level, we also offer individual-level results. The plots on the left side of figure 6 depict the distribution of $\alpha_{T,s}$ and $\beta_{T,s}$ for each subject s and treatment $T \in \{H, L\}$.

The CDF of $\alpha_{T,s}$ for the High Confidence treatment generally lies to the left of the corresponding CDF for the Low Confidence treatment. The CDF of $\beta_{T,s}$ for the High Confidence treatment generally lies to the right of the corresponding CDF for the Low Confidence treatment. Taken together, at the individual level, we generally observe larger values of β and smaller values of α with High Confidence treatment. Further, these graphs depict a mass at $\alpha = 0$ and $\beta = 1$ in the High Confidence treatment, which corresponds to subject who do not update their beliefs in any task in the High Confidence treatment. In fact about 35% (41 out of 118) of our subjects do not update their beliefs in the High Confidence treatment.³⁰

We also compute the within-subject difference of $\alpha_{T,s}$ and $\beta_{T,s}$ between treatments, as shown in the right-hand side plots of figure 6. Regarding $\alpha_{T,s}$, most subjects have a negative difference; these subjects place more weight on the signal in the Low Confidence treatment than in the High Confidence treatment. Regarding $\beta_{T,s}$, most subjects have a positive difference; these subjects place more weight on the prior in the High Confidence treatment than the Low Confidence treatment.

Result 4. *When updating, subjects place more weight on their prior and less weight on the signals in the High Confidence treatment than in the Low Confidence treatment.*

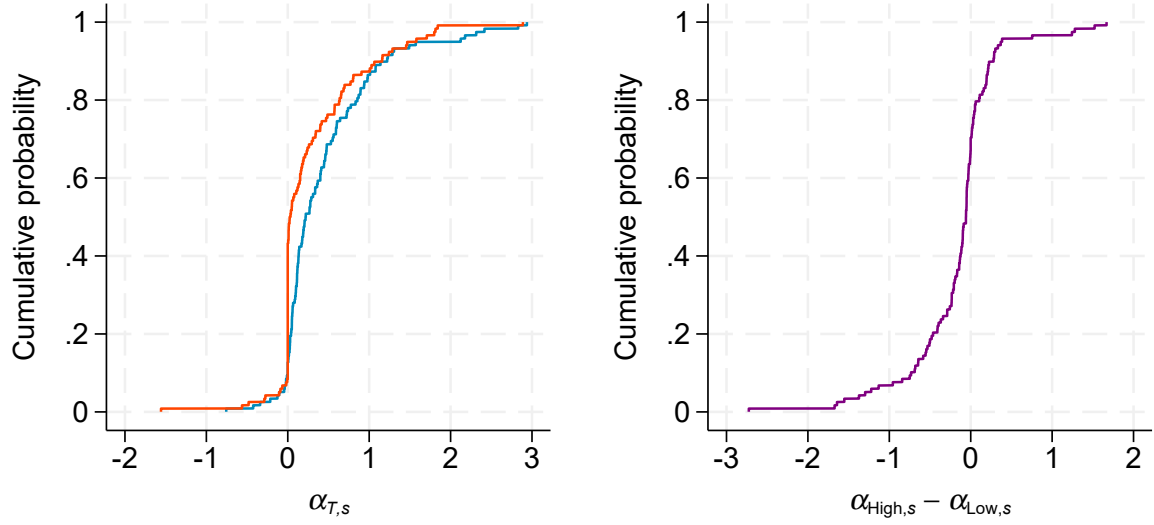
Overall our estimation of the Grether (1980) parameters shows that our subjects place more weight on their prior belief in the High Confidence treatment. Thus far our results are consistent with our hypotheses.

4.4 Noisy Cognition Models

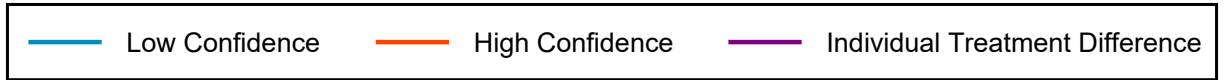
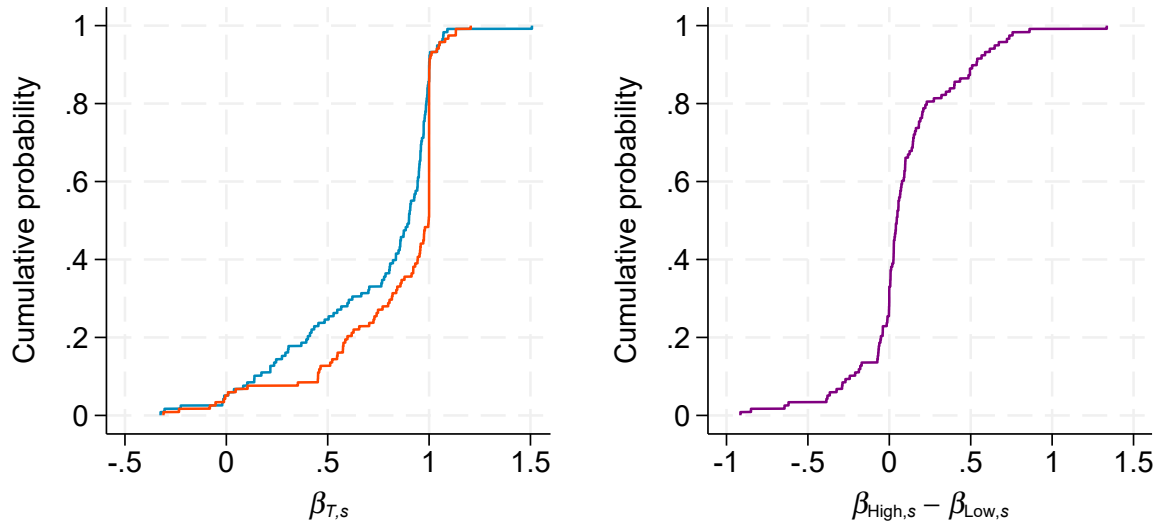
Augenblick, Lazarus, and Thaler (2024) and Ba, Bohren, and Imas (2022) observe considerable over- and under-updating in their experiments, which they rationalize with a noisy

³⁰No updating is a well-documented modal updating pattern. The survey of Benjamin (2019) notes that about one-third to one-half of individuals do not update whatsoever. Our results echo this empirical regularity.

Individual-level Estimates of $\alpha_{T,s}$



Individual-level Estimates of $\beta_{T,s}$



Notes: Each subject s completes ten updating tasks in each treatment T , permitting within-subject estimation of the Grether model for each treatment (on the left) and the resultant treatment effect (right).

Figure 6: Subject level estimates of α (weight on the signal) and β (weight on the prior) in the Grether (1980) model

cognition model. However, our results are not consistent with the existing noisy cognition model (Woodford 2020; Enke and Graeber 2023).

Noisy cognition models typically have a cognitive default (cognitive prior) over a variable, in which an agent observes an unbiased cognitive signal that reveals more information about this variable. The agent then updates their belief over this variable in a Bayesian manner. A standard assumption on the distribution of the cognitive default and the cognitive signal density is the “update towards the signal” property (Chambers and Healy 2012). This ensures that the average posterior belief of the variable is a convex combination of the prior mean and the realized signal value. The noisy cognition model is attractive because it maintains a Bayesian updating assumption and can explain behavioral attenuation (Enke et al. 2024).

Previous literature employs the noisy cognition model with two main approaches. The first approach applies noisy cognition to the primitives of the decision problem. Augenblick, Lazarus, and Thaler (2024) apply the noisy cognition to the signal conditional probability where the agent incorrectly perceives the strength of the signal.³¹ In our context, we will apply the noisy cognition to the priors, given that subjects do not accurately perceive the priors. Based on Proposition ??, we know that uncertainty over multiple priors is not going to affect the average Bayesian posterior belief. Because our over-updating measures are benchmarked against the Bayesian posterior belief as computed using the subjects’ reported prior, we should not expect a difference in the over-update ratio and the Grether (1980) weights if our subjects are updating in the same way across the High and Low Confidence treatments. However, we found that the over-update ratio is greater in the Low Confidence treatment, and our estimation of the Grether (1980) model shows that subjects place less weight on the prior and more weight on the signal in the Low Confidence treatment.

The second approach applies noisy cognition to the optimal action on a continuous action space (Enke and Graeber 2023), where the action is the report of updated beliefs. Conditioned on the perceived prior, the agent has a cognitive default over regarding the Bayesian posterior belief and receives an unbiased cognitive signal about the Bayesian posterior.³² In

³¹As another example in a non-belief updating context, Frydman and Jin (2022) apply the noisy cognition framework to lottery payoffs.

³²Ba, Bohren, and Imas (2022) employs a more general model; their cognitive signal is drawn from a distribution whose mean is the updated belief as computed using an updating rule that accommodates the representative heuristic (Kahneman and Tversky 1972b; Bordalo et al. 2020), which nests the Bayesian

the High Confidence treatment, we should expect a more precise cognitive signal about the Bayesian posterior. On average, we should expect subjects in the High Confidence treatment to be closer to the Bayesian benchmark (relative to the Low). Yet our data say otherwise: subjects in the High Confidence treatment are further from the Bayesian benchmark.

Given that existing noisy cognition models are unable to explain our results, we look toward alternative models. A simple extension to the Grether (1980) model accommodates our results—simply allow the exponential weight on the prior, β , to be an increasing function of the confidence over multiple priors.

4.5 Magnitude of Update

In our presentation of results thus far, we have only considered confidence across multiple priors. Our experiment also allows us to study our other notion of confidence, which is related to the dispersion of a belief distribution.³³ The actual proportion of successful projects represents this notion of confidence. Uncertainty regarding the true state (a success or failure) of the selected project is greater when the proportion of successful projects is closer to 50%.

Figure 7 depicts an inverse U-shape pattern in the Bayesian beliefs. A Bayesian agent would have the largest magnitude of update when the prior belief is close to 50%. This pattern of Bayesian updating is similar when the Bayesian update is computed using the reported prior belief or the actual prior. Further we see that the magnitude of our subjects' updates remains similar regardless of the actual prior.³⁴ This pattern of updating behavior is there even when we split our sample by signal accuracy as shown in figure 14 (see appendix B).

Result 5. *Inconsistent with Bayesian updating, subjects update by the same magnitude (in p.p.) regardless of the actual prior or confidence within the prior distribution.*

This suggests that our subjects are insensitive to the amount of uncertainty from a single prior belief distribution. It seems as if our subjects are using a heuristic where they are

posterior.

³³This is exploratory; we did not intend to look at this before collecting the data.

³⁴This explains the U-shaped pattern in our over-update measure in Figure 5.

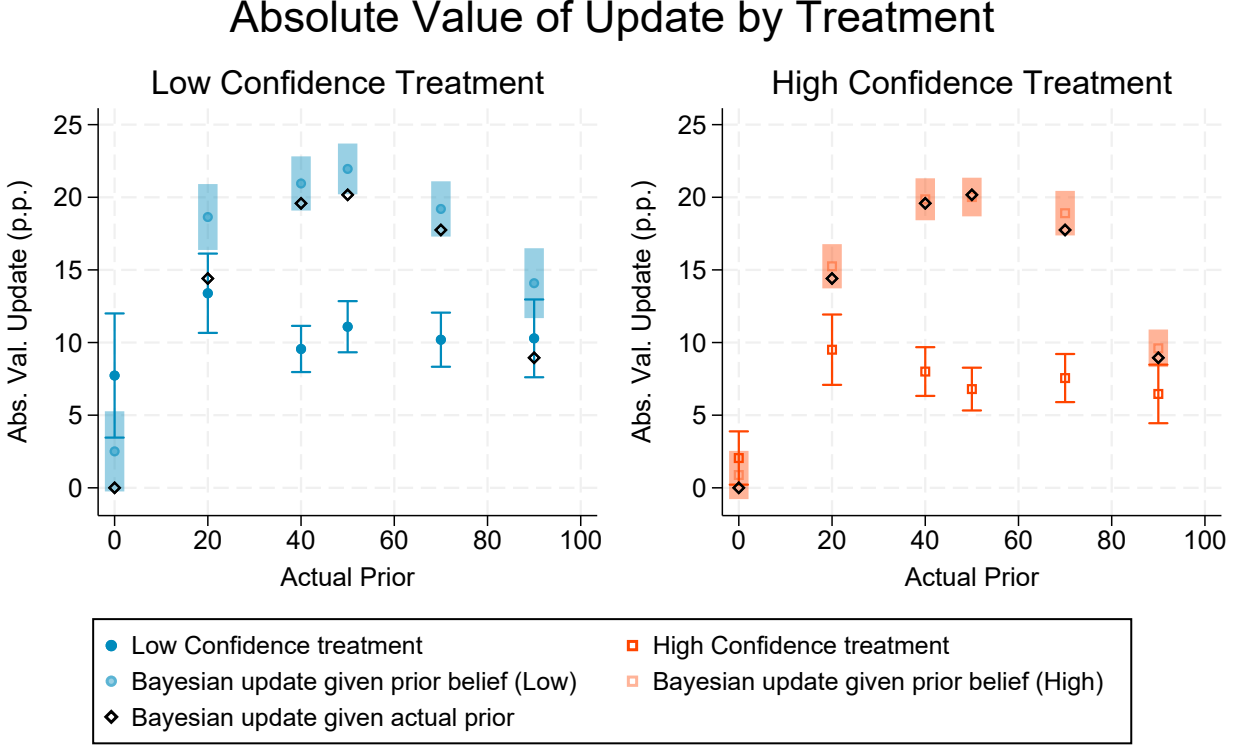


Figure 7: Mean Magnitude of Update by Treatment

updating their beliefs by a fixed amount regardless of the prior beliefs they started with. Our data show that people respond to uncertainty over multiple priors and they do not respond to the uncertainty over the belief distribution. Both of these results are at odds with extant theory.

4.6 Overconfidence in Stating Beliefs

Our final set of results categorizes subjects based on their over-confidence in their stated prior and updated beliefs. For each subject we compute the average reported confidence and the associated 95% confidence interval for their priors and updated beliefs. We then compute the proportion of all elicited beliefs (both prior and updated beliefs) that are within three percentage points of the actual prior or the true Bayesian posterior. We define a subject as over-confident if the proportion falls above the 95% confidence interval, under-confident if below, and neutral if within the interval.

Figures 8 and 9 demonstrate that subjects are largely over-confident when it comes

Confidence in Prior Belief by Treatment

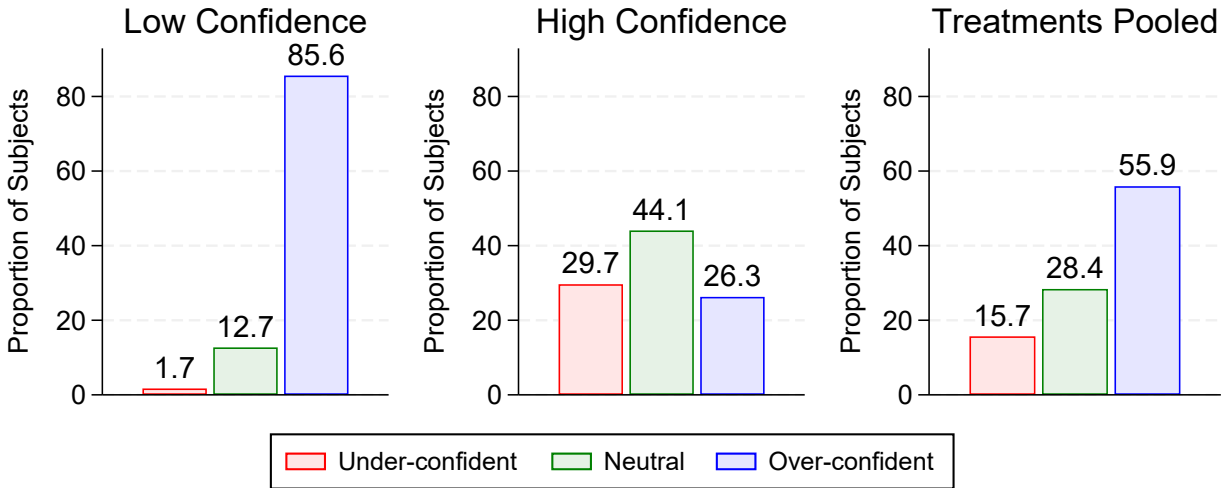


Figure 8: Subject proportions by confidence in prior belief

Confidence in Updated Belief by Treatment

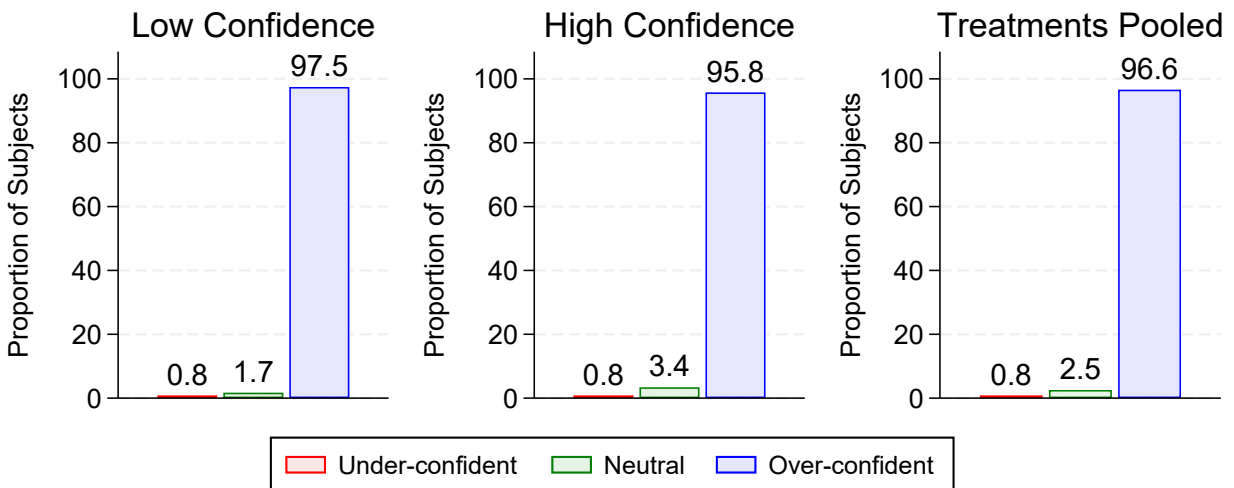


Figure 9: Subject proportions by confidence in updated belief

Mean Over-confidence by Treatment

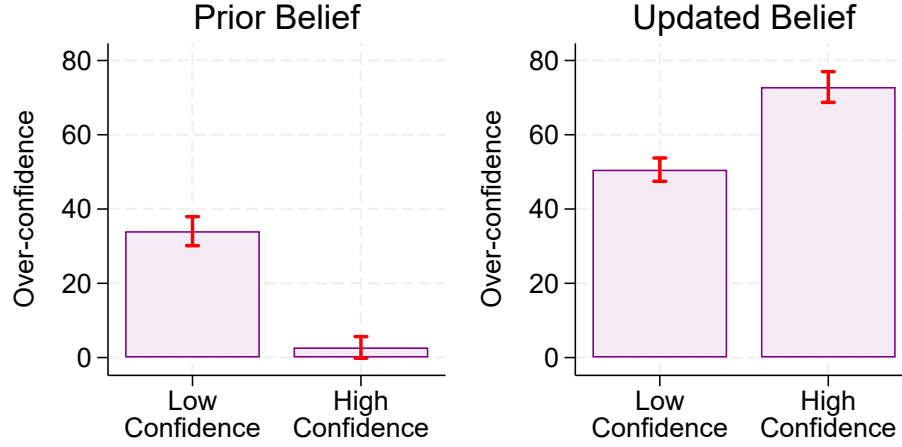


Figure 10: Mean (continuous) over-confidence measure with 95% confidence intervals

to reporting both their priors and their updated beliefs. Regarding prior beliefs (figure 8), however, the modal subject is neutral in the High Confidence treatment, with the proportions of over- and under-confident types being relatively balanced. The vast majority of subjects are over-confident regarding prior beliefs in the Low Confidence treatment. With respect to updated beliefs, figure 9 shows that subjects are overwhelmingly over-confident in both treatments relative to the Bayesian posterior. In general, our subjects could have increased their expected earnings by reporting lower confidence values.

We also define a continuous measure of confidence for each subject: we take the subject's mean stated confidence-in-beliefs and subtract the proportion of elicited beliefs (both prior and posterior) that are actually within three percentage points of the true value. A positive value indicates an over-confident subject and a negative value an under-confident subject.

Result 6. *Subjects are overly confident in their prior beliefs in the Low Confidence treatment. Subjects are over-confident in their updated beliefs in both treatments.*

Figure 10 shows the mean continuous over-confidence measure.³⁵ Results with the continuous measure are similar to those with the discrete measure. With respect to the prior, in the Low Confidence treatment, we see that the average confidence stated by our subjects

³⁵We also plot the distribution of the measure in figure 13 (see appendix B). The vast majority of our subjects are over-confident, echoing our results.

is about 37 percentage points more than the proportion in which their stated prior is within 3 percentage points of the actual value. For the High Confidence treatment, we see that our over-confidence measure is about 3 percentage points and this is not statistically different from zero ($p = 0.0697$).

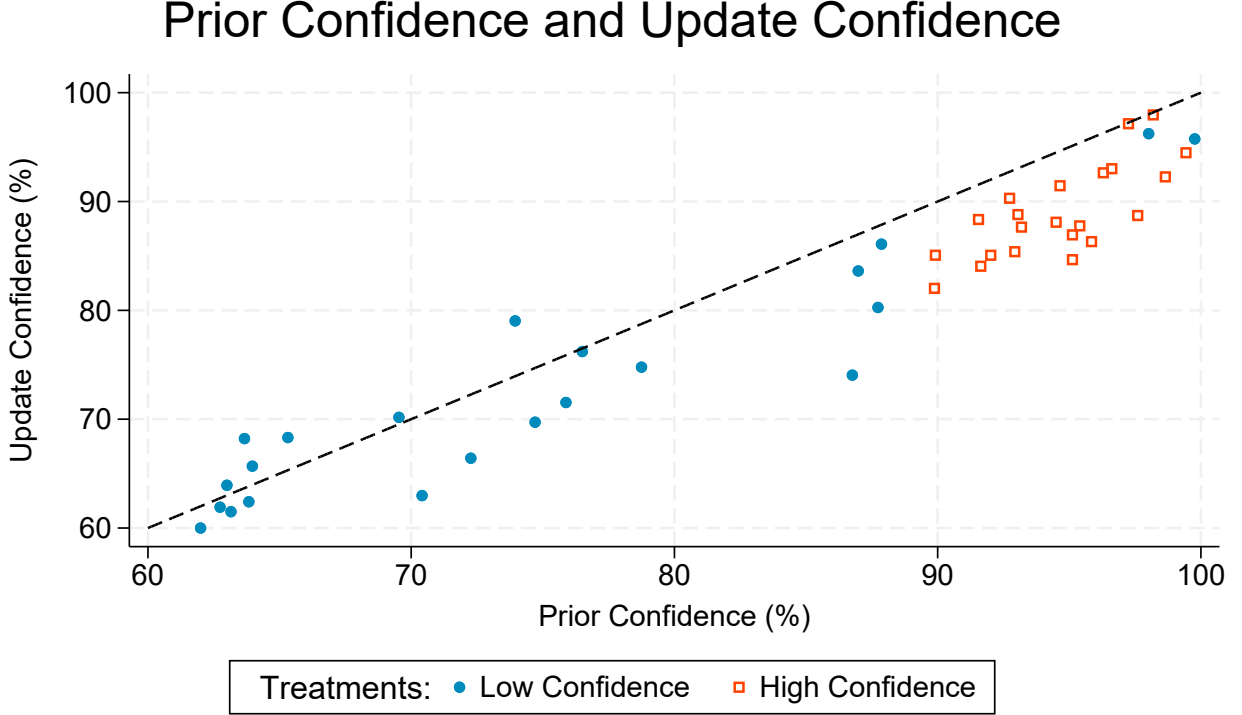
With respect to updated beliefs, we observe more over-confidence than in the case of prior beliefs. This suggests that people are making mistakes in updating their beliefs and are over-confident in their ability to update. In the High Confidence treatment, subjects state confidence that is on average 73 percentage points higher than the actual proportion of the times their stated beliefs are within three percentage points of the actual Bayesian belief, while in the Low Confidence treatment, subjects state confidence that is on average 51 percentage points more than the proportion of the times their stated beliefs are within three percentage of the actual Bayesian belief.

A surprising result is that the degree of over-confidence is significantly larger in the High Confidence treatment compared to the Low Confidence treatment ($p < 0.001$). In the High Confidence treatment subjects have more accurate priors *and* are more confident in their prior, yet they are more over-confident in their updated beliefs.

Result 7. *Subjects are more over-confident in stating their updated beliefs in the High Confidence treatment.*

To further examine this result, we plot the relationship between confidence in prior beliefs and confidence in updated beliefs. Figure 11 depicts a positive correlation for the mean value across subjects for each task. The fit of the 45-degree-line is striking ($R^2 = 0.9168$). This suggests that, on average, their confidence-in-prior-belief directly translates into confidence-in-updated-belief, implying that subjects may think that they are Bayesian or accurate in their belief updating process and that their margin of error is within three percentage points.³⁶ This provides suggestive evidence that the majority of the subjects think that they are Bayesians when in fact they are not.

³⁶Recall that we informed subjects that they would earn a \$3 bonus if their reported updated belief is within three percentage points of a value computed from a statistical process, Bayes' theorem. This suggests that the subjects believe that they update their beliefs consistent with Bayes' theorem.



Notes: Each point represents mean subject confidence for a given task.

Figure 11: Relationship between confidence in prior and confidence in updated beliefs

5 Conclusion

Our study is a first to shed light on the critical relationship between confidence in one’s prior beliefs and belief updating and test it against Bayesian predictions. We find that confidence over multiple priors matters when it shouldn’t and confidence in a belief distribution doesn’t matter when it should. Our subjects’ behavior deviates significantly from the Bayesian prediction.

Our novel design varies subjects’ confidence over multiple priors and elicits incentive-compatible confidence reports regarding beliefs. We find that greater confidence over multiple prior beliefs leads to more under-updating when a Bayesian agent should not respond to this notion of confidence. We find that subjects are insensitive to an alternative notion of confidence that corresponds to dispersion within a single belief distribution. Our results contribute to a broader understanding of belief updating, suggesting that confidence over multiple prior beliefs plays a role in how individuals update their beliefs through the weight placed on the priors and signals in the updating process.

Our results also present striking evidence of over-confidence when subjects report both their prior and updated beliefs. For updated beliefs, we observe a greater degree of over-confidence in the High Confidence treatment, in which subjects report prior beliefs with greater accuracy and are more confident in those prior beliefs. This suggests that having more accurate prior beliefs does not necessarily lead to increasingly accurate updating behavior. Our finding challenges the notion that having increasingly accurate priors, and confidence in those priors, inherently results in better-calibrated updated beliefs. We conclude by noting first our contribution to the elicitation of beliefs literature and second the stylized facts we present about the relationship between confidence and belief-updating. These facts should motivate further models of belief-formation processes.

References

- Agranov, Marina, and Pëllumb Reshidi. 2023. “Disentangling Suboptimal Updating: Complexity, Structure, and Sequencing.”
- Augenblick, Ned, and Eben Lazarus. 2023. *A new test of excess movement in asset prices*. Technical report.
- Augenblick, Ned, Eben Lazarus, and Michael Thaler. 2024. *Overinference from weak signals and underinference from strong signals*. Technical report. arXiv. <https://doi.org/10.48550/arXiv.2109.09871>.
- Augenblick, Ned, and Matthew Rabin. 2021. “Belief movement, uncertainty reduction, and rational updating.” *Quarterly Journal of Economics* 136 (2): 933–985.
- Ba, Cuimin, J Aislinn Bohren, and Alex Imas. 2022. “Over- and under-reaction to information.” <https://cuiminba.com/working-papers/overreaction/>.
- Becker, Gordon M, Morris H DeGroot, and Jacob Marschak. 1964. “Measuring utility by a single-response sequential method.” *Behavioral science* 9 (3): 226–232.
- Bénabou, Roland, and Jean Tirole. 2006. “Incentives and prosocial behavior.” *American Economic Review* 96 (5): 1652–1678.
- Benjamin, Daniel J. 2019. “Errors in probabilistic reasoning and judgment biases.” In *Handbook of Behavioral Economics: Applications and Foundations 1*, 69–186. Elsevier.
- Bianchi, Francesco, Cosmin Ilut, and Hikaru Saijo. 2024. “Diagnostic business cycles.” *Review of Economic Studies* 91 (1): 129–162.
- Bordalo, Pedro, Nicola Gennaioli, Yueran Ma, and Andrei Shleifer. 2020. “Overreaction in macroeconomic expectations.” *American Economic Review* 110 (9): 2748–2782.

- Bordalo, Pedro, Nicola Gennaioli, and Andrei Shleifer. 2018. "Diagnostic expectations and credit cycles." *The Journal of Finance* 73 (1): 199–227.
- Brier, Glenn W. 1950. "Verification of forecasts expressed in terms of probability." *Monthly Weather Review* 78 (1): 1–3.
- Brown, Sebastian, and Kenneth Chan. 2024. "How Much Can I Make? Insights on Belief Updating in the Labor Market."
- Cason, Timothy N, and Charles R Plott. 2014. "Misconceptions and game form recognition: Challenges to theories of revealed preference and framing." *Journal of Political Economy* 122 (6): 1235–1270.
- Chambers, Christopher P, and Paul J Healy. 2012. "Updating toward the signal." *Economic Theory* 50:765–786.
- Charness, Gary, James C Cox, Catherine C Eckel, Charles A Holt, and Brian Jabarian. 2023. *The Virtues of Lab Experiments*. Technical report.
- Charness, Gary, and Chetan Dave. 2017. "Confirmation bias with motivated beliefs." *Games and Economic Behavior* 104 (July): 1–23.
- Charness, Gary, and Martin Dufwenberg. 2006. "Promises and partnership." *Econometrica* 74 (6): 1579–1601.
- Charness, Gary, Ryan Oprea, and Sevgi Yuksel. 2021. "How do people choose between biased information sources? Evidence from a laboratory experiment." *Journal of the European Economic Association* 19 (3): 1656–1691.
- Conlon, John J, Laura Pilossoph, Matthew Wiswall, and Basit Zafar. 2018. *Labor market search with imperfect information and learning*. Technical report. National Bureau of Economic Research.
- Coutts, Alexander. 2019. "Good news and bad news are still news: Experimental evidence on belief updating." *Experimental Economics* 22 (2): 369–395.
- Cripps, Martin W. 2018. "Divisible updating." *Manuscript, UCL*.
- Danz, David, Lise Vesterlund, and Alistair J Wilson. 2022. "Belief elicitation and behavioral incentive compatibility." *American Economic Review* 112 (9): 2851–2883.
- DellaVigna, Stefano, and Devin Pope. 2018. "Predicting experimental results: who knows what?" *Journal of Political Economy* 126 (6): 2410–2456.
- Dufwenberg, Martin, and Uri Gneezy. 2000. "Measuring beliefs in an experimental lost wallet game." *Games and Economic Behavior* 30 (2): 163–182.
- Eil, David, and Justin M Rao. 2011. "The good news-bad news effect: asymmetric processing of objective information about yourself." *American Economic Journal: Microeconomics* 3 (2): 114–138.
- Ellsberg, Daniel. 1961. "Risk, ambiguity, and the Savage axioms." *The quarterly journal of economics* 75 (4): 643–669.

- Enke, Benjamin, and Thomas Graeber. 2023. “Cognitive uncertainty.” *Quarterly Journal of Economics* 138 (4): 2021–2067.
- Enke, Benjamin, Thomas Graeber, Ryan Oprea, and Jeffrey Yang. 2024. “Behavioral Attenuation.” *mimeo*.
- Enke, Benjamin, and Florian Zimmermann. 2019. “Correlation neglect in belief formation.” *The Review of Economic Studies* 86 (1): 313–332.
- Epstein, Larry G, Jawwad Noor, and Alvaro Sandroni. 2010. “Non-bayesian learning.” *The BE Journal of Theoretical Economics* 10 (1).
- Esponda, Ignacio, Ryan Oprea, and Sevgi Yuksel. 2023. “Seeing what is representative.” *Quarterly Journal of Economics* 138 (4): 2607–2657.
- Esponda, Ignacio, Emanuel Vespa, and Sevgi Yuksel. 2024. “Mental models and learning: The case of base-rate neglect.” *American Economic Review* 114 (3): 752–782.
- Eyster, Erik, and Georg Weizsacker. 2016. “Correlation neglect in portfolio choice: Lab evidence.”
- Frydman, Cary, and Lawrence J Jin. 2022. “Efficient coding and risky choice.” *The Quarterly Journal of Economics* 137 (1): 161–213.
- Gennaioli, Nicola, and Andrei Shleifer. 2010. “What comes to mind.” *Quarterly Journal of Economics* 125 (4): 1399–1433.
- Gilboa, Itzhak, and Massimo Marinacci. 2016. “Ambiguity and the Bayesian paradigm.” In *Readings in formal epistemology: Sourcebook*, 385–439. Springer.
- Gilboa, Itzhak, and David Schmeidler. 1993. “Updating ambiguous beliefs.” *Journal of economic theory* 59 (1): 33–49.
- Greiner, Ben. 2015. “Subject pool recruitment procedures: organizing experiments with ORSEE.” *Journal of the Economic Science Association* 1 (1): 114–125.
- Grether, David M. 1980. “Bayes rule as a descriptive model: The representativeness heuristic.” *Quarterly Journal of Economics* 95 (3): 537–557.
- Halevy, Yoram. 2007. “Ellsberg revisited: An experimental study.” *Econometrica* 75 (2): 503–536.
- Healy, Paul J. 2020. “Explaining the BDM – or any random binary choice elicitation mechanism – to Subjects.” *Mimeo*.
- Holt, Charles A, and Susan K Laury. 2002. “Risk aversion and incentive effects.” *American Economic Review* 92 (5): 1644–1655.
- Hossain, Tanjim, and Ryo Okui. 2013. “The binarized scoring rule.” *Review of Economic Studies* 80 (3): 984–1001.
- Kahneman, Daniel, and Amos Tversky. 1972a. *On the Psychology of Prediction*. Research Bulletin 12 (4). Oregon Research Institute, October.

- Kahneman, Daniel, and Amos Tversky. 1972b. "Subjective probability: A judgment of representativeness." *Cognitive Psychology* 3 (3): 430–454.
- Klibanoff, Peter, and Eran Hanany. 2007. "Updating preferences with multiple priors." *Theoretical Economics* 2 (3): 261–298.
- Kovach, Matthew. 2024. "Ambiguity and partial Bayesian updating." *Economic Theory* 78 (1): 155–180.
- Levy, Gilat, Inés Moreno de Barreda, and Ronny Razin. 2022. "Persuasion with correlation neglect: a full manipulation result." *American Economic Review: Insights* 4 (1): 123–138.
- Liang, Yucheng. 2024. "Learning from unknown information sources." *Management Science*, <https://doi.org/10.1287/mnsc.2021.03551>.
- Möbius, Markus M, Muriel Niederle, Paul Niehaus, and Tanya S Rosenblat. 2022. "Managing self-confidence: Theory and experimental evidence." *Management Science* 68 (11): 7793–7817.
- Nickerson, Raymond S. 1998. "Confirmation bias: A ubiquitous phenomenon in many guises." *Review of General Psychology* 2 (2): 175–220.
- Ortoleva, Pietro. 2012. "Modeling the change of paradigm: Non-Bayesian reactions to unexpected news." *American Economic Review* 102 (6): 2410–2436.
- . 2024. "Alternatives to bayesian updating." *Annual Review of Economics* 16.
- Phillips, Lawrence D, and Ward Edwards. 1966. "Conservatism in a simple probability inference task." *Journal of Experimental Psychology* 72 (3): 346.
- Pires, Cesaltina Pacheco. 2002. "A rule for updating ambiguous beliefs." *Theory and Decision* 53:137–152.
- Rabin, Matthew, and Joel L Schrag. 1999. "First impressions matter: A model of confirmatory bias." *Quarterly Journal of Economics* 114 (1): 37–82.
- Stock, James, and Motohiro Yogo. 2005. "Testing for Weak Instruments in Linear IV Regression." In *Identification and Inference for Econometric Models*, edited by Donald W.K. Andrews, 80–108. New York: Cambridge University Press. http://www.economics.harvard.edu/faculty/stock/files/TestingWeakInstr_Stock%5C%2BYogo.pdf.
- Thaler, Michael. 2021. "The supply of motivated beliefs." *arXiv preprint arXiv:2111.06062*.
- Woodford, Michael. 2020. "Modeling imprecision in perception, valuation, and choice." *Annual Review of Economics* 12 (1): 579–601.

A Model with n states and Continuum of Priors

Consider a finite state space $\Omega := \{\omega_1, \dots, \omega_n\}$,³⁷ where the randomly selected project is either a success or a failure. Suppose the agent considers a set Π_0 of possible priors, where the priors are indexed by an ordered set Θ . We define the continuous probability space $(\Theta, \mathcal{F}, K)$, where Θ is the state space, $\mathcal{F}$ is the σ -algebra and K is a probability measure. We define the probability of prior θ at state i as $\pi_{0,i}^\theta$. The agent has a belief distribution over Π_0 , assigning subjective probability density k_0^θ to prior θ . Ultimately the agent's average prior belief $\pi_{0,i}$ about the state i being realized is

$$\pi_{0,i} := \int_{\theta \in \Theta} k_0^\theta \pi_{0,i}^\theta.$$

Let S be the signal space. The agent receives an informative signal $\sigma \in S$ about the realized state. Define the conditional probability of observing signal σ conditioned on ω_i as $P(\sigma|\omega_i) \in [0, 1]$. The posterior belief is also a weighted average of posteriors. To obtain this posterior, the agent must update each prior π_0^θ to a posterior π_1^θ *as well as* update each respective prior weight k_0^θ to a posterior weight k_1^θ . Intuitively, the signal provides information about the state and which prior was the correct prior. However, for a Bayesian agent, updating with the average priors is equivalent to using the aforementioned procedure.

Proposition 2. *The average Bayesian posterior after observing signal σ is*

$$\pi_1^\theta(\sigma) = \int_{\theta \in \Theta} k_1^\theta(\sigma) \pi_{1,i}^\theta(\sigma),$$

which is the same as Bayesian posterior that is updated with the average prior belief.

Proof. Let us update each prior π_0^θ and its respective weight k_0^θ individually using Bayes' rule to obtain

$$\pi_{1,i}^\theta(\sigma) = \frac{\pi_{0,i}^\theta P(\sigma|\omega_i)}{P_\theta(\sigma)} \quad \text{and} \quad k_1^\theta(\sigma) := \frac{k_0^\theta P_\theta(\sigma)}{P(\sigma)},$$

where $P(\sigma) = \sum_{i=1}^n \int_{\theta \in \Theta} k_0^\theta \pi_{0,i}^\theta P(\sigma|\omega_i)$ is the probability of observing signal σ given the mixture prior, while $P_\theta(\sigma) = \int_{\theta \in \Theta} \pi_{0,i}^\theta P(\sigma|\omega_i)$ is the probability of observing signal σ given

³⁷We can also generalize this to a continuous state space as well; we define the continuous probability space $(\Omega, \mathcal{F}, P)$, where Ω is the state space, $\mathcal{F}$ is the σ -algebra and P is a probability measure.

prior π_0^θ . We now show that updating the beliefs of each prior and the mixture distribution over these is equivalent to simply updating the average prior belief.

$$\begin{aligned}
\pi_{1,i}^\theta(\sigma) &= \frac{\pi_{0,i}P(\sigma|\omega_i)}{P(\sigma)} \\
&= \frac{\int_{\theta \in \Theta} k_0^\theta \pi_{0,i}^\theta P(\sigma|\omega_i)}{P(\sigma)} \\
&= \frac{\int_{\theta \in \Theta} k_0^\theta P_j(\sigma) \pi_{0,i} P(\sigma|\omega_i) / P_j(\sigma)}{P(\sigma)} \\
&= \frac{\int_{\theta \in \Theta} k_0^\theta P_j(\sigma) \pi_{1,i}^\theta}{P(\sigma)} \\
&= \int_{\theta \in \Theta} k_1^\theta(\sigma) \pi_{1,i}^\theta(\sigma).
\end{aligned}$$

□

B Supplemental Results

Distributions of Over-updating and Overconfidence Figure 12 depicts the distribution of the average subject over-update and over-update-ratio measures. Figure 13 depicts the distribution of the average subject continuous over-confidence measure.

Absolute Update by Priors Figure 14 plots the magnitude of the absolute value of updates for each signal accuracy and treatment pair.

C Robustness Checks

C.1 Exclusion of Subjects Who Made Mistakes

As a robustness check, we exclude subjects who make a mistake. Specifically, we exclude subjects who in the updating task with a degenerate prior reported either (1) an incorrect prior or (2) a confidence level, as well as subjects who (3) updated in the wrong direction. Our over-update measures do well with updating in the wrong direction; this can be interpreted as severe under-updating. As shown in tables 4 to 6, our results are qualitatively similar

Table 4: OLS regressions of over-update measures (robustness check)

	Dependent variable			
	over-update		over-update-ratio	
	(1)	(2)	(3)	(4)
High Confidence treatment	−3.647 (0.613)	−3.647 (0.613)	−0.266 (0.096)	−0.267 (0.096)
Constant	−8.178 (0.940)	−8.178 (0.307)	−0.508 (0.096)	−0.507 (0.048)
Subject fixed-effects	No	Yes	No	Yes
R^2	0.010	0.013	0.003	0.003
Subjects	101	101	101	101
Observations	2020	2020	2019	2019

Notes: Standard errors (in parentheses) clustered on subject. Observations with undefined log-ratio dropped. Excludes 17 subjects who failed a comprehension check.

upon excluding subjects who made these mistakes.

C.2 Grether Regression with Ordinary Least Squares

Table 7 presents results for the Grether (1980) model using ordinary least squares. In the pooled regressions (columns 3 and 4) we find that β is larger by about 0.1 units in the High Confidence treatment compared to the Low treatment ($p = 0.002$), indicating that subjects place more weight on the prior in the High Confidence treatment.

We estimate the parameters for each signal accuracy in columns 1 and 2 of table 7. We see that the results are qualitatively similar. We find one notable exception: the weight placed on the signal (α) is larger when the signal is weaker ($p = 0.021$). Although we do not find over-reaction to the 60%-accurate signal, our result is similar to those of Augenblick, Lazarus, and Thaler (2024), who find that people tend to under-infer from stronger signals.³⁸ This validates findings across frames of belief-updating tasks.

³⁸Our design differs from Augenblick, Lazarus, and Thaler (2024) who used “balls and urns” framing.

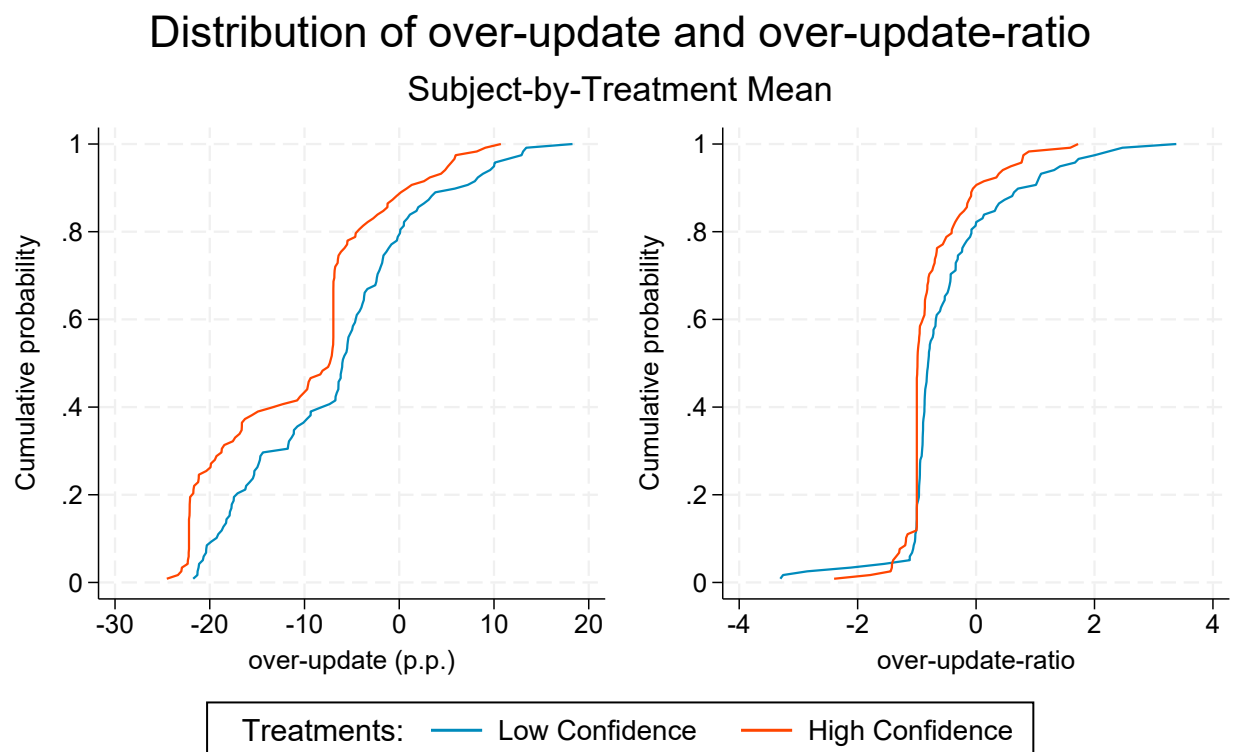
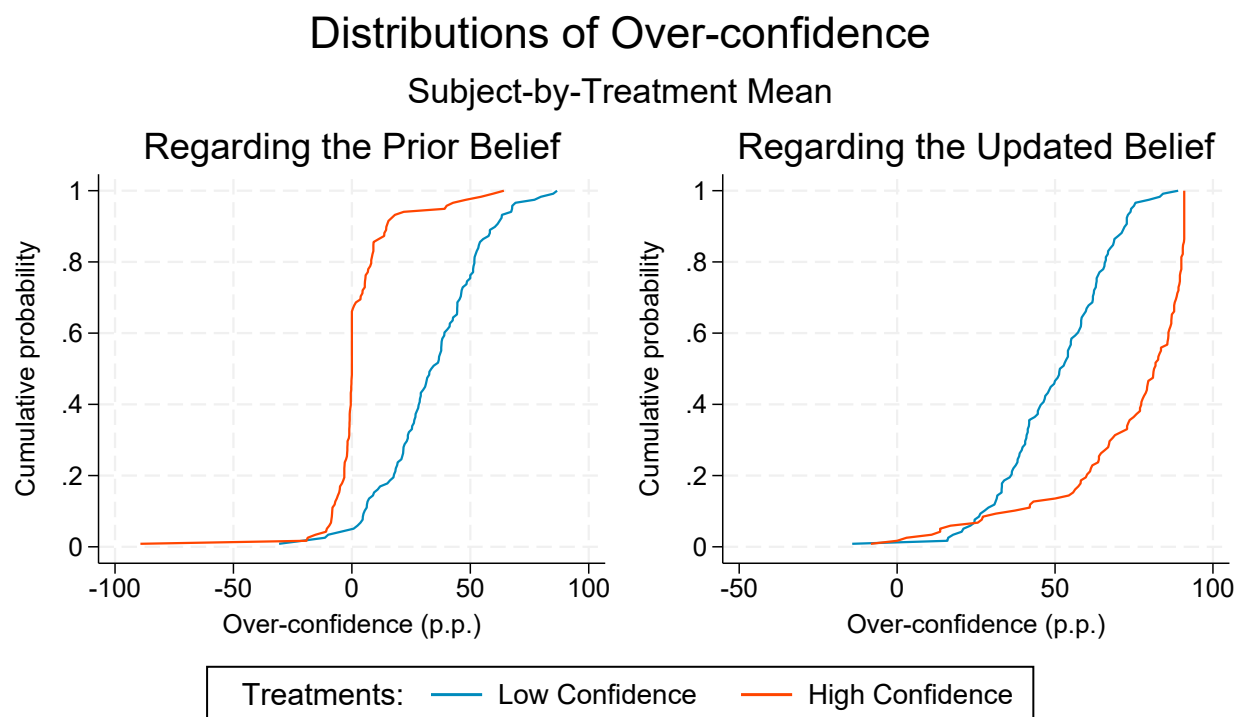


Figure 12: Distribution of the average subject over-update and over-update-ratio measures



Notes: Over-confidence is the difference between a subject's confidence-in-belief and accuracy (by treatment).

Figure 13: Distribution of the average subject continuous over-confidence measure

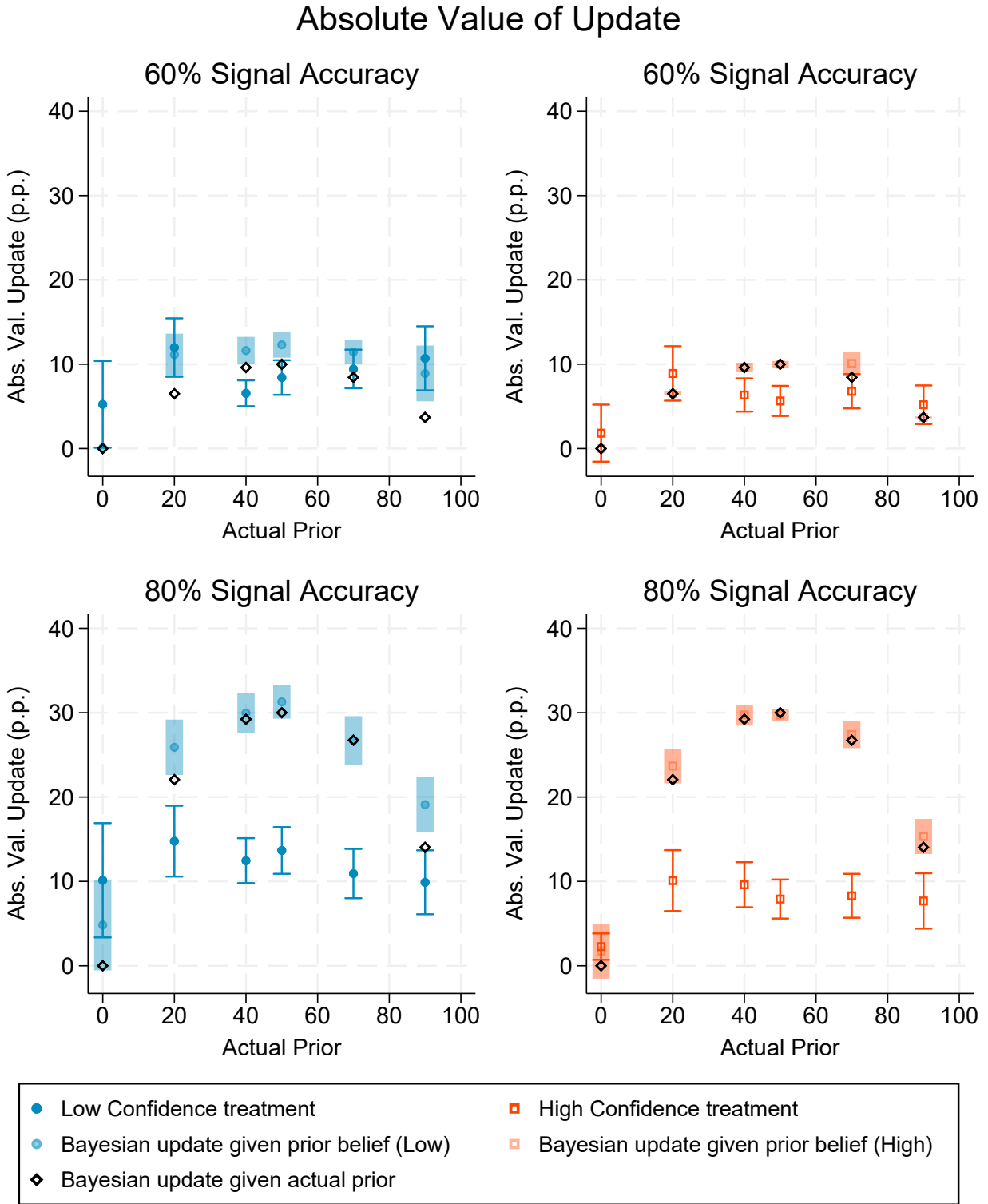


Figure 14: Magnitude of absolute value of updating by signal and treatment accuracy

Table 5: OLS and IV regressions of over-update measures (robustness check)

	Dependent variable			
	over-update		over-update-ratio	
	(1) OLS	(2) FE2SLS [†]	(3) OLS	(4) FE2SLS [†]
Prior confidence, q^*	0.015 (0.022)	−0.166 (0.031)	0.002 (0.002)	−0.012 (0.005)
Constant	−11.260 (1.860)	3.939 (2.630)	−0.773 (0.199)	0.377 (0.388)
First-stage F -stat		259.83		259.83
Subjects	101	101	101	101
Observations	2020	2020	2019	2019

Notes: Standard errors (in parentheses) clustered on subject. Subject fixed-effects included. Observations with undefined log-ratio dropped. Prior confidence q^* is measured in percentage points (between 0 and 100). Excludes 17 subjects who failed a comprehension check.

[†]Fixed-effect two-stage least squares (FE2SLS) regressions use an indicator of High Confidence treatment as the instrument.

Table 6: Grether model TSLS regressions (robustness check)

	Signal accuracy			
	60%	80%	Pooled	
	(1)	(2)	(3)	(4)
Reduced-form regression:				
α_L	0.586 (0.111)	0.320 (0.048)	0.341 (0.045)	0.348 (0.046)
β_L	0.747 (0.066)	0.772 (0.056)	0.759 (0.043)	0.763 (0.043)
$(\alpha_H - \alpha_L)$	-0.283 (0.116)	-0.093 (0.034)	-0.107 (0.032)	-0.118 (0.033)
$(\beta_H - \beta_L)$	0.156 (0.082)	0.040 (0.056)	0.099 (0.050)	0.097 (0.050)
Linear combinations:				
α_H	0.303 (0.101)	0.227 (0.044)	0.234 (0.041)	0.230 (0.041)
β_H	0.903 (0.040)	0.812 (0.054)	0.859 (0.034)	0.860 (0.034)
p -value of F -test:				
$\alpha_L = \beta_L = 1$	< 0.001	< 0.001	< 0.001	< 0.001
$\alpha_H = \beta_H = 1$	< 0.001	< 0.001	< 0.001	< 0.001
$\alpha_H = \alpha_L$	0.014	0.006	0.001	< 0.001
$\beta_H = \beta_L$	0.059	0.470	0.048	0.052
Subject fixed-effects	No	No	No	Yes
First-stage F -stat	2437.11	1835.54	4230.23	4244.52
R^2	0.669	0.596	0.627	0.628
Subjects	52	49	101	101
Observations	919	876	1795	1795

Notes: Standard errors (in parentheses) clustered on subject. Observations with undefined log-ratio dropped. The null hypothesis of Bayesian updating requires $\alpha = 1$ and $\beta = 1$. Fixed-effect two-stage least squares (FE2SLS) regressions use an indicator of High Confidence treatment as the instrument. Excludes 17 subjects who failed a comprehension check.

Table 7: Grether model OLS regressions of log updated belief ratio

	Signal accuracy			
	60%	80%	Pooled	
	(1)	(2)	(3)	(4)
Reduced-form regression:				
α_L	0.630 (0.114)	0.324 (0.041)	0.348 (0.040)	0.353 (0.040)
β_L	0.739 (0.045)	0.723 (0.047)	0.731 (0.033)	0.729 (0.034)
$(\alpha_H - \alpha_L)$	-0.245 (0.110)	-0.102 (0.031)	-0.112 (0.030)	-0.119 (0.030)
$(\beta_H - \beta_L)$	0.133 (0.041)	0.082 (0.040)	0.107 (0.029)	0.095 (0.029)
Linear combinations:				
α_H	0.386 (0.105)	0.222 (0.041)	0.236 (0.039)	0.234 (0.038)
β_H	0.872 (0.032)	0.806 (0.043)	0.838 (0.027)	0.824 (0.028)
p -value of F-test:				
$\alpha_L = \beta_L = 1$	< 0.001	< 0.001	< 0.001	< 0.001
$\alpha_H = \beta_H = 1$	< 0.001	< 0.001	< 0.001	< 0.001
$\alpha_H = \alpha_L$	0.030	0.002	< 0.001	< 0.001
$\beta_H = \beta_L$	0.002	0.047	< 0.001	0.002
Subject fixed-effects	No	No	No	Yes
R^2	0.674	0.584	0.622	0.638
Subjects	58	60	118	118
Observations	1025	1068	2093	2093

Notes: Standard errors (in parentheses) clustered on subject. Observations with undefined log-ratio dropped. The null hypothesis of Bayesian updating requires $\alpha = 1$ and $\beta = 1$.

On Prior Confidence and Belief Updating

Supplementary Materials

Kenneth Chan, Gary Charness, Chetan Dave, J. Lucas Reddinger

December 13, 2024

1 Interface

Figures 1 to 3 show the experimental interface and comprehension checks.

2 Instruments

Figures 4 to 26 show the slides that we personally presented to subjects using a projector. Subjects also retained a printed copy (two slides per page) for reference during the experiment.

Comprehension Check

Does the proportion of success and failure of projects vary across tasks?

- ☐ Yes
- ☐ No

Select all the statement(s) that is/are true about a test with 80% reliability?

- ☐ The test result will be Positive with 80% chance when the project is a success.
- ☐ The test result will be Positive with 80% chance when the project is a failure.
- ☐ The test result will be Negative with 80% chance when the project is a success.
- ☐ The test result will be Negative with 80% chance when the project is a failure.

Suppose a test has 80% reliability. After seeing a positive test result, the selected project is more likely to be a _____.

- ☐ Success
- ☐ Failure
- ☐ Not possible to tell

The more confident I am of my guess being within 3 percentage points of the actual value, I should

- ☐ Report a higher level of confidence
- ☐ Report a lower level of confidence
- ☐ What I report does not matter

----- Page Break -----

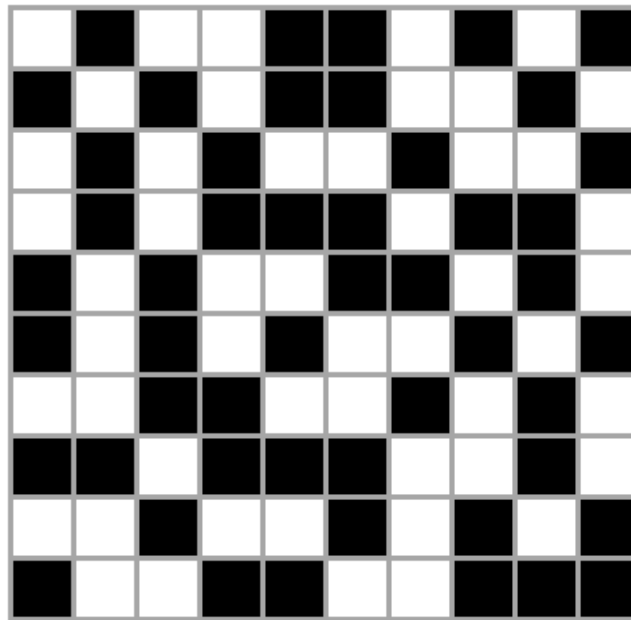
Figure 1: Interface, part 1 of 3

There are 100 projects arranged on a 10 by 10 grid.

- White Square = Success
- Black Square = Failure

Click the next button for the grid to appear on your screen.

----- Page Break -----



----- Page Break -----

One of the 100 projects is randomly selected for you to evaluate (with all projects having an equal chance of being selected). What is the chance that the selected project is a Success (white boxes)? Please input a number between 0-100 indicating the percentage of the project being a Success.

----- Page Break -----

Figure 2: Interface, part 2 of 3

You guessed that the chance of the randomly selected project is a success (white boxes) is _____%.

Please indicate the level of confidence, as a percentage between 0-100, you have that your guess is within 3 percentage points of the actual value.

----- Page Break -----

Your guess that the randomly selected project is a success: _____%.

To further aid your assessment, the computer will run a test on the selected project. The test has a reliability of 80%, so it will be correct four times out of five.

- If the selected project is a Success, the test result will be Positive with 80% chance (four times out of five) and the test result will be Negative with 20% chance (one time out of five).
- If the selected project is a Failure, the test result will be Positive with 20% chance (one time out of five) and the test result will be Negative with 80% chance (four times out of five).

Test result: **Positive**

After seeing the test result, what is the chance that the selected project is a Success? Please input a number between **0-100** indicating the percentage of the project being a Success.

----- Page Break -----

After seeing the test result, you guessed that the chance of the randomly selected project is a success (white boxes) is _____%.

Please indicate the level of confidence, as a percentage between 0-100, you have that your guess is within 3 percentage points of the statistical process.

Figure 3: Interface, part 3 of 3

WELCOME

- Welcome and thank you for participating in today's experiment.
- Please place all your personal belongings away so that we can have your complete attention.
- Please do not socialize or talk during the experiment.
- Please use the computers as instructed. Please do not attempt to browse the web or use programs unrelated to the experiment.
- You will be paid in private and with Venmo or Zelle at the end of the experiment.
- **The amount that you ultimately earn in the experiment depends on your decisions and chance.**

Figure 4: Slide 1 of 23

TODAY'S EXPERIMENT IS ABOUT EVALUATING THE SUCCESS OF A PROJECT

- In this experiment, you are playing the role of a manager evaluating the chances of a selected project being a success.
- There are a total of 22 tasks with 4 different parts each, comprising of 2 guesses and 2 choices in the following order
 1. Guess
 2. Choice
 3. Guess
 4. Choice
- For every part, one of the tasks will be randomly selected for your payment. For example, you can be paid for your response in task 6 part 1, task 20 part 2 and so on.

Figure 5: Slide 2 of 23

1ST GUESS

- In every task, there are 100 projects, some of these projects are successes and others are failures.
- The proportion of success and failures will vary across tasks.
- The 100 projects are arranged on a 10 by 10 grid.
- Each project is represented by a square on the grid.
- The color of the square determines if the project is a success or a failure.
- White Square = Success
- Black Square = Failure
- For the first 11 tasks, the grid will be flashed on the screen for 0.25 seconds.

Figure 6: Slide 3 of 23

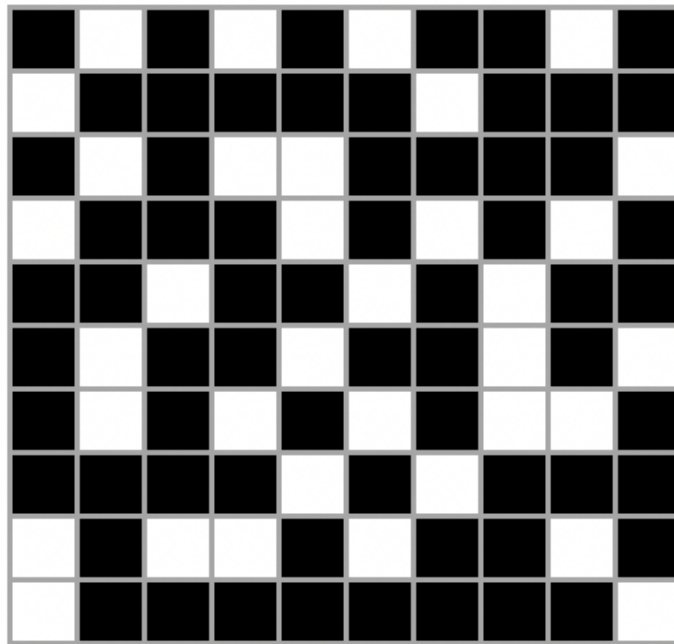


Figure 7: Slide 4 of 23

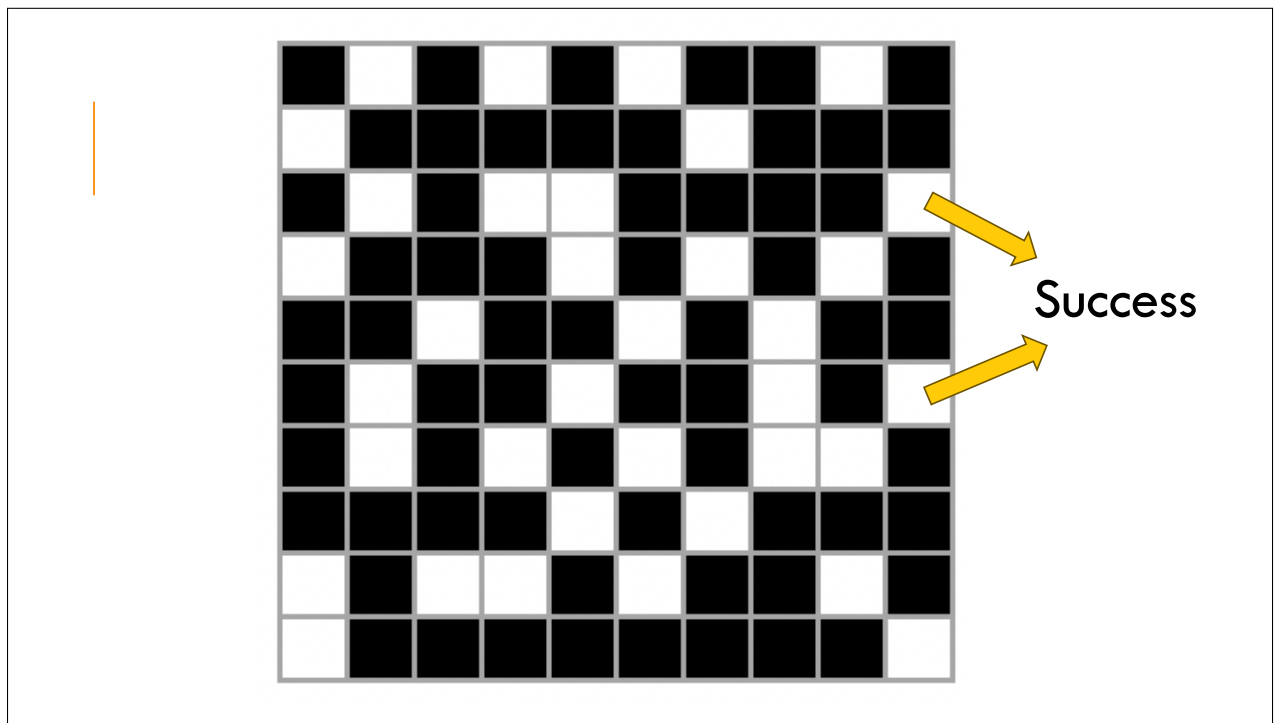


Figure 8: Slide 5 of 23

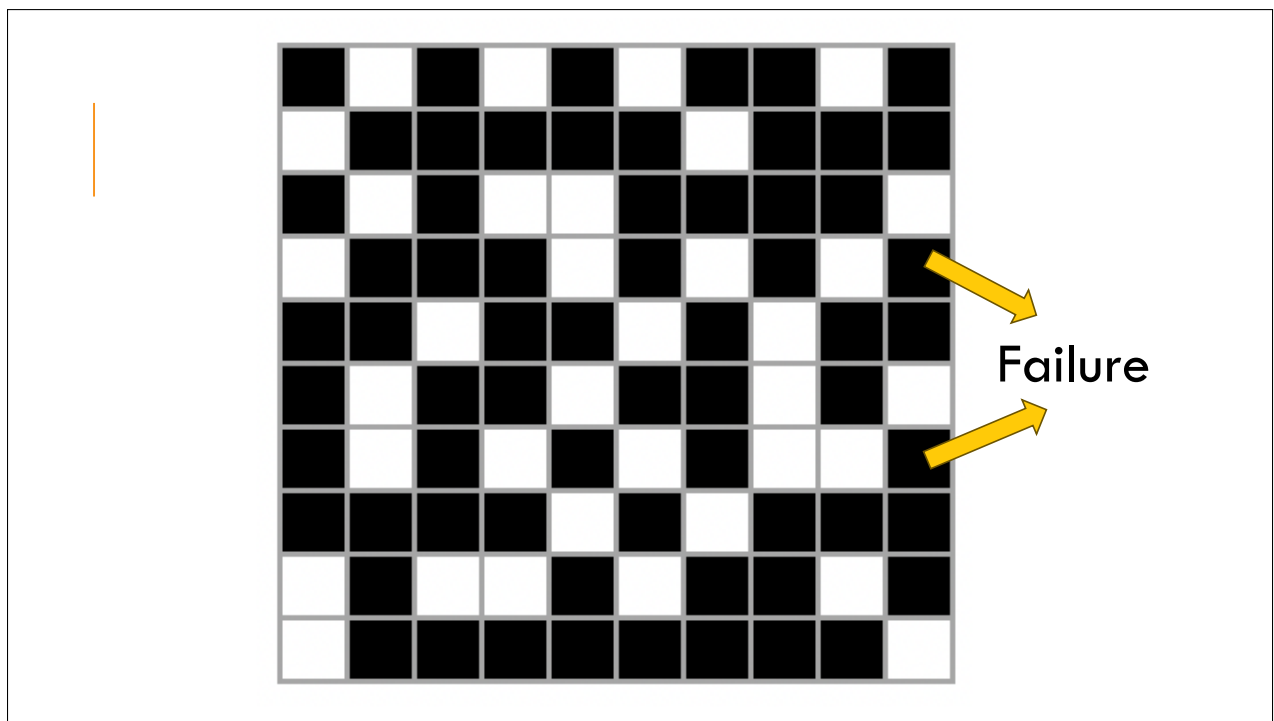


Figure 9: Slide 6 of 23

1ST GUESS

- One of the 100 projects from that grid is randomly selected for you to evaluate.
- With all projects having an **equal** chance of being selected.
- You will not know which specific project is selected.
- You will be asked to report the chance of the selected project being a success.

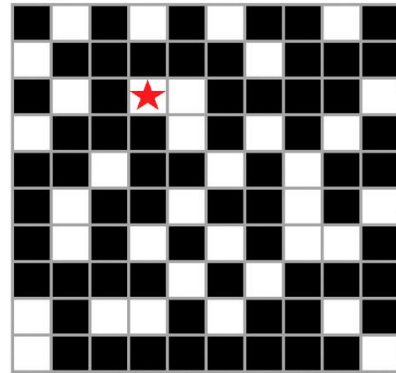


Figure 10: Slide 7 of 23

1ST GUESS

- For example, if you think there are 40 white squares
- This would mean 40% of the projects are successful (since there are 100 squares).
- There is a 40% chance that the randomly selected project is a success.
- You are not paid for whether the project is a success or a failure, but for how accurate you guess the chance of the selected project being a success.

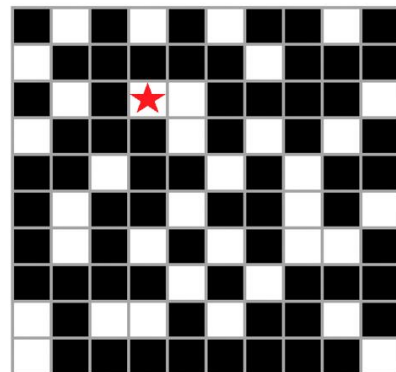


Figure 11: Slide 8 of 23

1ST GUESS

- For this part of each task, your earnings will be determined in the following manner:
- If your guess is **no more than 3 percentage points** from the actual value, you will **earn a bonus of \$3** for this part.
- For example, if the actual value is 50%, you will earn a bonus of \$3 if your response is between 47% and 53%, inclusive.

Figure 12: Slide 9 of 23

2ND GUESS

- To further aid your assessment on whether the selected project is a success or a failure:
- We will provide you with a computer test that has a reliability of 80% [60%].
- This test is not perfect. In 4 [3] out of 5 times, the computer will correctly predict whether the project is a success or a failure

Figure 13: Slide 10 of 23

2ND GUESS

- If the selected project is a Success, the test result will be Positive with 80% [60%] chance (four [three] times out of five) and the test result will be Negative with 20% [40%] chance (one [two] time out of five).
- If the selected project is a Failure, the test result will be Positive with 20% [40%] chance (one [two] time out of five) and the test result will be Negative with 80% [60%] chance (four [three] times out of five).
- You will be asked to report the chance of the selected project being a success after seeing the test result.

Figure 14: Slide 11 of 23

2ND GUESS

- For this part of each task, your earnings will be determined in the following manner:
- If your guess is **no more than 3 percentage points** from the value computed from a statistical process (we can show you this process after the experiment if you like), you will **earn a bonus of \$3** for this part.
- For example, if the value computed from the statistical process is 50%, you will earn a bonus of \$3 if your response is between 47% and 53%, inclusive.

Figure 15: Slide 12 of 23

CHOICES

- We will now explain how the choices that follow each guess works.
- After you have submitted your guess, you will have the opportunity to win another \$3 with the following choice you make:
- You will state your level of confidence (in percentage) that your guess is within 3 percentage points of the actual value or the statistical process.

Figure 16: Slide 13 of 23

CHOICES

- We are going to ask you the following list of questions:

Qn #		Option A		Option B
1	Would you rather have:	Stick to your guess	or	0% chance of \$3
2	Would you rather have:	Stick to your guess	or	1% chance of \$3
3	Would you rather have:	Stick to your guess	or	2% chance of \$3
⋮	⋮	⋮	⋮	⋮
100	Would you rather have:	Stick to your guess	or	99% chance of \$3
101	Would you rather have:	Stick to your guess	or	100% chance of \$3

Figure 17: Slide 14 of 23

CHOICES

- If you choose Option A we will pay you **\$3** for the choice if **your guess is within 3 percentage points** of the actual value/statistical process.
- If you choose Option B, you will be given a **lottery** that **pays \$3** with the **stated chance** and nothing otherwise.

Qn #		Option A		Option B
1	Would you rather have:	Stick to your guess	or	0% chance of \$3
2	Would you rather have:	Stick to your guess	or	1% chance of \$3
3	Would you rather have:	Stick to your guess	or	2% chance of \$3
⋮	⋮	⋮	⋮	⋮
100	Would you rather have:	Stick to your guess	or	99% chance of \$3
101	Would you rather have:	Stick to your guess	or	100% chance of \$3

Figure 18: Slide 15 of 23

CHOICES

- After you answer all 101 questions, we will randomly pick one question and pay you the option you chose on that one question. Each question is equally likely to be chosen for payment.

Qn #		Option A		Option B
1	Would you rather have:	Stick to your guess	or	0% chance of \$3
2	Would you rather have:	Stick to your guess	or	1% chance of \$3
3	Would you rather have:	Stick to your guess	or	2% chance of \$3
⋮	⋮	⋮	⋮	⋮
100	Would you rather have:	Stick to your guess	or	99% chance of \$3
101	Would you rather have:	Stick to your guess	or	100% chance of \$3

Figure 19: Slide 16 of 23

CHOICES

- We suspect that you may choose Option A in at least the first few questions, but at some point will switch to choosing Option B.
- So, to save time, just tell us at which point you'd switch. We can then 'fill out' your answers to all 101 questions based on your switch point

Qn #		Option A		Option B
1	Would you rather have:	Stick to your guess	or	0% chance of \$3
2	Would you rather have:	Stick to your guess	or	1% chance of \$3
3	Would you rather have:	Stick to your guess	or	2% chance of \$3
⋮	⋮	⋮	⋮	⋮
100	Would you rather have:	Stick to your guess	or	99% chance of \$3
101	Would you rather have:	Stick to your guess	or	100% chance of \$3

Figure 20: Slide 17 of 23

CHOICES

- Suppose your switch point is 75
- We will choose option A for all the rows before 76
- We will choose option B for row 76 and any rows that come after

Qn #		Option A		Option B
1	Would you rather have:	Stick to your guess	or	0% chance of \$3
2	Would you rather have:	Stick to your guess	or	1% chance of \$3
⋮	⋮	⋮	⋮	⋮
76	Would you rather have:	Stick to your guess	or	75% chance of \$3
⋮	⋮	⋮	⋮	⋮
101	Would you rather have:	Stick to your guess	or	100% chance of \$3

Figure 21: Slide 18 of 23

CHOICES

- This switch point is your level of confidence that your guess is within 3 percentage points of the actual value.

Qn #		Option A		Option B
1	Would you rather have:	Stick to your guess	or	0% chance of \$3
2	Would you rather have:	Stick to your guess	or	1% chance of \$3
⋮	⋮	⋮	⋮	⋮
76	Would you rather have:	Stick to your guess	or	75% chance of \$3
⋮	⋮	⋮	⋮	⋮
101	Would you rather have:	Stick to your guess	or	100% chance of \$3

Figure 22: Slide 19 of 23

CHOICES

- If you are 75% confident that your guess is within 3 percentage points of the actual value, you should only be willing to accept lotteries that pay \$3 at least 75% of the time.
- Reporting you are 75% confident will ensure you only get lotteries that pays \$3 at least 75% of the time.

Qn #		Option A		Option B
1	Would you rather have:	Stick to your guess	or	0% chance of \$3
2	Would you rather have:	Stick to your guess	or	1% chance of \$3
⋮	⋮	⋮	⋮	⋮
76	Would you rather have:	Stick to your guess	or	75% chance of \$3
⋮	⋮	⋮	⋮	⋮
101	Would you rather have:	Stick to your guess	or	100% chance of \$3

Figure 23: Slide 20 of 23

SUMMARY

- In summary, each task has 4 decisions that you will have to make in the following order:
 1. Guess the chance that the selected project is a success after seeing the grid
 2. State your confidence level for the earlier guess.
 3. Guess the chance that the selected project is a success after seeing a test result
 4. State your confidence level for the earlier guess.
- There are a total of 22 tasks in this experiment with 4 different parts each.
- For every part, one of the tasks will be randomly selected for your payment. For example, you can be paid for your response in task 6 part 1, task 20 part 2 and so on.

Figure 24: Slide 21 of 23

SUMMARY

- For the first 11 tasks, the grid will be flashed on your screen for 0.25 seconds
- We will provide more information about the last 11 tasks later in the experiment

Figure 25: Slide 22 of 23

EARNINGS FROM EXPERIMENT



- You have already earned \$7 for showing up on time.
- You will also receive earnings (in \$) depending on the responses you have provided in this experiment.

If there are no further questions, we will begin the experiment!

Figure 26: Slide 23 of 23

Department of Economics, University of Alberta Working Paper Series

2024-09: Merger Review under Asymmetric Information – Langinier, C. , Ray Chaudhuri, A.
2024-08: Environmental Regulations and Green Innovation: The Role of Trade and Technology Transfer – Langinier, C. , Martinez-Zarzoso, I., Ray Chaudhuri, A.
2024-07: Green Patents in an Oligopolistic Market with Green Consumers– Langinier, C. , Ray Chaudhuri, A.
2024-06: The Effect of Patent Continuation on the Patenting Process – Langinier, C.
2024-05: Applying Indigenous Approaches to Economics Instruction – Chavis, L., Wheeler, L.
2024-04: Electric Vehicles and the Energy Transition: Unintended Consequences of a Common Retail Rate Design – Bailey, M., Brown, D. , Myers, E., Shaffer, B., Wolak, F.
2024-03: Electricity Market Design with Increasing Renewable Generation: Lessons From Alberta – Brown, D. , Olmstead, D., Shaffer, B.
2024-02: Evaluating the Role of Information Disclosure on Bidding Behavior in Wholesale Electricity Markets – Brown, D. , Cajueiro, D., Eckert, A. , Silveira, D.
2024-01: Medical Ethics and Physician Motivations – Andrews, B.
2023-14: Employers’ Demand for Personality Traits and Provision of Incentives – Brencic, V. , McGee, A.
2023-13: Demand for Personality Traits, Tasks, and Sorting – Brencic, V. , McGee, A.
2023-12: Whoever You Want Me to Be: Personality and Incentives – McGee, A. , McGee, P.
2023-11: Gender Differences in Reservation Wages in Search Experiments – McGee, A. , McGee, P.
2023-10: Designing Incentive Regulation in the Electricity Sector – Brown, D. , Sappington, D.
2023-09: Economic Evaluation under Ambiguity and Structural Uncertainties – Andrews, B.
2023-08: Show Me the Money! Incentives and Nudges to Shift Electric Vehicle Charge Timing – Bailey, M., Brown, D. , Shaffer, B., Wolak, F.
2023-07: Screening for Collusion in Wholesale Electricity Markets: A Review of the Literature – Brown, D. , Eckert, A. , Silveira, D.
2023-06: Information and Transparency: Using Machine Learning to Detect Communication – Brown, D. , Cajueiro, D., Eckert, A. , Silveira, D.
2023-05: The Value of Electricity Reliability: Evidence from Battery Adoption – Brown, D. , Muehlenbachs, L.
2023-04: Expectation-Driven Boom-Bust Cycles – Brianti, M. , Cormun, V.
2023-03: Fat Tailed DSGE Models: A Survey and New Results – Dave, C. , Sorge, M.
2023-02: Evaluating the Impact of Divestitures on Competition: Evidence from Alberta's Wholesale Electricity Market – Brown, D. , Eckert, A. , Shaffer, B.
2023-01: Gender Specific Distortions, Entrepreneurship and Misallocation – Ranasinghe, A.
2022-12: The Impact of Wholesale Price Caps on Forward Contracting – Brown, D. , Sappington, D.
2022-11: Strategic Interaction Between Wholesale and Ancillary Service Markets – Brown, D. , Eckert, A. , Silveira, D.